



Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Application en amélioration des plantes

Marion Naveau

► To cite this version:

Marion Naveau. Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Application en amélioration des plantes. Statistiques [math.ST]. Université Paris-Saclay, 2024. Français. NNT : 2024UPASM031 . tel-04717855

HAL Id: tel-04717855

<https://theses.hal.science/tel-04717855v1>

Submitted on 2 Oct 2024

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.

Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Application en amélioration des plantes.

*High-dimensional variable selection procedures in non-linear
mixed effects models. Application in plant breeding.*

Thèse de doctorat de l'université Paris-Saclay

École doctorale de Mathématiques Hadamard (EDMH) n°574

Spécialité de doctorat: Mathématiques appliquées

Graduate School : Mathématiques. Référent : Faculté des sciences d'Orsay

Thèse préparée dans les unités de recherche MIA Paris-Saclay (Université Paris-Saclay, AgroParisTech, INRAE) et MaIAGE (Université Paris-Saclay, INRAE),
sous la direction de **Maud DELATTRE**, chargée de recherche,
et le co-encadrement de **Laure SANSONNET**, maîtresse de conférences.

Thèse soutenue à Paris-Saclay, le 27 septembre 2024, par

Marion NAVEAU

Composition du jury

Membres du jury avec voix délibérative

Pierre BARBILLON Professeur, Université Paris-Saclay, AgroParisTech, INRAE	Président
David CAUSEUR Professeur, Institut Agro, Université Rennes, CNRS	Rapporteur & Examineur
Robin RYDER Senior Lecturer, HDR, Imperial College London	Rapporteur & Examineur
Ismaël CASTILLO Professeur, Sorbonne Université	Examineur
Mélanie PRAGUE Chargée de recherche, Inria Bordeaux, Inserm Bordeaux Population Health, Université de Bordeaux	Examinatrice
Andrea RAU Directrice de recherche, Université Paris-Saclay, INRAE, AgroParisTech	Examinatrice

Titre: Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Application en amélioration des plantes.

Mots clés: Modèles non-linéaires à effets mixtes, Données de grande dimension, Sélection de variables bayésienne, Algorithme SAEM, Priors *spike-and-slab*, Contraction a posteriori.

Résumé: Les modèles à effets mixtes analysent des observations collectées de façon répétée sur plusieurs individus, attribuant la variabilité à différentes sources (intra-individuelle, inter-individuelle, résiduelle). Prendre en compte cette variabilité est essentiel pour caractériser sans biais les mécanismes biologiques sous-jacents. Ces modèles utilisent des covariables et des effets aléatoires pour décrire la variabilité entre individus : les covariables décrivent les différences dues à des caractéristiques observées, tandis que les effets aléatoires représentent la variabilité non attribuable aux covariables mesurées. Dans un contexte de grande dimension, où le nombre de covariables dépasse celui des individus, identifier les covariables influentes est difficile, car la sélection porte sur des variables latentes du modèle. De nombreuses procédures ont été mises au point pour les modèles linéaires à effets mixtes, mais les contributions pour les modèles non-linéaires sont rares et manquent de fondements théoriques. Cette thèse vise à développer une procédure de sélection de covariables en grande dimen-

sion pour les modèles non-linéaires à effets mixtes, en étudiant leurs implémentations pratiques et leurs propriétés théoriques. Cette procédure est basée sur l'utilisation d'un prior *spike-and-slab* gaussien et de l'algorithme SAEM (*Stochastic Approximation of Expectation Maximisation Algorithm*). Des taux de contraction a posteriori autour des vraies valeurs des paramètres dans un modèle non-linéaire à effets mixtes sous prior *spike-and-slab* discret ont été obtenus, comparables à ceux observés dans des modèles linéaires. Les travaux conduits dans cette thèse sont motivés par des questions appliquées en amélioration des plantes, où ces modèles décrivent le développement des plantes en fonction de leurs génotypes et des conditions environnementales. Les covariables considérées sont généralement nombreuses puisque les variétés sont caractérisées par des milliers de marqueurs génétiques, dont la plupart n'ont aucun effet sur certains traits phénotypiques. La méthode statistique développée dans la thèse est appliquée à un jeu de données réel relatif à cette application.

Title: High-dimensional variable selection procedures in nonlinear mixed effects models. Application in plant breeding.

Keywords: Nonlinear mixed effects models, High dimensional data, Bayesian variable selection, SAEM Algorithm, Spike-and-slab priors, Posterior contraction.

Abstract: Mixed-effects models analyze observations collected repeatedly from several individuals, attributing variability to different sources (intra-individual, inter-individual, residual). Accounting for this variability is essential to characterize the underlying biological mechanisms without bias. These models use covariates and random effects to describe variability among individuals: covariates explain differences due to observed characteristics, while random effects represent the variability not attributable to measured covariates. In high-dimensional context, where the number of covariates exceeds the number of individuals, identifying influential covariates is challenging, as selection focuses on latent variables in the model. Many procedures have been developed for linear mixed-effects models, but contributions for non-linear models are rare and lack theoretical foundations. This thesis aims to develop a high-dimensional covariate selection proce-

duce for non-linear mixed-effects models by studying their practical implementations and theoretical properties. This procedure is based on the use of a gaussian spike-and-slab prior and the SAEM algorithm (*Stochastic Approximation of Expectation Maximisation Algorithm*). Posterior contraction rates around true parameter values in a non-linear mixed-effects model under a discrete spike-and-slab prior have been obtained, comparable to those observed in linear models. The work in this thesis is motivated by practical questions in plant breeding, where these models describe plant development as a function of their genotypes and environmental conditions. The considered covariates are generally numerous since varieties are characterized by thousands of genetic markers, most of which have no effect on certain phenotypic traits. The statistical method developed in the thesis is applied to a real dataset related to this application.

Thèse préparée sous la direction de **Maud DELATTRE** (Université Paris-Saclay, INRAE, MaIAGE) et le co-encadrement de **Laure SANSONNET** (Université Paris-Saclay, AgroParis-Tech, INRAE, MIA Paris-Saclay).



Mathématique et Informatique Appliquées,
22 place de l'Agronomie,
91120 Palaiseau.



Mathématiques et Informatique Appliquées du
Génôme à l'Environnement,
Bât. 210, Domaine de Vilvert,
78352 Jouy-en-Josas.



Institut National de Recherche pour l'Agriculture,
l'Alimentation et l'Environnement



AgroParisTech, Université Paris-Saclay

Thèse financée par le projet ANR Stat4Plant ANR-20-CE45-0012, et soutenue par la FMJH.



Agence Nationale de la Recherche



Méthodes statistiques pour caractériser les in-
teractions entre la plante et son environnement



Fondation Mathématique Jacques Hadamard

À ma marraine.



Remerciements

Ce travail de recherche n'aurait pu voir le jour sans le soutien, les conseils et l'encouragement de nombreuses personnes. Je tiens à exprimer ma plus profonde gratitude à celles et ceux qui, de près ou de loin, ont contribué à l'élaboration de cette thèse.

Je tiens tout d'abord à remercier particulièrement mes directrices de thèse, Maud et Laure, pour leur accompagnement tout au long de ce travail. Votre expertise, vos précieux conseils et votre bienveillance ont été une véritable source d'inspiration et de motivation pour moi. Votre rigueur scientifique et votre grande disponibilité ont largement contribué à l'avancement et au succès de ce travail. Votre capacité à partager vos connaissances tout en respectant mon autonomie a été essentielle dans l'évolution de mes réflexions. Merci pour la confiance que vous m'avez accordée, notamment concernant le télétravail depuis Saint-Malo, ainsi que pour votre soutien sans faille dans les moments de doute, toujours accompagné de conseils éclairés et bienveillants, tant sur le plan scientifique que personnel. Grâce à vous, j'ai appris à gagner en confiance en moi (même s'il reste encore du chemin à parcourir !). Dès mon stage de M2, j'ai su que je ne pourrais pas rêver de meilleures directrices de thèse !

Ensuite j'aimerais remercier Guillaume Kon Kam King. Merci pour ta contribution à ma thèse, ton expertise m'a toujours impressionnée ! Tes réflexions et suggestions scientifiques ont toujours été pertinentes et ont grandement accéléré la maturation de mes recherches. Merci pour ton enthousiasme et ton aide précieuse dans ce travail. Ta présence dans ce jury était une évidence, merci d'avoir accepté notre invitation.

Je souhaite exprimer ma profonde gratitude à l'ensemble des membres du jury pour avoir accepté de donner de leur temps pour évaluer mon travail. Un merci tout particulier à David Causeur et Robin Ryder pour l'honneur qu'ils m'ont fait en rapportant cette thèse. Je suis très reconnaissante pour vos relectures attentives et approfondies, ainsi que pour vos compliments encourageants. Je tiens également à adresser mes sincères remerciements à Pierre Barbillon, Ismaël Castillo, Mélanie Prague et Andrea Rau pour leur implication et leur engagement dans l'examen de cette thèse. En particulier, je remercie Ismaël pour nos échanges enrichissants concernant la partie théorique de mon travail de thèse. Nos discussions ont été cruciales pour surmonter les impasses et ont considérablement éclairé ma compréhension des enjeux et difficultés de ce que je cherchais à démontrer.

Je souhaite également exprimer ma gratitude aux membres de mon comité de suivi de thèse : Julyan Arbel, Madalina Olteanu et Angelina Roche. Merci pour vos conseils, vos questions pertinentes et vos encouragements lors des comités. Ces réunions étaient pour moi l'occasion précieuse d'avoir un regard extérieur sur mes travaux, et vos retours m'ont permis de prendre confiance en moi et mon travail.

Je tiens également à remercier les membres du projet ANR *Stat4Plant* pour les nombreuses discussions enrichissantes à l'interface entre statistiques et biologie. Ces échanges ont renforcé ma conviction de poursuivre mes recherches à cette intersection. Je souhaite remercier tout particulièrement Estelle. Ton énergie débordante, ton soutien infaillible et tes conseils, tant sur le plan scientifique que personnel, ont été d'une grande aide. Merci de m'avoir pris sous ton aile, depuis les réunions "Point technique" durant mon stage de M2 jusqu'à tes conseils avisés dans mes recherches de Post-Doc. Et bien sûr, merci aussi pour les délicieux gâteaux d'anniversaire ! Je remercie aussi Renaud pour nos discussions sur les données biologiques qui m'ont permis de mieux comprendre les enjeux applicatifs de ces travaux de thèse. Merci également à Charlotte, Sarah et Jean-Benoist pour leurs suggestions et conseils sur mes travaux durant les réunions du projet ANR.

Un grand merci aussi à Veronica Beswick, ma mentore, pour tes conseils avisés, ton enthousiasme et d'avoir pris le temps de me partager ton expérience.

Je tiens également à remercier chaleureusement mes deux unités, MIA Paris-Saclay et MaIAGE, pour leur accueil et les nombreux moments conviviaux et scientifiques que nous avons partagés.

Bien que j'aie passé la majeure partie de mon temps à MaIAGE, j'ai toujours retrouvé une ambiance chaleureuse lors de mes passages à MIA ! Je tiens en particulier à remercier la belle équipe de doctorants : Marina, Mary, Armand, Bastien, Jérémy, Florian, Emré, Annaïg, Caroline, Barbara, Jeanne, Hayato, François. Je souhaite une belle aventure aux nouveaux arrivants également, mais je n'ai aucun doute, vous serez bien entourés. Un merci tout particulier à mes co-bureaux pour les discussions sur l'après-thèse, les parties de fléchettes, les conseils sur ChatGPT, et les éclats de rire en fin de journée, quand plus personne n'a envie de travailler ! J'en profite également pour remercier José pour son accueil dans son bureau dès mon stage de M2 quand on était encore sur Paris, car oui je pourrais dire que j'ai vécu ce déménagement de l'Agro sur le plateau ! Merci également à tous les permanents pour les repas et pauses café, toujours riches en échanges instructifs. J'en profite également pour remercier toute l'équipe de Rochebrune. Nous avons vraiment passé de bons moments ensemble, et j'espère avoir l'occasion de vous y rejoindre à nouveau dans le futur !

J'ai commencé mon stage de M2 dans les locaux de MaIAGE à Jouy, au bâtiment 210. Je tiens à remercier l'équipe Dynenvie pour leur accueil et soutien tout au long de cette thèse. Un merci tout particulier à Gildas pour notre séjour à Varsovie pour l'EMS avec la team des doctorants de stats, on aura bien rigolé ! Merci Laurent pour les discussions instructives à la pause café, et tes nombreux conseils toujours avisés. Je souhaite également remercier les post-doctorants Jeanne, Thibault, Auguste et Aurélien pour les encouragements et les discussions sur les concours et l'après-thèse notamment. Merci également aux nombreux stagiaires qui sont passés par là, je pense notamment à ma team Piou-Piou, ils se reconnaîtront ! Merci en particulier à la première étudiante que j'ai encadré, Viviana : merci pour ta confiance, j'ai été ravie de t'aiguiller dans ces travaux et je te souhaite le meilleur pour la suite. Et finalement les doctorants... Merci Henri pour ton soutien constant et tes mots toujours justes. Amandine, merci pour ta gentillesse. Marie, je te remercie pour ton accueil chaleureux à MaIAGE, ta bonne humeur contagieuse et ton soutien indéfectible. Madeleine, Tom, je suis tellement heureuse d'avoir partagé toutes ces étapes de la thèse à vos côtés. Merci à vous pour les moments de rire, de doute partagés et les efforts collectifs de re-motivation ! Un grand merci à Andréa, Antoine, Katia et Arthur pour l'énergie positive et la bonne humeur que vous apportez au labo. Merci en particulier à Antoine pour les debugging Latex ! Merci également à Julie, Anastasia et Eleonora pour tous les moments agréables passés ensemble au labo. Et bien sûr, je n'oublie pas les merveilleux moments passés en dehors du labo, que ce soit au Massy Pal, à Orsay, à Paris, en Bretagne ou lors des conférences à travers la France ! Je suis sincèrement triste de vous quitter, vous avez placé la barre très haut pour mon prochain labo... Vous allez beaucoup me manquer. J'espère que nous continuerons à organiser nos week-ends bretons et normands régulièrement ! Enfin, courage aux doctorants restants, je serai toujours là pour répondre à vos questions, même à distance (**#respo_admin**) !

Je tiens également à remercier mes amis de toujours : Axelle, Mathias, Julien, Joséphine. Merci pour votre soutien indéfectible, vos encouragements constants, et surtout pour les moments précieux passés ensemble, qui m'ont permis de m'évader un peu. Ça y est, je suis enfin arrivée au fameux bac+8000 ! Un grand merci également à mes amis de la fac d'Orsay : Claire, Taher, Léana, Marion, Donatien, Jeanne et Diane. Merci pour les soirées "filles", les éclats de rire partagés, et les échanges sur nos expériences d'enseignants. Un merci tout particulier à Claire pour nos longues discussions et tes paroles toujours justes face aux défis de la thèse, et de la vie en général d'ailleurs. Enfin, merci à Perrine pour ton écoute attentive, tes conseils précieux, et ton aide inestimable dans la gestion des complexités administratives de l'agrégation !

Un immense merci à toute ma famille ! Vous avez toujours cru en moi depuis le tout début. Maman, Papa, je sais que vous n'avez jamais douté de moi, et vous rendre fiers est, pour moi, la plus grande des réussites. Papa j'espère quand même que tu n'as pas apporté de banderole le jour J ! Laëtitia, Lucas, merci pour votre soutien indéfectible et les moments de rigolade en famille. Julie, Romain, merci pour vos encouragements constants et pour m'avoir honorée en me confiant le rôle de marraine de votre princesse. Un grand merci également à ma belle-famille pour leur soutien sans faille et les bons moments passés à Saint-Benoit pour déconnecter. Tonton je ne t'oublie pas, il est temps de faire tes comptes, tu as un voyage en Corse à organiser du coup ! Et pour info, je suis également ouverte à l'idée d'un road-trip aux États-Unis ou en Islande si jamais ! Ces quelques mots sont difficiles à rédiger étant donné les circonstances vous vous en doutez, mais je suis sûre que tata aurait été fière de tout ça. Ces moments difficiles renforcent encore plus nos liens et nous allons avancer encore plus soudés !

Et pour finir, le meilleur pour la fin comme on dit, je tiens à te remercier Thomas. Merci pour ta grannnnnde patience face à mes nombreux moments de doute et tes encouragements constants. Ta présence à mes côtés, ton soutien moral et ta capacité à me faire sourire, même dans les périodes les plus difficiles, ont été des piliers essentiels pour moi. Je suis profondément reconnaissante pour tout ce que tu as fait et pour la force que tu m'as donnée pour mener à bien ce projet. Merci pour les voyages partout dans le monde permettant de couper du quotidien. Je suis fière de t'avoir à mes côtés et impatiente de découvrir ensemble ce que l'avenir nous réserve.

Préambule

Cette thèse est structurée en trois parties principales.

La première partie sert d'introduction aux concepts fondamentaux nécessaires pour comprendre et contextualiser les contributions de la thèse. Elle est divisée en quatre chapitres. Le chapitre 1 expose le contexte et les motivations biologiques de ces recherches. Le chapitre 2 présente le modèle statistique central dans cette thèse, le modèle non-linéaire à effets mixtes, ainsi qu'un état de l'art des méthodes d'inférence disponibles. Le chapitre 3 aborde les défis de la sélection de variables en grande dimension, où le nombre de variables excède largement le nombre d'observations, et discute des solutions fréquentistes et bayésiennes. Le chapitre 4 offre un résumé des contributions de cette thèse, en mettant en avant leur originalité et les défis rencontrés.

La deuxième partie se concentre sur les travaux et résultats de la thèse, incluant des développements méthodologiques, algorithmiques, théoriques et une application sur des données réelles. Plus précisément, le chapitre 5 correspond à un article méthodologique publié dans *Statistics and Computing*, tandis que le chapitre 6 présente une pré-publication sur des avancées théoriques, accessible en ligne. Ces deux chapitres sont rédigés en anglais.

La dernière partie présente des travaux en cours et explore différentes perspectives. Le chapitre 7 introduit la notion de métamodélisation et son intérêt dans le domaine d'application de cette thèse. Le chapitre 8 propose diverses perspectives, tant méthodologiques que théoriques.

Table des matières

I	Introduction	4
1	Contexte et motivations biologiques	6
1.1	Projet ANR <i>Stat4Plant</i>	6
1.2	Sélection variétale	7
1.3	Modélisation pour la sélection assistée par marqueurs	8
2	Modèles non-linéaires à effets mixtes et méthodes d'inférence	10
2.1	Description du modèle	11
2.1.1	Spécificités de l'analyse de données répétées	11
2.1.2	Modèle statistique	12
2.1.3	Vraisemblance	13
2.2	Méthodes d'inférence dans les modèles non-linéaires à effets mixtes	13
2.2.1	Méthodes approchées	14
2.2.2	Méthodes exactes	16
2.3	Algorithme <i>Expectation-Maximisation</i> et ses extensions	18
2.3.1	Algorithme EM	18
2.3.2	Algorithme MCEM	19
2.3.3	Algorithme SAEM : une version stochastique de l'algorithme EM	20
2.3.4	Algorithme MCMC-SAEM : couplage avec une procédure MCMC	22
3	Sélection de variables en grande dimension	26
3.1	Sélection de variables en grande dimension	27
3.1.1	Méthodes pénalisées fréquentistes	28
3.1.2	Priors bayésiens induisant la parcimonie	30
3.2	Analyse fréquentiste des distributions a posteriori	35
3.2.1	Paradigme fréquentiste pour la convergence des distributions a posteriori	35
3.2.2	Taux de contraction de la posterior	36
3.2.3	Consistance en sélection	41
4	Positionnement et contributions de la thèse	44
4.1	Procédures de sélection de variables en grande dimension dans les NLMEM	45
4.1.1	État de l'art des méthodes de sélection de variables en grande dimension dans les modèles à effets mixtes	45
4.1.2	Contribution du Chapitre 5 : Méthode de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes	47
4.2	Taux de contraction a posteriori en grande dimension sous prior <i>spike-and-slab</i>	49
4.2.1	État de l'art des propriétés asymptotiques fréquentistes des priors <i>spike-and-slab</i>	49
4.2.2	Contribution du Chapitre 6 : Taux de contraction a posteriori en grande dimension dans un NLMEM	51
II	Travaux et résultats	56
5	Sélection de covariables bayésienne en grande dimension dans les modèles non-linéaires à effets mixtes à l'aide de l'algorithme SAEM	58
5.1	Introduction	59
5.2	Model description	61
5.2.1	Non-linear mixed-effects model and notations	61
5.2.2	Prior specification	63

5.3	Maximum <i>a posteriori</i> inference and thresholding	64
5.3.1	General description of SAEM algorithm	65
5.3.2	Central decomposition of the Q quantity in spike-and-slab non-linear mixed-effects models	67
5.3.3	Estimator thresholding	69
5.4	Covariate selection procedure	69
5.5	Numerical experiments	71
5.5.1	Comparison with a two-step approach	71
5.5.2	Impacts of the different parameters and collinearity between covariates	73
5.5.3	Comparison with an MCMC implementation	78
5.6	Application to plant senescence genetic marker identification	81
5.6.1	Data description and pre-processing	81
5.6.2	Modelling	82
5.6.3	Results and discussion	82
5.7	Conclusion and perspectives	84
5.8	Appendices	85
5.8.1	Algorithms: synthesis and implementation details	85
5.8.2	Results for scenarios with correlated covariates	91
5.8.3	Further details on real data	93
5.9	Application on a detailed example	95
5.9.1	Convergence of the MCMC-SAEM algorithm	95
5.9.2	Spike-and-slab regularisation plot and model selection	96
6	Taux de contraction a posteriori dans un modèle non-linéaire à effets mixtes de grande dimension	98
6.1	Introduction	99
6.2	Model description	101
6.2.1	Non-linear marginal mixed-effects model	101
6.2.2	Prior specification	102
6.2.3	Assumptions	102
6.3	Posterior contraction results	105
6.3.1	Support size control	105
6.3.2	Posterior contraction rates	105
6.4	Proofs of main theorems	108
6.4.1	Proof of Theorem 1	108
6.4.2	Proof of Theorem 2	110
6.4.3	Proof of Theorem 3	115
6.4.4	Proof of Theorem 4	117
6.5	Appendices	118
6.5.1	Proofs of technical lemmas	118
6.5.2	Useful lemmas	126
III	Travaux en cours, conclusion et perspectives	130
7	Travaux en cours : Métamodélisation pour la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes	132
7.1	Motivations pratiques	133
7.2	Outils de métamodélisation dans les modèles non-linéaires à effets mixtes complexes	134
7.2.1	Inférence dans les modèles non-linéaires à effets mixtes complexes	134
7.2.2	Méta-modèles par krigeage gaussien	135
7.3	Couplage SAEMVS et métamodélisation	137
7.3.1	Première approche	137

7.3.2	Plan de travail	139
8	Conclusion et perspectives	142
8.1	Perspectives méthodologiques	143
8.1.1	Quantification de l'incertitude	143
8.1.2	Sélection des effets aléatoires	143
8.1.3	Sélection par groupe de variables	144
8.1.4	Modèle hétéroscédastique	145
8.1.5	Ajout d'une structure d'apparement	145
8.2	Perspectives théoriques	146
8.2.1	Consistance en sélection	146
8.2.2	Extension au modèle non-linéaire à effets mixtes général	146
8.2.3	Extension au prior <i>spike-and-slab</i> continu	147

Partie I

Introduction

Contexte et motivations biologiques

Les travaux menés dans le cadre de cette thèse se concentrent sur l'analyse de données répétées, où une quantité d'intérêt est mesurée de manière répétée sur plusieurs individus. Ce type de données est particulièrement fréquent pour l'amélioration des plantes en agriculture, où les caractéristiques phénotypiques d'intérêt sont observées chez plusieurs variétés d'une espèce végétale dans différents environnements. Plus précisément, nous nous concentrons principalement, bien que pas exclusivement, sur la sélection de marqueurs génétiques associés à des performances phénotypiques d'intérêt. Dans le contexte du changement climatique, caractériser les interactions entre la plante et son environnement devient une réelle nécessité. En effet, l'agriculture est au cœur des préoccupations, à la fois comme l'une des causes du changement climatique, mais aussi pour les profondes perturbations auxquelles elle sera confrontée et devra s'adapter. Il est crucial de noter que même au sein d'une même espèce, les individus présentent des variations génétiques, ce qui se traduit par des différences dans leur capacité à résister à des stress environnementaux tels que la sécheresse ou les insectes ravageurs par exemple. L'objectif est donc d'identifier des leviers d'action biologiques afin de sélectionner les variétés les plus adaptées.

Ce chapitre aborde les contextes et motivations biologiques de cette thèse en trois sections distinctes. Premièrement, la section 1.1 introduit le projet ANR *Stat4Plant* qui finance ces travaux, en soulignant ses objectifs et son importance. Ensuite, la section 1.2 se concentre sur la sélection variétale en agriculture à travers l'identification de marqueurs génétiques, expliquant son rôle crucial dans l'amélioration des plantes. Finalement, la section 1.3 aborde les modèles statistiques utilisés dans ce domaine, mettant en lumière l'importance des modèles à effets mixtes.

Table des matières

1.1	Projet ANR <i>Stat4Plant</i>	6
1.2	Sélection variétale	7
1.3	Modélisation pour la sélection assistée par marqueurs	8

1.1 Projet ANR *Stat4Plant* : Méthodes statistiques pour caractériser les interactions entre la plante et son environnement

Actuellement, l'agriculture est confrontée à de nouveaux défis en raison de la demande croissante en nourriture à l'échelle mondiale, conjuguée aux effets du changement climatique et à la diminution des ressources naturelles telles que l'eau et le sol. Ces défis exigent des adaptations majeures de l'agriculture pour répondre aux nouvelles conditions, en particulier sur le plan environnemental. Comprendre des concepts clés tels que la diversité génétique des plantes et leurs interactions avec l'environnement est crucial pour cette adaptation. Dans cette optique, l'utilisation de modèles prédictifs basés sur les connaissances génétiques, écophysiologiques et la

modélisation mathématique se révèle très prometteuse (Onogi et al., 2016; Rincent et al., 2017; Messina et al., 2018; Rincent et al., 2019).

Le projet ANR *Stat4Plant*, qui finance cette thèse, a pour objectif le développement de nouvelles méthodes statistiques et d'outils algorithmiques afin de modéliser et d'analyser la variabilité génétique ainsi que les interactions entre les plantes et leur environnement. Ce projet réunit des chercheurs spécialisés en modélisation, en statistiques appliquées et en biologie végétale, bénéficiant d'une vaste expérience dans des collaborations interdisciplinaires. Il est structuré autour de quatre axes de recherche principaux, qui impliquent une étroite collaboration entre les statisticiens et les biologistes, et qui sont motivés par des questions biologiques spécifiques et des données réelles.

Ma thèse s'inscrit dans le troisième volet de ce projet, qui consiste à élaborer des méthodes permettant d'identifier parmi un ensemble important de covariables celles qui ont le plus d'influence sur un trait phénotypique. Pour cela, nous utilisons des modèles non-linéaires à effets mixtes qui combinent des modèles mécanistes de développement avec des modèles génétiques intégrant une multitude de covariables. Ces modèles servent à représenter la variabilité génotypique du trait d'intérêt. De plus, nous développons des techniques de sélection de variables adaptées au contexte non-linéaire pour identifier les facteurs génétiques cruciaux qui influent sur le trait. Le contexte biologique est discuté plus en détail dans les sections 1.2 et 1.3.

1.2 Sélection variétale : l'amélioration des plantes par l'identification de marqueurs génétiques

La sélection variétale joue un rôle essentiel dans l'amélioration des plantes en sélectionnant les individus ou génotypes les plus performants pour les générations à venir. Son objectif principal est de développer des variétés plus productives, résistantes aux maladies et aux ravageurs, adaptées à différents climats et types de sols, tout en répondant aux exigences de qualité nutritionnelle. Les avancées récentes en biotechnologie, et notamment l'émergence d'outils de génomique et de génotypage à haut débit, ont radicalement transformé ce domaine de la sélection végétale. Désormais, nous pouvons utiliser des informations génétiques pour orienter la sélection des plantes et améliorer leurs caractéristiques agronomiques. Plus précisément, cela repose sur l'identification de marqueurs génétiques, des séquences d'ADN impliquées dans la variation phénotypique des caractères d'intérêt. Cette approche permet une sélection plus rapide et précise dès les premiers stades de développement des plantes, réduisant ainsi la nécessité d'essais sur le terrain et économisant à la fois du temps et des ressources. Pour ces raisons, la sélection assistée par marqueurs est donc maintenant souvent favorisée par rapport à la sélection basée sur le phénotype (Moreau et al., 1998; Van Berloo and Stam, 1999; Jannink et al., 2001; Meuwissen et al., 2001).

Pour mettre en place la sélection de variétés assistée par marqueurs, il est donc d'abord nécessaire d'identifier les marqueurs génétiques, tels que des SNP (*Single Nucleotide Polymorphism*, mutations nucléotidiques simples), associés à un caractère phénotypique d'intérêt, tels que la croissance, le rendement ou la résistance aux maladies : cette étape est connue sous le nom de cartographie d'association (Collard et al., 2005). Cela nécessite des techniques statistiques avancées, comme celles qui sont décrites plus en détail par la suite.

1.3 Modélisation pour la sélection assistée par marqueurs

En sélection végétale, les variétés les plus performantes sont souvent différentes d'un environnement à l'autre. Ces phénomènes sont appelés interactions génotype $\times$ environnement. Ainsi, la principale difficulté de la sélection assistée par marqueurs réside dans la séparation de la part génétique de la variation du phénotype d'intérêt de celle due aux variations environnementales.

Pour résoudre ce problème, les données utilisées sont généralement des mesures répétées, comprenant plusieurs observations du caractère phénotypique sur plusieurs variétés d'une espèce végétale d'intérêt dans différents environnements. Pour tenir compte des différentes sources de variabilité présentes dans ces données (intra-individuelle, inter-individuelle et résiduelle), les modèles à effets mixtes, introduits par [Laird and Ware \(1982\)](#), offrent une solution naturelle. En effet, ce type de modèle est largement utilisé pour analyser des données génotype $\times$ environnement ([Piepho, 1997](#); [Smith et al., 2005](#); [Kang et al., 2008](#); [Schulz-Streeck et al., 2013](#); [Heslot et al., 2014](#)), car il permet de décrire le développement des plantes en fonction de leurs génotypes et des conditions environnementales. Ils permettent de comprendre le rôle des interactions entre le génotype et l'environnement dans l'évolution de la plante et sont utilisés pour prédire les performances de différentes variétés dans des conditions environnementales spécifiques. La structure hiérarchique des modèles à effets mixtes et l'incorporation d'effets aléatoires confèrent une plus grande flexibilité au processus de modélisation, comme cela sera plus largement détaillé dans le chapitre suivant (voir section 2.1.2, équations (2.1)-(2.2)). De plus, dans le cadre de l'identification de marqueurs génétiques, il a été démontré ([Yu et al., 2006](#); [Malosetti et al., 2007](#); [Zhao et al., 2007](#)) que les modèles à effets mixtes obtiennent moins de faux positifs que les méthodes précédemment utilisées telles que l'analyse en composantes principales ([Patterson et al., 2006](#)).

Par ailleurs, lors des expériences, de nombreuses covariables d'intérêt sont aussi collectées pour caractériser les variétés et le climat. Ces covariables peuvent être de grande dimension, comme les données génomiques. Elles sont utilisées pour prédire des caractéristiques phénotypiques d'intérêt, grâce à des modèles de culture ([Hammer et al., 2002](#); [Yin et al., 2004](#); [Vos et al., 2007](#)). Plus précisément, ces modèles simulent le développement de la plante spécifiquement pour chaque variété dans un environnement donné, en prenant en compte les paramètres génétiques et les covariables environnementales. Une fois que les paramètres génétiques sont estimés ou mesurés pour les variétés d'intérêt, le modèle de culture peut être utilisé pour prédire la performance de chaque variété. Pour les nouvelles variétés, un modèle de prédiction génomique basé sur des marqueurs moléculaires peut être appliqué pour prédire leurs paramètres génétiques, qui peuvent à leur tour être utilisés pour exécuter le modèle de culture ([Reymond et al., 2003](#)). Par exemple, le modèle de culture dénommé SIRIUS ([Jamieson et al., 1998](#)) permet de simuler la production de biomasse du blé à partir de covariables environnementales. Ainsi, pour tenir compte des interactions génotype $\times$ environnement, plusieurs travaux ont proposé de combiner les modèles de culture et de prédiction génomique au sein d'un modèle non-linéaire à effets mixtes afin d'effectuer des prédictions ([Technow et al., 2015](#); [Cooper et al., 2016](#); [Onogi et al., 2016](#); [Rincent et al., 2017](#); [Messina et al., 2018](#)).

Cependant, les variétés sont caractérisées par des milliers de covariables (par exemple, les marqueurs génétiques), et la plupart d'entre elles n'ont pas d'effet sur certains caractères phénotypique d'intérêt. Toutefois, à notre connaissance, aucune méthode ne propose de sélectionner les covariables associées au caractère d'intérêt dans un modèle de culture non-linéaire. Il sem-

ble donc particulièrement intéressant d'améliorer les approches précédentes en introduisant la sélection de variables. De plus, il est important de noter que le nombre de marqueurs est souvent beaucoup plus élevé que le nombre d'individus dans les données. Cette situation présente plusieurs défis pour l'analyse statistique des mesures répétées, notamment la gestion de la dépendance entre les observations, la sélection de variables dans un espace de grande dimension, et la non-linéarité du modèle. Dans cette thèse, nous avons opté pour une approche basée sur les statistiques bayésiennes pour résoudre ces problématiques, en exploitant la flexibilité offerte par la construction de modèles hiérarchiques. Notre objectif principal est de développer une méthode statistique efficace dans ces modèles pour sélectionner les marqueurs génétiques associés aux traits d'intérêt. Les détails de la modélisation statistique sont exposés dans le chapitre suivant.

Modèles non-linéaires à effets mixtes et méthodes d'inférence

Dans cette perspective de sélection variétale, il est nécessaire d'avoir un modèle capable de décrire un schéma global du comportement de la population tout en tenant compte de la diversité entre les individus. C'est pour ce type de contexte que les modèles à effets mixtes ont été introduits. En effet, de part leur structure hiérarchique et l'incorporation d'effets aléatoires, ces derniers permettent de prendre en compte les différentes sources de variabilité dans les données. Plus précisément, dans cette thèse nous nous concentrons sur les modèles non-linéaires à effets mixtes, où la fonction de régression est non-linéaire en les paramètres individuels.

Dans ce chapitre nous introduisons ces modèles, les défis associés à l'estimation de leurs paramètres et quelques solutions algorithmiques. Tout d'abord, la section 2.1 présente les modèles à effets mixtes comme outils d'analyse des données répétées. Ensuite, la section 2.2 explore les différentes méthodes d'inférence dans ces modèles. Enfin, la section 2.3 approfondit la description de la méthode d'estimation la plus couramment utilisée, à savoir l'algorithme *Expectation-Maximisation* et ses extensions.

Table des matières

2.1	Description du modèle	11
2.1.1	Spécificités de l'analyse de données répétées	11
2.1.2	Modèle statistique	12
2.1.3	Vraisemblance	13
2.2	Méthodes d'inférence dans les modèles non-linéaires à effets mixtes	13
2.2.1	Méthodes approchées	14
2.2.1.1	Approche en deux temps	14
2.2.1.2	Méthodes de linéarisation	15
2.2.2	Méthodes exactes	16
2.2.2.1	Méthodes d'intégration numérique	16
2.2.2.2	Méthodes bayésiennes ou d'intégration stochastique	16
2.2.2.3	Méthodes de maximisation de la vraisemblance	17
2.3	Algorithme <i>Expectation-Maximisation</i> et ses extensions	18
2.3.1	Algorithme EM	18
2.3.2	Algorithme MCEM	19
2.3.3	Algorithme SAEM : une version stochastique de l'algorithme EM	20
2.3.3.1	Description générale de l'algorithme	20
2.3.3.2	Version simplifiée de l'algorithme pour les modèles exponentiels courbes	21

2.3.4	Algorithme MCMC-SAEM : couplage avec une procédure MCMC . . .	22
2.3.4.1	Méthode MCMC et algorithme de Metropolis-Hastings . . .	22
2.3.4.2	Couplage MCMC-SAEM	23

2.1 Description du modèle

2.1.1 Spécificités de l'analyse de données répétées

Les modèles à effets mixtes se placent dans le contexte général d'observations collectées de façon répétée, soit au fil du temps, soit dans différentes conditions, environnementales par exemple, sur plusieurs individus au sein d'une population d'intérêt.

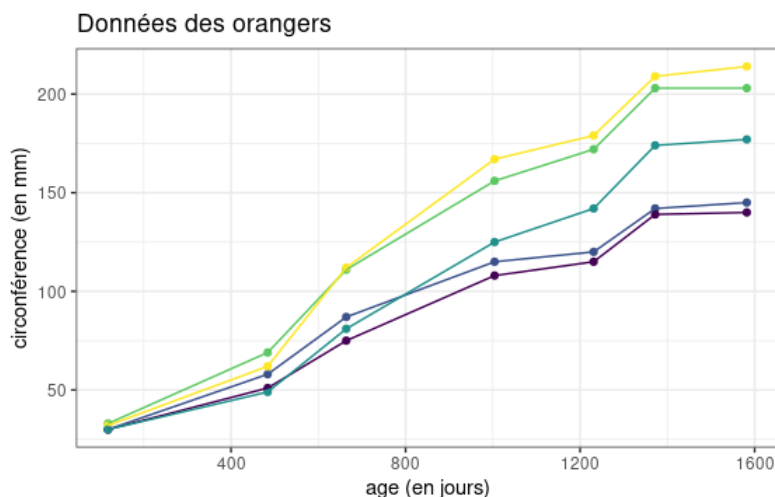


Figure 2.1: Exemple de mesures répétées de circonférences de tronc dans une population d'orangers : 5 individus sont représentés sur ce graphique, chacun par une couleur différente.

Les modèles à effets mixtes ont principalement été introduits pour modéliser les réponses d'individus ayant le même comportement global mais avec des variations individuelles. Par exemple, [Pinheiro and Bates \(2000\)](#) proposent de regarder la circonférence du tronc d'une population d'orangers au fil du temps (voir Figure 2.1). La variabilité intrinsèque aux données est alors attribuable à différentes sources (intra-individuelle, inter-individuelle, résiduelle) dont la prise en compte est essentielle pour caractériser sans biais les mécanismes biologiques à l'origine des observations. Dans ces modèles, on suppose que toutes les réponses suivent une forme fonctionnelle connue qui dépend d'effets inconnus. Ces effets peuvent être de deux types, soit fixes, c'est-à-dire qu'ils sont identiques pour tous les individus, soit aléatoires, c'est-à-dire qu'ils varient d'un individu à un autre. Les effets fixes décrivent les relations entre les covariables et le paramètre individuel d'intérêt pour une population entière, tandis que les effets aléatoires sont spécifiques à l'échantillon des données individuelles. Le terme "covariables" fait référence aux variables explicatives qui peuvent être pertinentes pour expliquer la variabilité inter-individuelle. Le formalisme est précisé dans la section suivante.

Les modèles linéaires à effets mixtes, introduits par (Laird and Ware, 1982), et leur version non-linéaire, proposée par (Lindstrom and Bates, 1990), sont utilisés dans divers domaines tels que, entre autres, la pharmacocinétique, la médecine, la sociologie ou la croissance biologique. Pour des exemples concrets illustrant leur application, le lecteur peut se référer aux illustrations fournies dans Pinheiro and Bates (2000) et Lavielle (2014). La section suivante s'appuie principalement sur ces deux références, mais le lecteur peut également consulter les ouvrages classiques suivants : Verbeke and Molenberghs (2000) et West et al. (2006) pour les modèles linéaires, et Davidian and Giltinan (1995) pour les modèles non-linéaires.

2.1.2 Modèle statistique

Les modèles à effets mixtes se décrivent selon deux niveaux décrivant l'ensemble des sources de variabilité. On note $n \in \mathbb{N}^*$ le nombre d'individus, $n_i \in \mathbb{N}^*$ le nombre d'observations de l'individu i , $p \in \mathbb{N}^*$ le nombre de covariables et $q \in \mathbb{N}^*$ le nombre de paramètres individuels.

Tout d'abord, **au niveau de l'individu**, pour décrire la variabilité intra-individuelle, on modélise les observations y_{ij} pour tout $i \in \{1, \dots, n\}$ et pour tout $j \in \{1, \dots, n_i\}$ par :

$$y_{ij} = f(\varphi_i, \psi, t_{ij}) + \varepsilon_{ij} \quad (2.1)$$

où $y_{ij} \in \mathbb{R}$ représente la réponse de l'individu i dans la condition t_{ij} (pouvant correspondre au temps ou à l'environnement), $\varphi_i \in \mathbb{R}^q$ sont les paramètres individuels de l'individu i et sont non observés, ψ est un vecteur qui désigne des paramètres de population inconnus constituant des effets fixes, et $(\varepsilon_{ij})_{i,j}$ sont des bruits gaussiens indépendants et identiquement distribués (iid) selon une loi normale $\mathcal{N}(0, \sigma^2)$, avec $\sigma^2 > 0$ la variance résiduelle des observations. La fonction f régit le comportement intra-individuel. Lorsque f est une fonction linéaire (respectivement non linéaire) en φ_i , on parle de modèle linéaire (respectivement non linéaire) à effets mixtes.

Ensuite, **au niveau de la population**, pour décrire la variabilité inter-individuelle, on modélise les paramètres individuels pour tout $i \in \{1, \dots, n\}$ par :

$$\varphi_i = X_i \beta + \xi_i \quad (2.2)$$

où $\beta \in \mathbb{R}^p$ est un vecteur d'effets fixes à estimer. X_i est une matrice de taille $q \times p$ de covariables connues de l'individu i et ξ_i est un bruit gaussien iid $\mathcal{N}_q(0, \Gamma)$, avec Γ la matrice de variance-covariance inter-individuelle. Cette modélisation caractérise la variabilité de φ_i à travers une association avec les covariables X_i , modélisée par β , et une variation inexpliquée par les covariables mesurées dans la population, représentée par ξ_i . Les paramètres du modèle à estimer sont alors $\theta = (\beta, \psi, \sigma^2, \Gamma)$, et appelés paramètres de population.

Exemple 1. En reprenant l'exemple des orangers de la section précédente (voir Figure 2.1), y_{ij} correspond à la circonférence du tronc de l'oranger numéro i au temps t_{ij} . De plus, une fonction f appropriée pour modéliser cette croissance est la fonction logistique :

$$f(\varphi_i, t_{ij}) = \frac{\varphi_{i1}}{1 + \exp\left(-\frac{t_{ij} - \varphi_{i2}}{\varphi_{i3}}\right)}.$$

On a alors que $\varphi_i = (\varphi_{i1}, \varphi_{i2}, \varphi_{i3})^\top$ avec φ_{i1} correspondant à la circonférence maximale limite du tronc, φ_{i2} correspondant au moment où la circonférence du tronc atteint la moitié de sa

circonférence maximale limite et φ_{i3} est un facteur d'échelle. Ces trois effets individuels peuvent, par exemple, être modélisés comme suit :

$$\begin{aligned}\varphi_{i1} &= \beta_{11} + \xi_{i1}, \\ \varphi_{i2} &= \beta_{21} + \beta_{22}x_i + \xi_{i2}, \\ \varphi_{i3} &= \beta_{31},\end{aligned}$$

où $\xi_i = (\xi_{i1}, \xi_{i2}, \xi_{i3})^\top \sim \mathcal{N}_3(0, \Gamma)$ avec :

$$X_i = \begin{pmatrix} 1 & x_i & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & x_i & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & x_i \end{pmatrix}, \quad \beta = \begin{pmatrix} \beta_{11} \\ 0 \\ \beta_{21} \\ \beta_{22} \\ \beta_{31} \\ 0 \end{pmatrix}, \quad \Gamma = \begin{pmatrix} \gamma_{11} & \gamma_{12} & 0 \\ \gamma_{12} & \gamma_{22} & 0 \\ 0 & 0 & 0 \end{pmatrix},$$

Les covariables mesurées x_i peuvent être des caractéristiques génétiques, comme par exemple la présence/absence de certains gènes.

2.1.3 Vraisemblance

Le cadre des modèles à effets mixtes fait intervenir ce qu'on appelle des *variables latentes*, c'est-à-dire des variables que l'on n'observe pas. Ici, le vecteur $\boldsymbol{\varphi} = (\varphi_1, \dots, \varphi_n)$ des paramètres individuels constitue l'ensemble des variables latentes du modèle. Ce type de modèle est aussi appelé modèle à données incomplètes. Ainsi, dans les modèles à données incomplètes, la vraisemblance des observations s'exprime comme l'intégrale de la vraisemblance conditionnelle des observations sachant les variables latentes par rapport à la distribution des variables latentes. Dans le modèle non-linéaire à effets mixtes, en notant $\mathbf{y}_i = (y_{ij})_{1 \leq j \leq n_i}$ le vecteur des observations de l'individu i , et $\mathbf{y} = (y_i)_{1 \leq i \leq n}$ le vecteur de toutes les observations, la forme générale de la vraisemblance des observations est donc la suivante :

$$p(\mathbf{y}; \theta) = \prod_{i=1}^n p(\mathbf{y}_i; \theta) = \prod_{i=1}^n \int p(\mathbf{y}_i | \varphi_i; \theta) p(\varphi_i; \theta) d\varphi_i, \quad (2.3)$$

par indépendance des individus. Les intégrales dans l'équation (2.3) ne sont généralement pas calculables, sauf si le modèle est linéaire. La vraisemblance des observations n'a alors pas d'expression analytique. Ainsi, le calcul de l'estimateur du maximum de vraisemblance, obtenu en maximisant la vraisemblance $p(\mathbf{y}; \theta)$ par rapport à θ , ne se fait pas de manière directe. Des solutions algorithmiques ont été proposées pour réaliser l'estimation des paramètres dans les modèles à variables latentes et sont détaillées dans la section suivante.

2.2 Méthodes d'inférence dans les modèles non-linéaires à effets mixtes

Une fois que le modèle est défini, l'objectif principal est d'évaluer la variabilité entre les individus ainsi que le profil caractéristique. En d'autres termes, il s'agit d'estimer la valeur du paramètre de

population θ qui décrit au mieux les données. La principale difficulté réside dans la complexité de l'expression de la vraisemblance (2.3). Cette section aborde différentes méthodes élaborées pour surmonter cet obstacle et permettre l'estimation des paramètres dans les modèles non-linéaires à effets mixtes. Cette section est principalement inspirée de Demidenko (2013) et Lavielle (2014).

2.2.1 Méthodes approchées

Une première catégorie de méthodes permettant de contourner le problème de la vraisemblance non explicite consiste à approcher le modèle par une version linéaire, soit par une approche en deux temps, soit par linéarisation du modèle.

2.2.1.1 Approche en deux temps

L'approche en deux temps (Pocock et al., 1981; Berkey and Laird, 1986) présentée dans cette section suppose que les individus ne partagent pas de paramètres fixes ψ . Dans ce cas, une approche un peu naïve consiste à estimer les paramètres individuels φ_i par la méthode des moindres carrés non-linéaires à partir de la première couche du modèle (2.1) individu par individu :

$$\hat{\varphi}_i = \operatorname{argmin}_{\varphi_i \in \mathbb{R}^q} \sum_{j=1}^{n_i} (y_{ij} - f(\varphi_i, t_{ij}))^2.$$

Lors de cette première étape, on résout les n problèmes de régression non-linéaires individuels par des techniques standard telles que l'algorithme de Gauss-Newton (Bard, 1974). Puis ces estimations sont traitées comme des observations pour la deuxième couche du modèle (2.2), et on se ramène ainsi à un problème d'estimation en régression linéaire gaussienne :

$$\hat{\varphi}_i = X_i \beta + \xi_i, \quad \xi_i \sim \mathcal{N}_q(0, \Gamma).$$

Cette méthode en deux temps (Pocock et al., 1981; Berkey and Laird, 1986) s'avère efficace dans des conditions où la variance résiduelle σ^2 est faible et où le nombre d'observations par individu est assez grand, permettant ainsi un ajustement précis des différents modèles individuels et donc une estimation précise des paramètres individuels $\hat{\varphi}_i$. Cependant, les limitations théoriques de cette approche découlent de sa sensibilité aux estimations des paramètres individuels. Plus précisément, les estimations des paramètres de population peuvent être affectées par des estimations incorrectes des paramètres individuels, qui sont utilisées dans leur calcul. Pour prendre en compte les erreurs d'estimations des paramètres individuels, plusieurs extensions ont été proposées, dont la méthode *Global two-stage* de Steimer et al. (1984). Le lecteur peut se référer à Davidian and Giltinan (1993) et Yeap and Davidian (2001) pour plus de discussions. Étant donné que les individus sont analysés de manière séparée, ces approches ne tirent pas pleinement profit de la fusion d'informations à l'échelle de la population offerte par les modèles à effets mixtes.

Il convient de noter qu'en présence d'effets fixes ψ , cette approche pourrait être généralisée par la minimisation de la somme totale des carrés :

$$\sum_{i=1}^n \sum_{j=1}^{n_i} (y_{ij} - f(\varphi_i, \psi, t_{ij}))^2.$$

Cependant, dans ce cas, le problème d'optimisation ne peut pas être séparé en n problèmes des moindres carrés non-linéaires distincts, ce qui augmente considérablement les coûts computationnels.

2.2.1.2 Méthodes de linéarisation

Parmi les premières techniques adaptées à l'estimation dans les modèles non-linéaires à effets mixtes, figurent celles qui reposent sur le principe de linéarisation. Cela peut impliquer soit la linéarisation de la fonction de régression, soit celle de la fonction de vraisemblance. L'idée de ces méthodes est de construire une approximation du modèle pour obtenir une expression analytique de la vraisemblance, qui est alors ensuite maximisée en utilisant un algorithme de type Newton-Raphson. Celui-ci est détaillé dans la section 2.2.2.3.

Premièrement, [Sheiner et al. \(1972, 1977\)](#) proposent d'effectuer une approximation au premier ordre de la fonction de régression non-linéaire autour de la moyenne des paramètres individuels, appelée méthode FOA (*First-Order Approximation*). Ensuite, [Sheiner and Beal \(1980, 1981, 1983\)](#) proposent une analyse de population basée sur la méthode FOA, appelée NONMEM (*NONlinear Mixed-Effects Models*).

Plus précisément, la méthode FOA repose sur le développement limité à l'ordre 1 de la fonction de régression f . Définissons $f_i(\varphi_i, \psi) = (f(\varphi_i, \psi, t_{i1}), \dots, f(\varphi_i, \psi, t_{in_i}))^\top$ et $\varepsilon_i = (\varepsilon_{ij})_{1 \leq j \leq n_i}$ de sorte que l'équation (2.1) se réécrit :

$$\mathbf{y}_i = f_i(\varphi_i, \psi) + \varepsilon_i, \quad i = 1, \dots, n.$$

Alors en remplaçant $f_i(X_i\beta + \xi_i, \psi)$ par son approximation linéaire à l'ordre 1 autour de $\mathbb{E}[\varphi_i] = X_i\beta$, le modèle se réduit à un *modèle marginal non-linéaire à effets mixtes* ([Demidenko, 2013](#)) :

$$\mathbf{y}_i = f_i(X_i\beta, \psi) + Z_i(\beta)\xi_i + \varepsilon_i,$$

avec $Z_i(\beta) = \partial f_i(X_i\beta, \psi) / \partial \varphi_i$. La différence majeure entre le modèle non-linéaire à effets mixtes (2.1)-(2.2) et le modèle marginal non-linéaire à effets mixtes est que dans ce dernier l'effet aléatoire est introduit de manière linéaire dans le modèle. Ainsi, en particulier, le qualificatif "marginal" provient du fait que, contrairement aux modèles non-linéaires à effets mixtes, l'espérance et la covariance des observations possèdent une expression explicite en fonction du paramètre de population, ce qui facilite l'estimation par maximum de vraisemblance : $\mathbb{E}[\mathbf{y}_i] = f_i(X_i\beta, \psi)$, $Cov(\mathbf{y}_i) = Z_i(\beta)\Gamma Z_i(\beta)^\top + \sigma^2 I_{n_i}$, où I_{n_i} désigne la matrice identité de taille n_i . Cette approximation constitue le support des développements théoriques proposés dans le chapitre 6.

Cependant, malgré l'avantage inhérent à la méthode NONMEM de ne pas exiger d'estimations des paramètres individuels, l'approche FOA utilisée dans NONMEM tend à générer des estimations biaisées des paramètres. Ce biais est particulièrement marqué en présence d'une variabilité inter-individuelle significative, puisque la linéarisation entraîne que la réponse chez un individu "moyen" se confond avec la réponse "moyenne" d'un individu. [Lindstrom and Bates \(1990\)](#) proposent la méthode FOCE (*First-Order Conditional Estimation*), et tentent d'améliorer l'approximation de premier ordre en linéarisant non plus autour de $\mathbb{E}[\varphi_i] = X_i\beta$, mais autour des moyennes spécifiques à l'individu $\mathbb{E}[\varphi_i] + \hat{\xi}_i$, où $\hat{\xi}_i$ est la moyenne a posteriori dans le cas linéaire et en utilisant l'approximation gaussienne de la posterior de [Stiratelli et al. \(1984\)](#) pour estimer les paramètres de variance. Bien que la méthode FOCE soit plus précise, elle est aussi plus coûteuse

en termes de temps de calcul que la méthode FOA, déjà assez exigeante en ressources computationnelles, notamment si un algorithme de type Newton-Raphson est employé pour maximiser la vraisemblance.

D'autres stratégies reposent sur une approximation de la vraisemblance du modèle. Par exemple, [Wolfinger \(1993\)](#) montre que l'approche de [Lindstrom and Bates \(1990\)](#) peut être dérivée comme une approximation de la vraisemblance par l'approximation de Laplace :

$$\int e^{nl(\theta)} d\theta \approx (2\pi/n)^{r/2} |l''(\hat{\theta})|^{-1/2} e^{nl(\hat{\theta})},$$

où l est la log-vraisemblance, θ est le vecteur des paramètres de dimension r , $\hat{\theta}$ maximise $e^{nl(\theta)}$, et n est grand. Toutefois, aucune de ces méthodes n'est étayée par des résultats théoriques démontrant leur convergence vers le maximum de vraisemblance. En particulier, [Vonesh \(1996\)](#) montre que les estimateurs FOA et ceux basés sur l'approximation de Laplace ne sont pas consistants quand le nombre d'observations par individu croît moins vite que le nombre d'individus.

2.2.2 Méthodes exactes

Le défaut principal des méthodes précédentes réside dans le fait qu'elles cherchent à maximiser une vraisemblance approchée plutôt que la vraisemblance exacte. L'approche de maximisation exacte de la vraisemblance du modèle demeure une stratégie cruciale. À la fois des méthodes par intégration numérique, mais aussi des méthodes par intégration stochastique ont été proposées pour le calcul exact de la vraisemblance (2.3) dans le cadre des modèles non-linéaires à effets mixtes. Dans ces approches, la vraisemblance obtenue est qualifiée d'"exacte" dans le sens où l'approximation numérique ou stochastique peut être rendue aussi précise que nécessaire moyennant un coût computationnel accru.

2.2.2.1 Méthodes d'intégration numérique

Une méthode d'intégration numérique basée sur la quadrature de Gauss pour calculer la vraisemblance (2.3) a par exemple été proposée par [Davidian and Gallant \(1993\)](#), puis améliorée par [Pinheiro and Bates \(1995\)](#) sous le nom de quadrature gaussienne adaptative. La maximisation de la vraisemblance s'obtient alors en utilisant un algorithme de type Newton-Raphson en calculant à chaque itération la vraisemblance par approximation numérique. Cependant, ces algorithmes présentent des problèmes de convergence et de lenteur informatique.

2.2.2.2 Méthodes bayésiennes ou d'intégration stochastique

Les méthodes bayésiennes se présentent comme une alternative traditionnelle aux approches fondées sur le maximum de vraisemblance. Dans ce cadre, le paramètre θ est traité comme une variable aléatoire, associée à une loi a priori. L'objectif des méthodes bayésiennes consiste à évaluer la distribution a posteriori, reflétant notre connaissance mise à jour après l'observation des données, définie par l'expression suivante :

$$\pi(\theta|\mathbf{y}) = \frac{p(\mathbf{y};\theta)\pi(\theta)}{p(\mathbf{y})},$$

où $p(\mathbf{y};\theta)$ représente la vraisemblance des observations $\mathbf{y}$, $\pi(\theta)$ le prior, et $p(\mathbf{y})$ la densité marginale de $\mathbf{y}$, $p(\mathbf{y}) = \int p(\mathbf{y};\theta)\pi(\theta)d\theta$. Comme précisé précédemment, dans notre contexte,

la vraisemblance (2.3) ne possède pas d'expression analytique explicite, ce qui rend l'évaluation de la distribution a posteriori complexe. Pour surmonter cette difficulté, l'utilisation de méthodes de Monte Carlo par chaînes de Markov (MCMC) a été proposée pour approcher numériquement l'intégrale multidimensionnelle de la vraisemblance (2.3) (Wakefield et al., 1994). Plus précisément, les méthodes MCMC permettent d'obtenir des réalisations d'une variable aléatoire de loi π spécifiée par l'utilisateur. Cette méthode peut être utilisée pour l'approximation stochastique d'intégrales en utilisant que :

$$\int \phi(x)\pi(x)dx = \mathbb{E}[\phi(X)], \quad \text{avec } X \sim \pi.$$

Ces méthodes MCMC sont plus largement détaillées dans la section 2.3.4.1.

2.2.2.3 Méthodes de maximisation de la vraisemblance

L'algorithme de Newton-Raphson (Verbeke and Cools, 1995) est un algorithme classique de minimisation basé sur la recherche d'un zéro du gradient de la fonction à minimiser. Il s'agit d'un algorithme itératif qui met à jour la valeur du paramètre à chaque itération selon le schéma suivant :

$$\theta_{k+1} = \theta_k + H(\theta_k)^{-1}J(\theta_k),$$

où les fonctions $J(\theta) = \frac{\partial \log p(\mathbf{y}; \theta)}{\partial \theta}$ et $H(\theta) = \frac{\partial^2 \log p(\mathbf{y}; \theta)}{\partial \theta \partial \theta^\top}$ sont respectivement la jacobienne et la hessienne de la log-vraisemblance. Comme discuté précédemment, il peut être appliqué en utilisant à chaque itération l'approximation de la vraisemblance obtenue par l'une des méthodes précédentes. Sinon, en reliant les fonctions $J(\theta)$ et $H(\theta)$ à la vraisemblance des données complètes $p(\mathbf{y}, \boldsymbol{\varphi}; \theta)$ par les formules de Fisher (Fisher, 1922) et de Louis (Louis, 1982), respectivement :

$$J(\theta) = \mathbb{E} \left[\frac{\partial \log p(\mathbf{y}, \boldsymbol{\varphi}; \theta)}{\partial \theta} \middle| \mathbf{y}; \theta \right],$$

$$H(\theta) = \mathbb{E} \left[\frac{\partial^2 \log p(\mathbf{y}, \boldsymbol{\varphi}; \theta)}{\partial \theta \partial \theta^\top} \middle| \mathbf{y}; \theta \right] + \mathbb{V} \left[\frac{\partial \log p(\mathbf{y}, \boldsymbol{\varphi}; \theta)}{\partial \theta} \middle| \mathbf{y}; \theta \right],$$

on peut directement approcher ces fonctions par des schémas numériques. Par exemple, McCulloch (1997) propose d'utiliser des sommes de Monte-Carlo basées sur des réalisations simulées des paramètres individuels $\boldsymbol{\varphi}$ obtenues à l'aide d'un algorithme MCMC. Pour alléger les temps de calculs, Gu and Li (1998) proposent d'utiliser un schéma d'approximation stochastique.

Parmi ces méthodes basées sur la maximisation exacte de la vraisemblance, l'algorithme EM (*Expectation-Maximization*) (Dempster et al., 1977) et ses variantes, telles que l'algorithme SAEM (*Stochastic Approximation EM*) (Delyon et al., 1999), sont très répandus. L'algorithme EM est un algorithme itératif qui vise à construire une suite de valeurs de paramètres $(\theta^{(k)})_k$ convergeant vers l'estimateur du maximum de vraisemblance. À noter que ces algorithmes possèdent des propriétés de convergence seulement dans le cas où la vraisemblance complète appartient à la famille exponentielle courbe (voir équation (2.4)). Cet algorithme et ses principales extensions sont les plus couramment utilisés pour estimer les paramètres dans les modèles à variables latentes. Ils sont examinés en détail dans la section suivante.

En outre, il existe des méthodes de descente de gradient pour résoudre le problème de maximisation de la vraisemblance (2.3) :

$$\theta_{k+1} = \theta_k + \rho_{k+1} \nabla \log p(\mathbf{y}; \theta),$$

où ρ_k est le pas et ∇h est le gradient d'une fonction h . Notamment il existe des versions stochastiques de l'algorithme de descente du gradient qui ont été spécifiquement élaborées pour les modèles à variables latentes (Fang and Li, 2021; Baey et al., 2023), sans exiger que la vraisemblance complète appartienne à la famille exponentielle courbe.

2.3 Algorithme Expectation-Maximisation et ses extensions

Dans cette section, on considère le contexte plus large des modèles à variables latentes, incluant notamment les modèles non-linéaires à effets mixtes. Nous considérons des observations $\mathbf{y} = (y_i)_{1 \leq i \leq n}$ et des variables latentes $Z = (Z_i)_{1 \leq i \leq m}$ qui caractérisent la distribution des observations $\mathbf{y}$, où $n \in \mathbb{N}^*$ et $m \in \mathbb{N}^*$. On suppose que la vraisemblance complète f de $(\mathbf{y}, Z)$ appartient à une famille paramétrique $\{f_\theta(\mathbf{y}, Z), \theta \in \Theta\}$.

2.3.1 Algorithme EM

L'algorithme EM (*Expectation Maximisation*) est un algorithme itératif, proposé par Dempster et al. (1977), qui permet d'approcher l'estimateur du maximum de vraisemblance (ou l'estimateur du maximum *a posteriori* dans un cadre bayésien) du paramètre d'un modèle probabiliste lorsque celui-ci n'admet pas d'expression explicite. Il est particulièrement utilisé pour de nombreux modèles à variables latentes.

Dans un cadre fréquentiste, notre intérêt se porte sur la maximisation de la vraisemblance des observations $g_\theta(\mathbf{y})$:

$$\hat{\theta}^{EMV} = \operatorname{argmax}_{\theta \in \Theta} g_\theta(\mathbf{y}) = \operatorname{argmax}_{\theta \in \Theta} \int f_\theta(\mathbf{y}, Z) dZ.$$

Cependant, l'intégrale précédente étant rarement connue analytiquement, l'idée clé de l'algorithme EM est de la maximiser en maximisant itérativement une quantité plus facile à calculer :

$$Q(\theta|\theta') = \mathbb{E}_{\theta'}[\log f_\theta(\mathbf{y}, Z)|\mathbf{y}] = \int \log(f_\theta(\mathbf{y}, Z)) p_{\theta'}(Z|\mathbf{y}) dZ,$$

l'espérance conditionnelle de la log-vraisemblance des données complètes $(\mathbf{y}, Z)$ sachant les observations $\mathbf{y}$ par rapport à la valeur courante de l'estimation du paramètre θ' . Cette approche est intéressante car elle tire avantage du fait que la log-vraisemblance des données complètes $f_\theta(\mathbf{y}, Z)$ est connue explicitement.

L'algorithme EM repose sur le fait que tout accroissement de l'espérance conditionnelle de la log-vraisemblance des données complètes entraîne systématiquement une augmentation de la vraisemblance des observations $g_\theta(\mathbf{y})$ que l'on cherche à maximiser (Dempster et al., 1977). Les étapes de l'algorithme EM sont résumées dans le pseudo-code de l'Algorithme 1.

L'algorithme EM possède une propriété de convergence garantissant que, sous certaines hypothèses de régularité sur le modèle, la suite $(\theta^{(k)})_k$ ainsi construite converge vers un point stationnaire $\hat{\theta}_g$ de la vraisemblance observée $g_\theta(\mathbf{y})$ (Dempster et al., 1977; Wu, 1983). De plus, sous des hypothèses plus strictes, Delyon et al. (1999) ont établi la convergence de cette suite vers un maximum local de la vraisemblance observée. De fait, l'algorithme EM étant une procédure déterministe, le choix du point initial θ_0 est crucial. En pratique, il est courant de lancer

Algorithm 1 Expectation Maximisation (EM)

Entrée : $K \in \mathbb{N}^*$ nombre d'itérations de l'algorithme, $\theta^{(0)}$ valeur initiale du paramètre.

Pour $k = 0$ à $K - 1$

1. **Étape E (Expectation) :** Calculer l'espérance conditionnelle de la log-vraisemblance des données complètes $(\mathbf{y}, Z)$ sachant $\mathbf{y}$ par rapport à la valeur courante du paramètre $\theta^{(k)}$:

$$Q(\theta|\theta^{(k)}) = \mathbb{E}_{\theta^{(k)}}[\log f_{\theta}(\mathbf{y}, Z)|\mathbf{y}]$$

2. **Étape M (Maximisation) :** Mettre à jour la valeur de θ en maximisant la quantité Q selon

$$\theta^{(k+1)} = \operatorname{argmax}_{\theta \in \Theta} Q(\theta|\theta^{(k)})$$

Sortie : $\hat{\theta}^{EMV} = \theta^{(K)}$, pour K assez grand.

l'algorithme plusieurs fois avec des initialisations différentes en espérant trouver ainsi le maximum global.

Dans le cadre bayésien, un estimateur classique est l'estimateur du maximum a posteriori (MAP) défini par :

$$\hat{\theta}^{MAP} = \operatorname{argmax}_{\theta \in \Theta} \pi(\theta|\mathbf{y})$$

où $\pi(\theta|\mathbf{y})$ est la loi *a posteriori* de θ . L'algorithme EM s'adapte très bien à ce cas (Dempster et al., 1977; Green, 1990) et avec les mêmes propriétés de convergence, en remplaçant la fonction $Q(\theta|\theta^{(k)})$ de l'étape E par :

$$Q(\theta|\theta^{(k)}) = \mathbb{E}[\log \pi(\theta, Z|\mathbf{y})|\mathbf{y}, \theta^{(k)}].$$

Cependant, dans certaines situations, l'étape E n'est pas traitable car la quantité $Q(\theta|\theta^{(k)})$ n'a pas d'expression analytique. Cette quantité a bien une forme fermée pour les modèles de mélange ou les modèles linéaires à effets mixtes, par exemple, mais pas pour les modèles non-linéaires à effets mixtes. Ainsi, plusieurs alternatives ont été proposées dans la littérature pour approcher numériquement cette espérance conditionnelle à chaque itération. Ces méthodes sont présentées dans les sections suivantes.

2.3.2 Algorithme MCEM

Wei and Tanner (1990) proposent une approximation de la quantité Q de l'étape E de l'algorithme EM par une somme de Monte-Carlo en générant à chaque itération un grand nombre de réalisations indépendantes des variables latentes à partir de la distribution conditionnelle sachant les observations $p_{\theta^{(k)}}(Z|\mathbf{y})$. Cet algorithme, appelé MCEM (*Monte-Carlo Expectation Maximisation*), est détaillé dans le pseudo-code ci-après, Algorithme 2.

En pratique, la convergence de l'algorithme MCEM est conditionnée par le nombre L de réalisations indépendantes des variables latentes. Plus précisément, la convergence du MCEM exige que L augmente à chaque itération (Fort and Moulines, 2003), ce qui peut engendrer des coûts considérables en termes de temps de calcul. Il est à noter que la série de réalisations $(Z_{\ell}^{(k)})_{1 \leq \ell \leq L}$ générées à chaque itération est systématiquement abandonnée. Afin de tirer un

Algorithm 2 Monte Carlo EM (MCEM)

Entrée : $K \in \mathbb{N}^*$ nombre d'itérations de l'algorithme, $L \in \mathbb{N}^*$ nombre de réalisations des variables latentes, $\theta^{(0)}$ valeur initiale du paramètre.

Pour $k = 0$ à $K - 1$

1. **Étape S (Simulation) :** Simuler L réalisations des variables latentes Z , notées $(Z_\ell^{(k)})_{1 \leq \ell \leq L}$, selon la distribution conditionnelle sachant les observations par rapport à la valeur courante du paramètre $p_{\theta^{(k)}}(Z|\mathbf{y})$.
2. **Étape E (Expectation) :** Approcher $Q(\theta|\theta^{(k)})$ selon la somme de Monte-Carlo suivante :

$$Q(\theta|\theta^{(k)}) \approx \tilde{Q}(\theta|\theta^{(k)}) = \frac{1}{L} \sum_{\ell=1}^L \log f_\theta(\mathbf{y}, Z_\ell^{(k)})$$

3. **Étape M (Maximisation) :** Mettre à jour la valeur de θ en maximisant la quantité $\tilde{Q}$ selon

$$\theta^{(k+1)} = \underset{\theta \in \Theta}{\operatorname{argmax}} \tilde{Q}(\theta|\theta^{(k)})$$

Sortie : $\hat{\theta}^{EMV} = \theta^{(K)}$ pour K assez grand.

meilleur parti de ces réalisations, [Delyon et al. \(1999\)](#) ont suggéré l'introduction d'un schéma d'approximation stochastique visant à estimer la quantité Q . Cet algorithme, appelé SAEM (*Stochastic Approximation version of EM*), est détaillé dans la section suivante.

2.3.3 Algorithme SAEM : une version stochastique de l'algorithme EM

2.3.3.1 Description générale de l'algorithme

L'algorithme SAEM (*Stochastic Approximation version of EM*), introduit par [Delyon et al. \(1999\)](#), est une alternative à l'algorithme EM lorsque l'étape E n'est pas traitable. L'idée de l'algorithme SAEM est d'approcher l'espérance conditionnelle de la log-vraisemblance complète par une procédure d'approximation stochastique. Plus précisément, l'étape E de l'algorithme EM est remplacée par deux étapes, appelées étape S et étape SA, pour approcher itérativement la quantité $Q(\theta|\theta^{(k)})$. Un pseudo-code de l'algorithme SAEM est présenté dans l'algorithme 3.

Il convient de noter que si le nombre d'observations est petit, la convergence de l'algorithme SAEM peut être amélioré en simulant $L > 1$ réalisations des variables latentes $(Z_\ell^{(k)})_{1 \leq \ell \leq L}$ à l'étape S ([Lavielle, 2014](#)), et donc l'étape SA combine les approches par approximation stochastique et par somme de Monte-Carlo comme suit :

$$Q_{k+1}(\theta) = Q_k(\theta) + \gamma_k \left(\frac{1}{L} \sum_{\ell=1}^L \log f_\theta(\mathbf{y}, Z_\ell^{(k)}) - Q_k(\theta) \right).$$

On peut remarquer que l'algorithme SAEM est la généralisation de deux autres versions stochastiques de l'algorithme EM. En effet, si $\gamma_k = 1$ pour tout k , alors il n'y a pas de mémoire dans l'approximation stochastique et alors $Q_{k+1}(\theta) = \log f_\theta(\mathbf{y}, Z^{(k)})$. Cet algorithme est appelé Stochastic EM (SEM, se référer à [Celeux and Diebolt \(1986\)](#); [Diebolt and Ip \(1996\)](#)). Il consiste à itérativement simuler une réalisation des variables latentes $Z^{(k)}$ selon la conditionnelle

Algorithm 3 Stochastic Approximation of EM (SAEM)

Entrée : $K \in \mathbb{N}^*$ nombre d'itérations de l'algorithme, $\theta^{(0)}$ valeur initiale du paramètre, $Q_0 = 0$ et $(\gamma_k)_k$ une suite de pas positive décroissante vers 0 telle que $\forall k, \gamma_k \in [0, 1]$, $\sum_k \gamma_k = \infty$ et $\sum_k \gamma_k^2 < \infty$.

Pour $k = 0$ à $K - 1$

1. **Étape S (Simulation) :** Simuler une réalisation $Z^{(k)}$ des variables latentes selon la loi conditionnelle sachant les observations $p_{\theta^{(k)}}(Z|\mathbf{y})$
2. **Étape SA (Stochastic Approximation) :** Mettre à jour $Q_{k+1}(\theta)$, approximation stochastique de $Q(\theta|\theta^{(k+1)}) = \mathbb{E}_{\theta^{(k+1)}}[\log f_{\theta}(\mathbf{y}, Z)|\mathbf{y}]$, selon :

$$Q_{k+1}(\theta) = Q_k(\theta) + \gamma_k(\log f_{\theta}(\mathbf{y}, Z^{(k)}) - Q_k(\theta)).$$

3. **Étape M (Maximisation) :** Mettre à jour la valeur de θ selon

$$\theta^{(k+1)} = \operatorname{argmax}_{\theta \in \Theta} Q_{k+1}(\theta).$$

Sortie : $\hat{\theta}^{EMV} = \theta^{(K)}$ pour K assez grand.

$p_{\theta^{(k)}}(Z|\mathbf{y})$, puis de mettre à jour $\theta^{(k+1)}$ en maximisant la distribution jointe $f_{\theta}(\mathbf{y}, Z^{(k)})$. De même, si $\gamma_k = 1$ pour tout k et $L \gg 1$, alors on retrouve l'algorithme MCEM décrit dans la section précédente.

Delyon et al. (1999) ont montré que l'algorithme SAEM converge avec probabilité 1 vers un maximum local de la vraisemblance des observations g sous des hypothèses de régularité sur le modèle. Par conséquent, en pratique, il est nécessaire d'exécuter l'algorithme avec diverses conditions initiales et de comparer les résultats obtenus afin de déterminer le maximum global. Du point de vue de l'implémentation, l'algorithme SAEM ne nécessite qu'une seule réalisation des variables latentes (ou un petit nombre) par itération. Ainsi, le coût computationnel est nettement moins élevé que dans le cas de l'algorithme MCEM. Un autre avantage réside dans le fait qu'en pratique, seulement un nombre faible d'itérations est nécessaire pour que l'algorithme atteigne la convergence. De plus, de même que pour l'algorithme EM, l'algorithme SAEM s'adapte à un contexte bayésien (Delyon et al., 1999).

2.3.3.2 Version simplifiée de l'algorithme pour les modèles exponentiels courbes

L'implémentation de l'algorithme SAEM peut être simplifiée dans le cas où la vraisemblance des données complètes appartient à une famille de modèles exponentiels courbes. En effet, l'une des hypothèses de régularité nécessaire à la preuve de la convergence de cet algorithme (Delyon et al., 1999) est l'appartenance du modèle à la famille exponentielle courbe. Plus précisément, on suppose que la vraisemblance complète peut s'écrire sous la forme :

$$f_{\theta}(\mathbf{y}, Z) = \exp \left(-\psi(\theta) + \langle S(\mathbf{y}, Z), \phi(\theta) \rangle \right) \quad (2.4)$$

où ψ et ϕ sont deux fonctions sur l'espace des paramètres Θ , $S(\mathbf{y}, Z)$ est une *statistique exhaustive* et avec $\langle \cdot, \cdot \rangle$ le produit scalaire canonique. Dans ce cas, la quantité Q se réécrit donc :

$$Q(\theta|\theta^{(k)}) = \mathbb{E}_{\theta^{(k)}}[\log f_{\theta}(\mathbf{y}, Z)|\mathbf{y}] = -\psi(\theta) + \langle \mathbb{E}_{\theta^{(k)}}[S(\mathbf{y}, Z)|\mathbf{y}], \phi(\theta) \rangle.$$

Il suffit donc de se concentrer sur l'espérance conditionnelle de la statistique exhaustive $\mathbb{E}_{\theta^{(k)}}[S(\mathbf{y}, Z)|\mathbf{y}]$ pour avoir une approximation stochastique de $Q(\theta|\theta^{(k)})$. Plus précisément, le pseudo-code de l'algorithme SAEM dans le cas de la famille exponentielle courbe est présenté dans l'algorithme 4.

Algorithm 4 SAEM (famille exponentielle courbe)

Entrée : $K \in \mathbb{N}^*$ nombre d'itérations de l'algorithme, $\theta^{(0)}$ valeur initiale du paramètre, $S_0 = 0$ et $(\gamma_k)_k$ une suite de pas positive décroissante vers 0 telle que $\forall k, \gamma_k \in [0, 1]$, $\sum_k \gamma_k = \infty$ et $\sum_k \gamma_k^2 < \infty$.

Pour $k = 0$ à $K - 1$

1. **Étape S (Simulation) :** Simuler une réalisation $Z^{(k)}$ des variables latentes selon la loi conditionnelle sachant les observations $p_{\theta^{(k)}}(Z|\mathbf{y})$
2. **Étape SA (Stochastic Approximation) :** Mettre à jour s_{k+1} , approximation stochastique de $\mathbb{E}_{\theta^{(k+1)}}[S(\mathbf{y}, Z)|\mathbf{y}]$, selon :

$$s_{k+1} = s_k + \gamma_k (S(\mathbf{y}, Z^{(k)}) - s_k).$$

3. **Étape M (Maximisation) :** Mettre à jour la valeur de θ selon

$$\theta^{(k+1)} = \underset{\theta \in \Theta}{\operatorname{argmax}} L(s_{k+1}, \theta),$$

où $L(s, \theta) := -\psi(\theta) + \langle s, \phi(\theta) \rangle$.

Sortie : $\hat{\theta}^{EMV} = \theta^{(K)}$ pour K assez grand.

2.3.4 Algorithme MCMC-SAEM : couplage avec une procédure MCMC

L'étape S de l'algorithme SAEM requiert de pouvoir calculer facilement la loi conditionnelle des variables latentes Z sachant les observations $\mathbf{y}$. Cependant, dans de nombreuses situations, telles que les modèles non-linéaires à effets mixtes, la loi conditionnelle de Z n'est pas explicite et donc l'étape S n'est pas faisable en l'état. Ainsi, [Kuhn and Lavielle \(2004\)](#) ont proposé une version améliorée de cet algorithme qui consiste à coupler la méthode SAEM avec une procédure MCMC (*Markov Chain Monte Carlo*) lorsque l'étape de simulation n'est pas réalisable du fait qu'on ne connaît la loi conditionnelle de Z sachant $\mathbf{y}$ qu'à constante multiplicative près. Plus précisément, cette alternative implique, par exemple, l'utilisation d'un algorithme de Metropolis-Hastings pour générer une chaîne de Markov ayant la loi conditionnelle comme loi stationnaire.

2.3.4.1 Méthode MCMC et algorithme de Metropolis-Hastings

L'idée d'une méthode de Monte-Carlo par chaîne de Markov est de construire une chaîne de Markov $(X_n)_n$ convergente ayant pour loi stationnaire une certaine loi d'intérêt h que l'on souhaite simuler. On laisse ce processus évoluer jusqu'à un temps N assez grand pour que la distribution de X_N soit assez proche de la distribution stationnaire h . La variable X_N est alors utilisée comme échantillon de la loi h .

Un algorithme classique utilisant cette méthode est l'*algorithme de Metropolis-Hastings* ([Hastings, 1970](#)). L'idée est d'explorer la loi cible h avec une marche aléatoire utilisant une loi de proposition pour se déplacer. Cette loi de proposition doit être choisie telle qu'elle soit

facile à simuler. Ainsi, une étape de cet algorithme itératif se déroule en deux phases : premièrement on propose un candidat en simulant une réalisation sous la loi de proposition, puis on accepte ou on refuse ce candidat avec une probabilité adaptée qui assure que la loi cible soit bien la loi stationnaire de la chaîne de Markov que l'on est en train de construire. Plus précisément, l'algorithme 5 présente un pseudo-code de l'algorithme de Metropolis-Hastings.

Algorithm 5 Metropolis-Hastings

Entrée : $R \in \mathbb{N}^*$ nombre d'itérations de l'algorithme, X_0 valeur initiale de la chaîne de Markov, une loi de proposition q .

Pour $r = 0$ à $R - 1$

1. Simuler $X^{(c)}$ selon la loi de proposition $q(X|X_r)$
2. Calculer la probabilité d'acceptation :

$$\alpha(X^{(c)}, X_r) = \min \left(1, \frac{h(X^{(c)})}{q(X^{(c)}|X_r)} \frac{q(X_r|X^{(c)})}{h(X_r)} \right)$$

3. Construire $X_{r+1} = \begin{cases} X^{(c)} & \text{avec probabilité } \alpha(X^{(c)}, X_r) \\ X_r & \text{avec probabilité } 1 - \alpha(X^{(c)}, X_r) \end{cases}$

Sortie : Retourner $(X_N, X_{N+1}, \dots, X_R)$ pour N assez grand.

La probabilité d'acceptation est choisie telle que l'on accepte toujours un candidat qui fait croître le ratio "loi cible sur loi de proposition". En effet,

$$\alpha(X^{(c)}, X_r) = 1 \iff \frac{h(X^{(c)})}{q(X^{(c)}|X_r)} \geq \frac{h(X_r)}{q(X_r|X^{(c)})}$$

Ce ratio est choisi pour vérifier les équations de balance détaillée des chaînes de Markov, permettant de montrer la stationnarité d'une loi (Hastings, 1970). De plus, les performances de cet algorithme dépendent du choix de la loi de proposition, laquelle doit notamment avoir un support plus large que la distribution cible. La taille des pas parcourus dans la chaîne dépend également de cette loi de proposition, ce qui influence la vitesse de convergence vers la loi stationnaire. À noter que puisque seuls les ratios interviennent, il suffit de ne connaître la loi cible qu'à une constante multiplicative près.

2.3.4.2 Couplage MCMC-SAEM

Dans l'étape S de l'algorithme SAEM, notre loi d'intérêt est $p_\theta(Z|\mathbf{y})$. Par la formule de Bayes, le ratio d'acceptation se réécrit donc :

$$\begin{aligned} \alpha(Z^{(c)}, Z_r) &= \min \left(1, \frac{p_\theta(Z^{(c)}|\mathbf{y})}{q(Z^{(c)}|Z_r)} \frac{q(Z_r|Z^{(c)})}{p_\theta(Z_r|\mathbf{y})} \right) \\ &= \min \left(1, \frac{p_\theta(\mathbf{y}|Z^{(c)})p_\theta(Z^{(c)})}{p_\theta(\mathbf{y}|Z_r)p_\theta(Z_r)} \frac{q(Z_r|Z^{(c)})}{q(Z^{(c)}|Z_r)} \right) \end{aligned}$$

où $p_\theta(\mathbf{y}|Z)$ désigne la loi conditionnelle de $\mathbf{y}$ sachant Z et $p_\theta(Z)$ désigne la loi de Z dans le modèle statistique. Ainsi on remarque qu'en choisissant comme loi de proposition la loi de Z ,

$q(Z|Z') = p_\theta(Z)$, le ratio d'acceptation se simplifie en :

$$\alpha(Z^{(c)}, Z_r) = \min \left(1, \frac{p_\theta(\mathbf{y}|Z^{(c)})}{p_\theta(\mathbf{y}|Z_r)} \right).$$

L'étape S du couplage MCMC-SAEM consiste alors à simuler une réalisation des variables latentes $Z^{(k)}$ en utilisant le résultat de M itérations d'une procédure MCMC ayant pour loi cible $p_{\theta^{(k)}}(Z|\mathbf{y})$. Kuhn and Lavielle (2005) recommandent de choisir seulement $M = 5$ ou 10 en pratique, car une plus grande valeur de M n'améliore pas la convergence de l'algorithme. Ce couplage MCMC-SAEM permet donc d'effectuer l'étape de simulation même dans le cas où la loi conditionnelle des variables latentes sachant les observations n'est connue qu'à une constante multiplicative près. Kuhn and Lavielle (2004) ont montré que ce couplage conserve les propriétés de convergence de l'algorithme SAEM classique. De plus, Kuhn and Lavielle (2005) traitent le cas particulier des modèles non-linéaires à effets mixtes et montrent les bonnes performances de cet algorithme dans ce cadre. Notamment, il est à noter que dans l'étape de maximisation on dispose dans ce cas de formules explicites pour la mise à jour des paramètres qui dépendent uniquement des statistiques exhaustives.

Sélection de variables en grande dimension

Dans le chapitre précédent, nous avons examiné en détail le processus d'inférence dans les modèles à effets mixtes, sans toutefois aborder la sélection de variables dans ce contexte. En effet, l'identification des covariables les plus pertinentes pour décrire la variabilité inter-individuelle est essentielle pour assurer une description précise des sources de variabilité. Cette thèse met l'accent sur le processus de sélection des covariables dans les modèles non-linéaires à effets mixtes, afin d'établir des liens entre la variabilité inter-individuelle et les caractéristiques individuelles mesurées dans la population. Par ailleurs, dans le contexte de la thèse, les covariables considérées sont généralement nombreuses puisque les variétés sont caractérisées par des milliers de covariables génétiques, dont on sait que la plupart d'entre elles n'ont aucun effet sur certains traits phénotypiques. La grande dimension des données génomiques implique d'abord la sélection de variables dans un cadre où le nombre de covariables est plus grand que le nombre d'individus. Plus précisément, la sélection de variables dans les modèles non-linéaires à effets mixtes porte sur la deuxième couche du modèle défini par l'équation (2.2). Il est intéressant de noter que si les paramètres individuels étaient observés, cette couche correspondrait à un modèle classique de régression linéaire gaussienne. Dans ces modèles plus classiques, de nombreuses méthodes de sélection et d'études théoriques ont déjà été proposées. Celles-ci sont présentées dans ce chapitre.

Dans ce chapitre, nous proposons d'introduire les défis inhérents à la sélection de variables dans un contexte de grande dimension dans le cadre de la régression linéaire gaussienne. La première section 3.1 est dédiée aux procédures de sélection de variables en grande dimension, couvrant à la fois les méthodes pénalisées fréquentistes et les priors bayésiens favorisant la parcimonie, avec une attention particulière sur les priors de mélange *spike-and-slab*. Puis, la deuxième section 3.2 se concentre sur l'analyse fréquentiste des distributions a posteriori bayésiennes sous contrainte de parcimonie. Plus précisément, elle introduit le paradigme fréquentiste pour la convergence des distributions a posteriori, dont les notions de taux de contraction et de consistance en sélection de la posterior.

Table des matières

3.1	Sélection de variables en grande dimension	27
3.1.1	Méthodes pénalisées fréquentistes	28
3.1.1.1	Régularisation de la vraisemblance	28
3.1.1.2	Sélection de modèles	29
3.1.2	Priors bayésiens induisant la parcimonie	30
3.1.2.1	Extension de la régularisation fréquentiste	31
3.1.2.2	Les priors de mélange <i>spike-and-slab</i>	32
3.1.2.3	Autres formes de priors favorisant la parcimonie	34

3.2	Analyse fréquentiste des distributions a posteriori	35
3.2.1	Paradigme fréquentiste pour la convergence des distributions a posteriori	35
3.2.2	Taux de contraction de la posterior	36
3.2.2.1	Définition	36
3.2.2.2	Lemme fondamental	37
3.2.2.3	Théorie générale	39
3.2.3	Consistance en sélection	41
3.2.3.1	Définition	41
3.2.3.2	Schéma classique de preuve	41

3.1 Sélection de variables en grande dimension

Dans ce chapitre, nous considérons le modèle de régression linéaire suivant :

$$Y = X\beta + \varepsilon, \quad \varepsilon \sim \mathcal{N}_n(0, \sigma^2 I_n), \quad (3.1)$$

avec $Y \in \mathbb{R}^n$ le vecteur des observations, $X \in \mathbb{R}^{n \times p}$ une matrice contenant la collection des p variables explicatives, appelée matrice de design, $\beta \in \mathbb{R}^p$ le vecteur des coefficients de régression, et $\sigma^2 > 0$ la variance du bruit gaussien. Supposons que l'on dispose d'un n -échantillon $\{(y_1, x_1), \dots, (y_n, x_n)\}$ généré à partir du modèle précédent pour une vraie valeur du vecteur de régression β_0 et de la variance σ_0^2 , où $y_i \in \mathbb{R}$ est une observation de la variable à expliquer et $x_i \in \mathbb{R}^p$ est un vecteur d'observations des p variables explicatives. On note :

$$Y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}, \quad X = \begin{pmatrix} x_1^\top \\ \vdots \\ x_n^\top \end{pmatrix} = \begin{pmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{pmatrix},$$

et X_ℓ la ℓ -ième colonne de X , qui correspond aux observations sur les n individus de la covariable $\ell \in \{1, \dots, p\}$. À partir de ce n -échantillon, le problème de sélection de variables consiste à identifier les coordonnées de β_0 qui sont non-nulles, ou autrement dit d'estimer le support de β_0 , noté S_0 :

$$S_0 = \left\{ \ell \in \{1, \dots, p\} \mid (\beta_0)_\ell \neq 0 \right\}.$$

Dans la suite, une variable ℓ est dite "pertinente" si $\ell \in S_0$, c'est-à-dire s'il s'agit d'une variable réellement influente.

Nous commençons par présenter les méthodes de sélection de variables fréquentistes qui utilisent une approche de pénalisation de la vraisemblance pour sélectionner les variables les plus pertinentes tout en atténuant le risque de sur-ajustement. Cette première section 3.1.1 est principalement inspirée de Giraud (2021). Ensuite, dans un cadre bayésien, la loi a priori, ou prior, revêt une importance cruciale en permettant naturellement une régularisation (Bickel et al., 2006). Ainsi, dans un second volet, nous introduisons les priors qui induisent la parcimonie dans le vecteur de régression, en mettant particulièrement l'accent sur les priors de type *spike-and-slab* qui seront explorés tout au long de cette thèse. Cette seconde section 3.1.2 est principalement inspirée de Tadesse and Vannucci (2021).

3.1.1 Méthodes pénalisées fréquentistes

Supposons la variance résiduelle σ^2 connue pour simplifier, et concentrons nous uniquement sur le paramètre d'intérêt de la sélection de variables : le vecteur de régression β . Classiquement en fréquentiste, il est courant d'estimer β par la méthode des moindres carrés ordinaires :

$$\hat{\beta}^{MCO} = \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} \|Y - X\beta\|_2^2.$$

Cet estimateur est en fait équivalent à l'estimateur du maximum de vraisemblance dans le cadre de la régression linéaire gaussienne. Ce problème de minimisation possède une unique solution lorsque X est de rang p , dont la forme explicite est donnée par : $\hat{\beta}^{MCO} = (X^\top X)^{-1} X^\top Y$. En effet, l'hypothèse d'inversibilité de $X^\top X \in \mathbb{R}^{p \times p}$ est nécessaire pour avoir une unique solution, et celle-ci fait défaut en grande dimension lorsque $p > n$. Ainsi, pour rendre le modèle identifiable dans un contexte de grande dimension, nous devons faire une hypothèse de régularité sur le vecteur β .

Il est important de remarquer que dans les données de grande dimension, la répartition n'est généralement pas uniforme dans l'espace $\mathbb{R}^p$, mais plutôt regroupée autour de structures de dimension inférieure. Ces structures sont dues à la complexité relativement faible des systèmes produisant les données. Cependant, ces structures sont souvent inconnues et leur identification est un défi majeur des statistiques en grande dimension. L'hypothèse classiquement considérée sur la structure des données est de supposer que la plupart des coordonnées β_ℓ du vecteur de régression β sont nulles, impliquant ainsi que la variable réponse dépend seulement d'un nombre limité de variables explicatives. Cette hypothèse de parcimonie dans β est cohérente avec le contexte biologique, où le nombre de marqueurs génétiques réellement influant pour un trait phénotypique donné est faible devant le nombre total de marqueurs génétiques de la plante.

3.1.1.1 Régularisation de la vraisemblance

Ainsi, pour intégrer cette hypothèse de parcimonie dans la sélection des variables, l'approche classiquement utilisée consiste à pénaliser les moindres carrés par une fonction convexe de β , notée F (Bühlmann and Van De Geer, 2011). Ainsi, pour un paramètre $\lambda > 0$ donné, le problème d'optimisation devient :

$$\hat{\beta}_\lambda = \underset{\beta \in \mathbb{R}^p}{\operatorname{argmin}} \{ \|Y - X\beta\|_2^2 + \lambda F(\beta) \}.$$

La résolution de ce problème est explicite, du fait de la linéarité du modèle, ce qui permet de faire des études théoriques sur cet estimateur. Cette procédure vise à trouver le meilleur compromis entre la recherche de la meilleure combinaison linéaire des variables X_ℓ expliquant Y et assurer le rétrécissement (ou *shrinkage*) des coordonnées du vecteur β . Différentes pénalités existent dans la littérature dont voici les plus connues :

- **Pénalité Lasso** (Least Absolute Shrinkage and Selection Operator, Tibshirani (1996)) : $F(\beta) = \|\beta\|_1 = \sum_{\ell=1}^p |\beta_\ell|$,
- **Pénalité Ridge** (Hoerl and Kennard, 1970) : $F(\beta) = \|\beta\|_2^2 = \sum_{\ell=1}^p |\beta_\ell|^2$,
- **Pénalité Elastic-Net** (Zou and Hastie, 2005) : $F(\beta) = \alpha \|\beta\|_1 + (1 - \alpha) \|\beta\|_2^2$, $\alpha \in]0, 1[$.

La pénalité Lasso exploite l'irrégularité de la boule pour la norme ℓ_1 , permettant de forcer certains coefficients proches de zéro à valoir exactement zéro, ce qui aboutit à la sélection de variables. Tout l'enjeu réside dans la détermination adéquate du paramètre λ , qui contrôle la taille du support de l'estimateur. En effet, lorsque $\lambda = 0$, l'estimateur $\hat{\beta}_\lambda$ correspond à l'estimateur des moindres carrés, sélectionnant toutes les variables avec un support complet. À mesure que λ augmente, le nombre de coordonnées non-nulles de β diminue, ce qui peut réduire la qualité de l'ajustement du modèle aux données. Lorsque λ devient suffisamment grand, l'estimateur de β tend vers le vecteur nul, et aucune variable n'est sélectionnée. Cependant, en présence de corrélation entre certaines variables, la pénalité Lasso a tendance à sélectionner aléatoirement une seule variable parmi elles, plutôt que de toutes les sélectionner ou de ne sélectionner aucune.

L'avantage majeur de la pénalité Ridge est la différentiabilité de la norme ℓ_2 , mais elle ne sélectionne pas de variables. En effet, même si elle réduit considérablement la valeur des coefficients, les petites valeurs des coefficients ne sont pas mises à zéro. En revanche, la pénalité Elastic-Net, une combinaison convexe des pénalités Lasso et Ridge, offre un compromis en permettant à la fois la sélection de variables grâce au Lasso, tout en conservant une certaine régularité grâce au Ridge. Ainsi, l'Elastic-Net favorise les effets de groupe : les variables corrélées sont soit toutes sélectionnées, soit aucune d'entre elles n'est sélectionnée. À noter également que ces méthodes de pénalisation présentent un biais car elles ont tendance à réduire vers zéro aussi les coefficients non-nuls (en faisant tendre λ vers l'infini) pour minimiser la fonction de coût.

D'autres variantes du Lasso ont été développées pour différents types de parcimonie (Hastie et al., 2015; Giraud, 2021), telles que *Group-Lasso* pour la parcimonie par groupe de variables (Yuan and Lin, 2006) ou *Fused-Lasso* pour les situations où nous souhaitons limiter les variations dans les coordonnées du vecteur de régression (Tibshirani et al., 2005). Une synthèse des méthodes de pénalisation pour la sélection de variables en grande dimension est proposée par Fan and Lv (2010), comprenant des approches telles que SCAD (*Smoothly Clipped Absolute Deviation*) basé sur une pénalisation non convexe (Fan and Li, 2001) ou *Adaptive Lasso* qui utilise une pénalité ℓ_1 pondérée (Zou, 2006), par exemple.

3.1.1.2 Sélection de modèles

En pratique, cette procédure de vraisemblance pénalisée est réalisée pour une grille de valeurs finies de λ , notée Λ . Nous obtenons ainsi une collection d'estimateurs $(\hat{\beta}_\lambda)_{\lambda \in \Lambda}$ de support respectif $(\hat{S}_\lambda)_{\lambda \in \Lambda}$, et on sélectionne le "meilleur" λ par un critère de sélection de modèle :

$$\hat{\lambda} = \operatorname{argmin}_{\lambda \in \Lambda} \left\{ -2 \log \left(p(Y; \tilde{\beta}_\lambda) \right) + \operatorname{pen}(|\hat{S}_\lambda|) \right\},$$

avec :

- $\log(p(Y; \beta))$ la log-vraisemblance du modèle,
- $\tilde{\beta}_\lambda$ l'estimateur des moindres carrés dans le modèle défini par le support $\hat{S}_\lambda$:

$$\tilde{\beta}_\lambda = \operatorname{argmin}_{\beta \in \mathbb{R}^p : \operatorname{supp}(\beta) = \hat{S}_\lambda} \left\{ \|Y - X\beta\|_2^2 \right\},$$

avec $\operatorname{supp}(\beta) = \left\{ \ell \in \{1, \dots, p\} \mid \beta_\ell \neq 0 \right\}$ le support de β ,

- $\text{pen}(|\hat{S}_\lambda|)$ une fonction de pénalité sur la taille du support $|\hat{S}_\lambda|$.

Le support sélectionné est alors $\hat{S}_{\hat{\lambda}}$. Il existe différentes pénalités comme :

- **Pénalité AIC** (*Akaike Information Criterion*, Akaike (1973)) :

$$\text{pen}_{\text{AIC}}(|\hat{S}_\lambda|) = 2|\hat{S}_\lambda|,$$

- **Pénalité BIC** (*Bayesian Information Criterion*, Schwarz (1978)) :

$$\text{pen}_{\text{BIC}}(|\hat{S}_\lambda|) = |\hat{S}_\lambda| \log(n),$$

- **Pénalité eBIC** (*extended BIC*, Chen and Chen (2008)) :

$$\text{pen}_{\text{eBIC}}(|\hat{S}_\lambda|) = |\hat{S}_\lambda| \log(n) + 2\gamma \log \left(\binom{p}{|\hat{S}_\lambda|} \right),$$

avec $\gamma \in [0, 1]$.

En grande dimension, les critères AIC et BIC ont tendance à sélectionner un nombre trop important de variables et tombent dans le problème de sur-apprentissage. L'une des raisons est l'absence de prise en compte du fait que le nombre de supports avec m variables augmente très vite quand m augmente : il y a beaucoup plus de supports avec 10 variables (il y en a $\binom{p}{10}$) que de supports avec une seule variable (il y en a p). Ainsi, là où le critère BIC suppose une loi a priori uniforme sur tous les supports, le critère eBIC suppose plutôt :

$$\pi(\hat{S}_\lambda) = \binom{p}{|\hat{S}_\lambda|}^{-\gamma},$$

ce qui permet de prendre en compte la complexité du support. Notamment, Chen and Chen (2008) montrent que le critère eBIC entraîne une légère perte dans le taux de vrais positifs par rapport au critère BIC, mais il contrôle plus fortement le taux de faux positifs, une propriété souhaitable dans de nombreuses applications, notamment en génétique. Plus précisément, ils ont montré que pour $\gamma = 1$ et p de l'ordre d'une puissance de n , la probabilité de sélectionner un mauvais support tend vers zéro quand n tend vers l'infini.

3.1.2 Priors bayésiens induisant la parcimonie

En réalité, la plupart des techniques de régularisation développées initialement dans un contexte fréquentiste ont également été proposées dans un cadre bayésien (Casella et al., 2010). Ces adaptations sont exposées dans cette section, ainsi que les principaux priors bayésiens favorisant la parcimonie.

Comme pour la section précédente, nous introduisons les priors bayésiens dans le cas du modèle de régression linéaire (3.1), où l'on note $\theta = (\beta, \sigma^2)$ l'ensemble des paramètres à estimer. Le paradigme bayésien considère les paramètres d'un modèle non plus comme des quantités fixes inconnues mais comme des variables aléatoires auxquelles on associe une loi dite a priori, ou prior. Le choix de cette distribution a priori et des hyperparamètres associés (*i.e.* les paramètres du prior) peut être influencé par les connaissances expertes initiales sur ces paramètres avant de

voir les données. En combinant l'information a priori avec les données observées Y à l'aide du théorème de Bayes, nous obtenons la distribution a posteriori des paramètres θ , ou posterior, qui représente notre connaissance mise à jour après avoir vu les données :

$$\pi(\theta|Y) = \frac{p(Y|\theta)\pi(\theta)}{p(Y)},$$

où $p(Y|\theta)$ est la vraisemblance des observations Y , $\pi(\theta)$ est le prior sur les paramètres θ , et $p(Y) = \int p(Y|\theta)\pi(\theta)d\theta$ est la densité marginale de Y . Cette distribution a posteriori $\pi(\theta|Y)$ est utilisée pour estimer les paramètres, en choisissant comme estimateur par exemple la moyenne ou le maximum de cette posterior. En outre, l'une des raisons de l'attrait des méthodes d'inférence bayésiennes est leur capacité intrinsèque à quantifier l'incertitude une fois que la distribution a posteriori a été calculée puisque cette dernière contient toutes les informations nécessaires pour répondre à diverses questions statistiques.

Cette section est consacrée dans un premier temps à l'analyse de l'extension des méthodes de régularisation fréquentistes dans le cadre bayésien. Ensuite, nous présenterons les priors de type *spike-and-slab*, une approche largement adoptée en statistique bayésienne. Enfin, nous aborderons d'autres formes de priors bayésiens qui favorisent la parcimonie dans le vecteur de régression.

3.1.2.1 Extension de la régularisation fréquentiste

Le choix du prior est essentiel en grande dimension car son influence ne disparaît pas complètement asymptotiquement (Rousseau, 2016), et va induire naturellement une certaine forme de régularisation (Bickel et al., 2006). Plus précisément, les priors agissent comme des termes de pénalité dans l'approche fréquentiste. Prenons l'exemple du Lasso bayésien, ou *Bayesian Lasso*, introduit par Park and Casella (2008). Ils ont proposé un échantillonnage de Gibbs pour le Lasso avec le prior Laplace dans le modèle hiérarchique. Plus précisément, ils ont envisagé une analyse entièrement bayésienne utilisant un prior conditionnel Laplace de la forme :

$$\pi(\beta|\sigma^2) = \prod_{\ell=1}^p \frac{\lambda}{2\sqrt{\sigma^2}} e^{-\lambda|\beta_\ell|/\sqrt{\sigma^2}}, \quad \pi(\sigma^2) = \frac{1}{\sigma^2}.$$

Ils ont souligné que le conditionnement par rapport à σ^2 est important puisqu'il garantit un posterior complet unimodal. L'absence d'unimodalité ralentit la convergence de l'échantillonneur de Gibbs et rend les estimations ponctuelles moins significatives. Sous ce modèle, on peut montrer que le logarithme de la distribution a posteriori de β est, à une constante additive indépendante de β près :

$$-\frac{1}{2\sigma^2}\|Y - X\beta\|_2^2 - \frac{\lambda}{\sigma}\|\beta\|_1.$$

Donc, pour toutes valeurs de σ^2 et λ fixées, le maximum de la distribution a posteriori (estimateur MAP) de β coïncide avec l'estimateur Lasso (avec un paramètre de régularisation de $2\sigma\lambda$). De même, on peut retrouver les régularisations Ridge et Elastic-Net en utilisant respectivement un prior gaussien sur β ou un compromis entre les priors Laplace et gaussien (Zou and Hastie, 2005). À noter cependant que les estimateurs associés au Lasso bayésien ne sont pas parcimonieux (Park and Casella, 2008). Il en est de même pour les versions bayésiennes du Ridge et d'Elastic-Net.

3.1.2.2 Les priors de mélange *spike-and-slab*

Ainsi, en voyant les priors comme une régularisation automatique, il est naturel d'assigner un prior différent pour les coefficients nuls et non-nuls de β afin d'induire la parcimonie dans ce vecteur mais tout en évitant de réduire vers zéro les coefficients non-nuls. Les vrais coefficients non-nuls dans le vecteur β sont appelés coefficients de signal, tandis que les autres sont appelés coefficients de nuisance. Ainsi, pour offrir un contrôle séparé des coefficients de signal et de nuisance, les priors de mélange de type *spike-and-slab* ont été introduit par [Mitchell and Beauchamp \(1988\)](#), puis étudiés par [George and McCulloch \(1993, 1997\)](#). Pour les décrire, nous introduisons un vecteur latent de variables indicatrices binaires $\delta = (\delta_1, \dots, \delta_p)$, qui indexe les 2^p sous-ensembles possibles de variables pertinentes, avec :

$$\delta_\ell = \begin{cases} 1 & \text{si la covariable } \ell \text{ est pertinente,} \\ 0 & \text{sinon.} \end{cases}$$

Ce vecteur binaire correspond alors exactement au support du vecteur β . Ainsi, un modèle hiérarchique bayésien où un prior est assigné au vecteur δ suivi d'un prior sur le vecteur de régression β conditionnellement à δ permet à la fois de réaliser l'estimation de β , en utilisant sa distribution a posteriori sachant les données $\pi(\beta|Y)$, mais aussi la sélection de variables, en maximisant la distribution a posteriori du vecteur binaire δ sachant les données $\pi(\delta|Y)$. Le prior de mélange *spike-and-slab* peut s'écrire sous la forme générique suivante pour la régression linéaire (3.1) ([Tadesse and Vannucci, 2021](#)) :

$$\begin{aligned} \pi(\beta|\delta, \sigma^2) &= \prod_{\ell=1}^p \left[(1 - \delta_\ell) \phi_0(\beta_\ell|\sigma^2) + \delta_\ell \phi_1(\beta_\ell|\sigma^2) \right], \\ \pi(\delta|\alpha) &= \prod_{\ell=1}^p \alpha^{\delta_\ell} (1 - \alpha)^{1-\delta_\ell}, \\ \alpha &\sim \pi(\alpha), \\ \sigma^2 &\sim \pi(\sigma^2), \end{aligned}$$

où ϕ_0 est une densité très concentrée autour de zéro utilisée pour modéliser les coefficients de nuisance, appelée distribution "spike", ϕ_1 est une densité continue diffuse autorisant des valeurs intermédiaires et grandes de β_ℓ pour modéliser les coefficients de signal, appelée distribution "slab". Une représentation des distributions *spike* et *slab* dans le cas d'un mélange de gaussiennes est proposée sur la figure 3.1. Le paramètre $\alpha \in [0, 1]$ correspond à une proportion de mélange, ou plus précisément à la proportion de covariables pertinentes dans le modèle, auquel on associe un prior $\pi(\alpha)$. Pour favoriser la parcimonie dans un contexte de grande dimension, le choix du prior sur α est donc primordial. Classiquement, en suivant la recommandation de [Castillo and van der Vaart \(2012\)](#), on choisit une distribution Beta $\pi(\alpha) \propto \alpha^{a-1} (1 - \alpha)^{b-1}$ avec $a > 0$ petit et $b > 0$ grand pour lui donner une variance faible. Pour la variance résiduelle σ^2 , on choisit typiquement un prior conjugué inverse-gamma ou un prior impropre de Jeffreys $\pi(\sigma^2) \propto \sigma^{-2}$.

On peut distinguer deux classes de prior *spike-and-slab* en fonction du choix de la distribution *spike* ϕ_0 : les priors *spike-and-slab* discrets (ou *discrete spike-and-slab*) qui utilisent une masse de Dirac en zéro ([Mitchell and Beauchamp, 1988](#)), et les priors *spike-and-slab* continus (ou *continuous spike-and-slab*) qui choisissent une distribution continue mais très concentrée autour

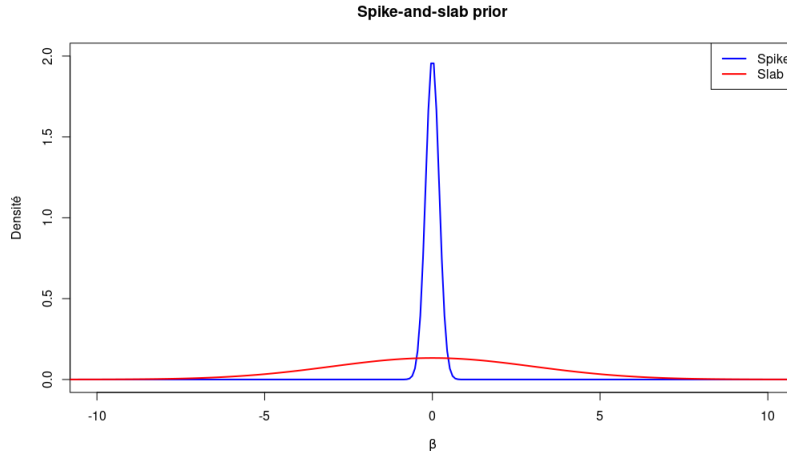


Figure 3.1: Illustration des distributions *spike* et *slab* dans un cas gaussien continu.

de zéro (George and McCulloch, 1993). Le choix entre l'une ou l'autre approche implique des compromis. Les priors *spike-and-slab* discrets offrent une représentation correcte du problème de parcimonie en attribuant vraiment la valeur zéro aux coefficients de nuisance. De ce point de vue, ils sont souvent considérés comme un idéal théorique dans la littérature (Carvalho et al., 2009). Cependant, ils rencontrent des difficultés d'un point de vue algorithmique, notamment liées à l'exploration de la posterior complète sur l'ensemble des modèles de covariables possibles et à la complexité combinatoire de mise à jour des variables indicatrices, ce qui peut-être computationnellement prohibitif en grande dimension. Pour pallier ce problème, des stratégies ont été développées pour identifier rapidement des régions de forte probabilité a posteriori telles que *Shotgun Stochastic Search* (Hans et al., 2007), ou basées sur des approches variationnelles comme le *Variational Bayes* (Ray and Szabó, 2022) qui reformulent l'approximation de la posterior en un problème d'optimisation. En parallèle, l'alternative qui consiste à choisir une distribution *spike* continue adopte le point de vue que les coefficients de nuisance peuvent être négligeables sans pour autant être exactement égaux à zéro. Par exemple, dans le contexte de la sélection génétique, l'utilisation de ce prior offre plus de flexibilité ainsi qu'une représentation plus fidèle de la réalité. En effet, il est raisonnable de postuler que de nombreux gènes exercent un effet faible mais non nul sur le caractère d'intérêt. Cependant, notre objectif principal demeure l'identification et l'isolement des gènes qui ont un impact significatif sur ce caractère spécifique. Ainsi, les priors *spike-and-slab* continus peuvent être intéressants d'un point de vue algorithmique comparés au cas discret, mais les estimations ne sont pas véritablement parcimonieuses. La sélection de variables doit donc être opérée de manière "artificielle", typiquement grâce à l'application de techniques de seuillage des probabilités d'inclusion des covariables. En effet, il est plus avantageux d'utiliser les probabilités marginales a posteriori $\pi(\delta_\ell = 1|Y)$, pour $\ell = 1, \dots, p$, car nous avons besoin de moins d'itérations de Gibbs pour estimer seulement ces p posterior, au lieu des 2^p posterior $\pi(\delta = \mathbf{m}|Y)$ pour toutes combinaisons de variables possibles $\mathbf{m} \in \{0, 1\}^p$ (Tadesse and Vannucci, 2021). En pratique, la règle du seuillage médian est la plus souvent utilisée (*median probability model*, Barbieri and Berger (2004)) : $\pi(\delta_\ell = 1|Y) > 0.5$.

Donnons quelques exemples de prior de mélange *spike-and-slab* développés dans la littéra-

ture dans le modèle de régression linéaire gaussienne (3.1). Pour le cas des priors *spike-and-slab* discrets, Geweke (1996) propose d'utiliser une loi gaussienne (potentiellement tronquée) comme distribution *slab* et il utilise un échantillonneur de Gibbs pour construire les moments de la posterior. Castillo et al. (2015) se concentrent sur une loi Laplace comme distribution *slab* et montrent qu'un tel prior donne de bon résultats à la fois en termes d'estimation, de prédiction mais aussi de sélection. Par ailleurs, George and McCulloch (1993) sont les premiers à remplacer la distribution Dirac par un prior *spike* continu. Plus précisément, ils utilisent comme distributions *spike* et *slab* des densités gaussiennes avec respectivement une petite et une grande variance :

$$\pi(\beta|\delta, \sigma^2) = \prod_{\ell=1}^p [(1 - \delta_\ell)\mathcal{N}(0, \sigma^2\tau_0^2) + \delta_\ell\mathcal{N}(0, \sigma^2\tau_1^2)],$$

avec $0 < \tau_0^2 \ll \tau_1^2$. Ils développent une procédure appelée *Stochastic Search Variable Selection* (SSVS) basée sur un échantillonneur de Gibbs et le seuillage des probabilités d'inclusion a posteriori $\pi(\delta_\ell = 1|Y)$, $\ell = 1, \dots, p$. Narisetty and He (2014) ont montré que ce prior *spike-and-slab* gaussien atteint de bonnes performances en sélection de variables en grande dimension si les variances τ_0^2 et τ_1^2 dépendent de la taille de l'échantillon n de manière appropriée. Plus récemment, Ročková and George (2014) ont proposé une alternative déterministe à la recherche stochastique SSVS, appelée *Expectation Maximisation Variable Selection* (EMVS), basée sur l'algorithme EM et permettant d'identifier rapidement les ensembles de variables de grande probabilité a posteriori sous le prior *spike-and-slab* gaussien. Pour cela, ils se concentrent uniquement sur le mode de la distribution a posteriori pour estimer leurs paramètres grâce à l'algorithme EM. On note ce mode $(\hat{\beta}, \hat{\alpha}, \hat{\sigma}^2)$. Ils opèrent ensuite la sélection de variables par une version locale du modèle de probabilité médiane de Barbieri and Berger (2004) : $\pi(\delta_\ell = 1|\hat{\beta}, \hat{\alpha}, \hat{\sigma}^2) > 0.5$. À noter qu'en transformant le problème d'échantillonnage des posterior en un problème d'optimisation, Ročková and George (2014) ont drastiquement réduit les coûts computationnels. Pour finir avec les priors *spike-and-slab* continus, on trouve aussi dans la littérature des mélanges de loi Laplace (Ročková and George, 2018) : plus précisément on choisit comme distribution *spike* (respectivement *slab*) une loi Laplace de paramètre de régularisation λ_0 (respectivement λ_1), avec typiquement $\lambda_0 \gg \lambda_1$. Ce prior est appelé *Spike-and-Slab Lasso* car pour $\lambda_0 = \lambda_1$ on retrouve la pénalité classique Lasso. Les performances théoriques de tous ces types de prior *spike-and-slab* sont résumés dans le livre Tadesse and Vannucci (2021) et seront plus largement discutées dans la section 4.2.1 du chapitre suivant.

3.1.2.3 Autres formes de priors favorisant la parcimonie

Bien que les priors *spike-and-slab* soient au coeur des préoccupations principales de cette thèse, il convient de noter l'existence d'autres formes de priors bayésiens capables d'induire la parcimonie : les priors de rétrécissement global-local, ou *global-local shrinkage* prior, prenant la forme générique suivante (Tadesse and Vannucci, 2021) :

$$\beta_\ell|\lambda_\ell^2, \tau^2, \sigma^2 \sim \mathcal{N}(0, \sigma^2\tau^2\lambda_\ell^2), \quad \lambda_\ell^2 \sim \pi(\lambda_\ell^2), \quad \ell = 1, \dots, p, \quad (3.2)$$

où π est un prior spécifique qui détermine les différentes variétés de prior de rétrécissement global-local. Par exemple, le prior *Horseshoe* (Carvalho et al., 2010) correspond au prior *half-Cauchy* pour π . Il convient de noter que les priors *spike-and-slab* gaussiens sont aussi un cas particulier de ce type de prior en choisissant : $\pi(\lambda_\ell^2) = \delta_\ell\delta_{\lambda_\ell^2=\tau_0^2} + (1 - \delta_\ell)\delta_{\lambda_\ell^2=\tau_1^2}$ et $\tau^2 = 1$. Le

paramètre τ^2 , commun à tous les β_ℓ , détermine le niveau global de parcimonie, tandis que les paramètres $(\lambda_\ell^2)_{1 \leq \ell \leq p}$ sont des composantes locales qui permettent soit d'atténuer la réduction vers zéro lorsque λ_ℓ^2 est grand, soit l'amplifier lorsque λ_ℓ^2 est proche de zéro. Des priors sur les paramètres τ et σ^2 peuvent aussi être considérés. Les garanties théoriques et les procédures MCMC pour ces priors de rétrécissement global-local sont résumés dans [Tadesse and Vannucci \(2021\)](#).

3.2 Analyse fréquentiste des distributions a posteriori

Dans cette thèse, nous nous concentrons également sur l'étude théorique de l'utilisation du prior *spike-and-slab* pour la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Cette section présente alors les outils nécessaires pour comprendre le comportement des distributions a posteriori bayésiennes dans un contexte de grande dimension.

Une compréhension approfondie des propriétés théoriques des posteriors est essentielle pour justifier la pertinence des algorithmes et méthodologies employés en pratique, ainsi que pour orienter le choix des priors, certains étant démontrés comme sous-optimaux ou moins appropriés que d'autres. Nous mettrons l'accent sur ce qu'on appelle *l'analyse fréquentiste des distributions a posteriori*. Tout d'abord, la section 3.2.1 introduit les notations et le paradigme fréquentiste pour la convergence des distributions a posteriori bayésiennes. La notion de taux de contraction de la posterior et la théorie générale associée sont introduites dans la section 3.2.2. Ces deux premières sections sont principalement inspirées de [Castillo \(2024\)](#), et de [Ghosal and Van der Vaart \(2017\)](#). Ensuite, la section 3.2.3 présente la propriété de consistance en sélection de la posterior et un schéma classique de démonstration.

3.2.1 Paradigme fréquentiste pour la convergence des distributions a posteriori

Considérons un modèle statistique $\mathcal{P} = \{P_\theta^{(n)}, \theta \in \Theta\}$, indexé par un paramètre θ qui est soit de dimension infinie, typiquement $\theta = f$ une fonction, soit de grande dimension, typiquement $\dim(\theta) \gg n$, avec $n \geq 1$ une mesure de l'information des données. Dans ce dernier cas, le modèle est paramétrique pour chaque valeur de n fixée, mais la dimension du paramètre augmente avec n . On munit Θ d'une (semi-)métrique d_n , et on note $Y^{(n)}$ des observations de ce modèle. Le modèle bayésien associé s'écrit :

$$\theta \sim \Pi, \tag{3.3}$$

$$Y^{(n)} | \theta \sim P_\theta^{(n)}, \tag{3.4}$$

où Π est le prior sur l'espace mesurable des paramètres $(\Theta, \mathcal{B})$. Supposons que pour tout $\theta \in \Theta$, pour tout $n \geq 1$, $P_\theta^{(n)}$ est absolument continu par rapport à $\mu^{(n)}$, pour $\mu^{(n)}$ une mesure σ -finie, et notons $p_\theta^{(n)}(Y^{(n)}) = dP_\theta^{(n)}(Y^{(n)})/d\mu^{(n)}$ la vraisemblance des observations. On note $\mathbb{E}_\theta^{(n)}$ l'espérance sous la loi $P_\theta^{(n)}$. Ainsi, la distribution a posteriori, qui quantifie l'incertitude restante

sur le paramètre après l'observation des données, s'obtient par la formule de Bayes :

$$\Pi(B|Y^{(n)}) := \frac{\int_B p_{\theta}^{(n)}(Y^{(n)}) d\Pi(\theta)}{\int_{\Theta} p_{\theta}^{(n)}(Y^{(n)}) d\Pi(\theta)}, \quad \text{pour tout } B \in \mathcal{B}. \quad (3.5)$$

Une critique fréquente à l'égard de l'approche bayésienne réside dans la subjectivité du choix du prior (Robert, 2006). En effet, bien que le paradigme bayésien offre une méthode naturelle pour incorporer des connaissances a priori concrètes sur les paramètres, cette approche devient problématique lorsque ces connaissances sont vagues ou incomplètes. Dans de telles situations, la spécification d'un prior justifié devient difficile, ce qui peut conduire à des résultats divergents entre deux statisticiens ayant choisi des priors différents. En effet, la complexité de l'ensemble des paramètres Θ implique que l'influence du prior est important et ne disparaît pas complètement asymptotiquement (Diaconis and Freedman, 1986; Rousseau, 2016). Pour répondre à cette préoccupation, plutôt que de considérer les observations $Y^{(n)}$ comme provenant de la distribution marginale de $Y^{(n)}$ résultant de (3.3)-(3.4), nous faisons l'hypothèse qu'il existe une "vraie" valeur θ_0 du paramètre et que les données observées suivent $Y^{(n)} \sim P_{\theta_0}^{(n)}$ comme en fréquentiste. Ensuite, l'idée est d'étudier la posterior (3.5), $\Pi(\cdot|Y^{(n)})$, en probabilité sous cette hypothèse, ce qui constitue l'essence de l'analyse fréquentiste des distributions a posteriori (Rousseau, 2016; Castillo, 2024). Son objectif est de trouver des conditions sur le modèle telles que, avec un nombre suffisant de données n , la suite des distributions a posteriori se concentre, dans un sens qui sera clarifié ci-après, autour du vrai θ_0 . Dans la suite, nous introduisons les différentes propriétés asymptotiques fréquentistes de la distribution a posteriori qui sont particulièrement intéressantes dans le cadre de la sélection de variables en grande dimension.

3.2.2 Taux de contraction de la posterior

Pour rappel, nous supposons dans la suite que les observations $Y^{(n)}$ sont distribuées selon $P_{\theta_0}^{(n)}$, avec θ_0 fixé. On note $p_{\theta_0}^{(n)}(Y^{(n)}) = dP_{\theta_0}^{(n)}(Y^{(n)})/d\mu^{(n)}$ la densité associée. On peut alors réécrire la distribution a posteriori (3.5) comme suit :

$$\Pi(B|Y^{(n)}) := \frac{\int_B \frac{p_{\theta}^{(n)}}{p_{\theta_0}^{(n)}}(Y^{(n)}) d\Pi(\theta)}{\int_{\Theta} \frac{p_{\theta}^{(n)}}{p_{\theta_0}^{(n)}}(Y^{(n)}) d\Pi(\theta)}, \quad \forall B \in \mathcal{B}. \quad (3.6)$$

3.2.2.1 Définition

Afin de valider une méthode d'estimation bayésienne d'un point de vue fréquentiste, il est essentiel que la distribution a posteriori attribue une proportion significative de sa masse aux éléments de l'espace paramétrique Θ proches de la vraie valeur du paramètre θ_0 . La définition suivante quantifie cet propriété comme énoncé dans Ghosal and Van der Vaart (2017).

Définition 1 (Taux de contraction). *La distribution a posteriori se concentre au taux ϵ_n autour de θ_0 pour une certaine métrique $d_n(\cdot, \cdot)$ définie sur Θ s'il existe une constante $M > 0$ telle que :*

$$\mathbb{E}_{\theta_0}^{(n)} \left[\Pi \left(\theta : d_n(\theta, \theta_0) > M\epsilon_n \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

La suite ϵ_n (souvent tendant vers zéro quand n tend vers l'infini) est alors appelée taux de contraction de la distribution a posteriori.

En d'autres termes, cette définition signifie essentiellement que la distribution a posteriori concentre la plupart de sa masse dans une boule centrée autour de θ_0 , avec un rayon de plus en plus petit $M\epsilon_n$, à condition que ϵ_n tende vers zéro à mesure que n augmente. Ainsi, à mesure que la taille de l'échantillon n augmente, les informations provenant des données devraient prévaloir sur l'impact du prior. Il est intéressant de noter que la contraction de la posterior au taux ϵ_n implique que les estimateurs ponctuels bayésiens, tels que la moyenne ou la médiane a posteriori, ont une vitesse de convergence au moins aussi rapide que ϵ_n au sens fréquentiste (Ghosal et al., 2000). Par ailleurs, ϵ_n dépend typiquement des caractéristiques de θ_0 et des propriétés du prior (Rousseau, 2016). Par exemple, dans un contexte de sélection de variables en grande dimension, ϵ_n va dépendre de la taille du support de θ_0 . Étant donné que la vraie valeur du paramètre θ_0 est inconnue, notre objectif est d'obtenir des taux de contraction de la distribution a posteriori qui sont uniformes sur tout l'espace des paramètres $\theta_0 \in \Theta$, ou restreints à un sous-ensemble $\Theta_0 \subset \Theta$ dans les cas complexes :

$$\sup_{\theta_0 \in \Theta_0} \mathbb{E}_{\theta_0}^{(n)} \left[\Pi \left(\theta : d(\theta, \theta_0) > M\epsilon_n \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Typiquement, lorsque la variance résiduelle n'est pas connue, les résultats de contraction de posterior obtenus dans la littérature sont restreints à une certaine région de l'ensemble des paramètres $\Theta_0 \subset \Theta$ (Ning et al., 2020; Jeong and Ghosal, 2021b).

Un taux de contraction de la posterior sera dit "optimal" s'il coïncide avec le taux minimax :

$$\bar{\epsilon}_n = \inf_T \sup_{\theta \in \mathcal{R}} \mathbb{E}_{\theta}^{(n)} [d_n(T, \theta)],$$

où l'infimum est calculé sur tous les estimateurs possibles $T = T(Y^{(n)})$ de θ , et θ_0 est supposé appartenir à un ensemble régulier $\mathcal{R}$.

3.2.2.2 Lemme fondamental

Au premier abord, l'obtention de taux de contraction peut sembler difficile lorsque la distribution a posteriori ne possède pas une forme explicite. Cependant, grâce aux travaux de Ghosal et al. (2000); Ghosal and Van Der Vaart (2007), également exposés dans l'ouvrage de référence Ghosal and Van der Vaart (2017), une théorie générale a émergé, fournissant des conditions suffisantes pour obtenir ce type de résultats. Avant d'aborder cette théorie, nous présentons un lemme essentiel pour contrôler la distribution a posteriori. Pour ce faire, des définitions supplémentaires s'avèrent nécessaires.

Définition 2 (Divergence et variation de Kullback-Leibler). *Pour deux densités f et g , on note*

$$K(f, g) = \int f(x) \log(f(x)/g(x)) dx,$$

la divergence de Kullback-Leibler (KL), et

$$V(f, g) = \int f(x) |\log(f(x)/g(x)) - K(f, g)|^2 dx,$$

la variation de Kullback-Leibler.

Définition 3 (Voisinage de type KL). *Pour tout $\epsilon > 0$, on définit :*

$$B_n^{KL}(\theta_0, \epsilon) := \left\{ \theta \in \Theta : K(p_{\theta_0}^{(n)}, p_{\theta}^{(n)}) \leq n\epsilon^2, V(p_{\theta_0}^{(n)}, p_{\theta}^{(n)}) \leq n\epsilon^2 \right\},$$

un voisinage de type Kullback-Leibler autour de la vraie valeur $\theta_0 \in \Theta$.

Cette théorie générale s'appuie sur le lemme fondamental suivant, appelé *Evidence lower bound* dans la littérature (voir Lemme 8.21 de Ghosal and Van der Vaart (2017), ou encore Lemme 1.1 de Castillo (2024) par exemple), qui permet de minorer la quantité du dénominateur dans la formule de Bayes de la distribution a posteriori (3.6).

Lemme 1. *Pour toute distribution de probabilité Π sur Θ , pour toutes constantes $C > 0$ et $\epsilon > 0$, on a :*

$$P_{\theta_0}^{(n)} \left(\int_{\Theta} \frac{p_{\theta}^{(n)}}{p_{\theta_0}^{(n)}} (Y^{(n)}) d\Pi(\theta) \leq \Pi(B_n^{KL}(\theta_0, \epsilon)) e^{-(1+C)n\epsilon^2} \right) \leq \frac{1}{C^2 n \epsilon^2}.$$

Comme discuté notamment dans Castillo (2024), le lemme 1 implique que si la masse a priori d'un ensemble est très petite, alors sa masse a posteriori l'est aussi. Plus précisément, nous avons le corollaire suivant.

Corollaire 1. *Soient une suite $(\epsilon_n)_n$, vérifiant que $n\epsilon_n^2$ tend vers l'infini quand n tend vers l'infini, et une suite d'ensembles mesurables $(\mathcal{A}_n)_n$ telles que,*

$$\frac{\Pi(\mathcal{A}_n)}{e^{-2n\epsilon_n^2} \Pi(B_n^{KL}(\theta_0, \epsilon_n))} \xrightarrow{n \rightarrow \infty} 0.$$

Alors, on a

$$\mathbb{E}_{\theta_0}^{(n)} \left[\Pi(\mathcal{A}_n | Y^{(n)}) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Ainsi, en assurant un contrôle approprié la masse a priori de $\mathcal{A}_n$ par rapport à celle d'un voisinage de type KL de θ_0 , la masse a posteriori de $\mathcal{A}_n$ reste contrôlée.

3.2.2.3 Théorie générale

Cette section introduit la théorie générale de Ghosal et al. (2000); Ghosal and Van Der Vaart (2007) pour obtenir des taux de contraction a posteriori, basée sur le construction de tests exponentiellement puissants. Rappelons qu'un test est une fonction mesurable $\phi_n : \mathcal{Y}^n \rightarrow \{0, 1\}$. Dans la suite, on note $\Theta_n^c = \Theta \setminus \Theta_n$ le complémentaire d'un sous-ensemble Θ_n de Θ .

Théorème 1 (Contraction de posterior). *Soit $(\epsilon_n)_n$ une suite telle que $n\epsilon_n^2 \rightarrow \infty$ quand $n \rightarrow \infty$. Supposons qu'il existe une constante $C > 0$ telle que :*

(a) *Condition de Kullback-Leibler :*

$$\Pi(B_n^{KL}(\theta_0, \epsilon_n)) \geq e^{-Cn\epsilon_n^2}.$$

(b) *Condition de tests : il existe $M > 0$, $\Theta_n \subset \Theta$ et une suite de fonctions de test $\phi_n = \phi_n(Y^{(n)})$ tels que :*

$$(i) \quad \Pi(\Theta_n^c) \leq e^{-(C+4)n\epsilon_n^2},$$

$$(ii) \quad \mathbb{E}_{\theta_0}^{(n)}[\phi_n] \xrightarrow{n \rightarrow \infty} 0, \quad \sup_{\theta \in \Theta_n : d_n(\theta, \theta_0) > M\epsilon_n} \mathbb{E}_{\theta}^{(n)}[1 - \phi_n] \leq e^{-(C+4)n\epsilon_n^2}.$$

Alors, la distribution a posteriori se concentre au taux ϵ_n autour de θ_0 :

$$\mathbb{E}_{\theta_0}^{(n)} \left[\Pi \left(\theta : d_n(\theta, \theta_0) > M\epsilon_n \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

La première hypothèse de ce théorème semble assez naturelle. En effet, l'hypothèse (a) stipule que pour avoir concentration de la distribution a posteriori autour de θ_0 , il faut que la loi a priori ne soit pas trop mauvaise au début, c'est-à-dire qu'elle place une masse suffisante autour de la vraie valeur θ_0 . C'est pour satisfaire cette condition de masse qu'il est nécessaire, dans certaines situations, de contraindre le vrai paramètre à appartenir à un ensemble borné $\Theta_0 \subset \Theta$ (Ning et al., 2020). Concernant la condition de tests (b), l'hypothèse (i) permet de se concentrer par la suite sur un sous-ensemble $\Theta_n \subset \Theta$, ce qui est particulièrement intéressant lorsque Θ est complexe et/ou de grande dimension. En effet, en combinant les hypothèses (a) et (b)(i) nous obtenons que :

$$\frac{\Pi(\Theta_n^c)}{\Pi(B_n^{KL}(\theta_0, \epsilon_n))} \leq e^{-4n\epsilon_n^2},$$

et donc par le Corollaire 1 on a que $\mathbb{E}_{\theta_0}^{(n)} [\Pi(\Theta_n^c | Y^{(n)})] \xrightarrow{n \rightarrow \infty} 0$. La condition (b)(ii) est plus difficile à interpréter et la structure du test est atypique : on teste un point contre le complémentaire d'une boule, au lieu de tester contre une boule. Une méthode couramment utilisée pour construire de tels tests consiste à construire des fonctions de tests "locales" $\varphi_n^{\theta_1}$ pour comparer la vraie valeur du paramètre θ_0 contre une boule centrée en un certain $\theta_1 \in \Theta_n$ suffisamment éloigné de θ_0 : $H_0 : \theta = \theta_0$ contre $H_1 : d_n(\theta, \theta_1) \leq a\epsilon_n$, avec $0 < a < 1$ et $d_n(\theta_0, \theta_1) \geq \epsilon_n$ satisfaisant :

$$\mathbb{E}_{\theta_0}^{(n)}[\varphi_n^{\theta_1}] \leq e^{-Kn\epsilon_n^2}, \quad \sup_{\theta \in \Theta_n : d_n(\theta, \theta_1) < a\epsilon_n} \mathbb{E}_{\theta}^{(n)}[1 - \varphi_n^{\theta_1}] \leq e^{-Kn\epsilon_n^2}, \quad (3.7)$$

pour une certaine constante $K > 0$. Par exemple, l'existence de tels tests a été démontrée pour l'estimation de densité pour la distance d'Hellinger et la distance $\mathbb{L}^1$ (Le Cam, 2012), et en

régression gaussienne pour la distance ℓ_2 (Birgé, 2006). Le lemme suivant (Castillo, 2024), basé sur le lemme D.5 de Ghosal and Van der Vaart (2017), permet d'utiliser ces tests locaux pour satisfaire la condition (b)(ii) du Théorème 1.

Définition 4 (Nombre minimal de recouvrement et entropie). *Le nombre minimal de boules de rayon ϵ pour la distance d nécessaire pour recouvrir un ensemble $\mathcal{A}$ est noté $N(\epsilon, \mathcal{A}, d)$. On définit l'entropie de cet ensemble $\mathcal{A}$ comme le logarithme de $N(\epsilon, \mathcal{A}, d)$.*

Lemme 2. *Supposons qu'il existe une suite d'ensembles mesurables Θ_n , une suite ϵ_n avec $n\epsilon_n^2 \geq 1$, et des constantes a et K satisfaisant l'hypothèse de tests locaux (3.7) pour tout $\theta_1 \in \Theta_n$ tel que $d_n(\theta_0, \theta_1) > \epsilon_n$. Supposons également la condition d'entropie suivante :*

$$\log N(\epsilon_n, \Theta_n, d_n) \leq Dn\epsilon_n^2,$$

pour une constante $D > 0$. Alors, pour une constante $c > 0$ donnée, il existe une constante $M = M(a, K, c)$ assez grande et des tests ϕ_n tels que :

$$\mathbb{E}_{\theta_0}^{(n)}[\phi_n] \xrightarrow{n \rightarrow \infty} 0, \quad \sup_{\theta \in \Theta_n: d_n(\theta, \theta_0) > M\epsilon_n} \mathbb{E}_{\theta}^{(n)}[1 - \phi_n] \leq e^{-cn\epsilon_n^2}.$$

Intuitivement, cette condition d'entropie est une mesure de la complexité du modèle. Pour résumer cette discussion, le théorème 1 peut se reformuler comme suit.

Théorème 2 (Contraction de posterior - version entropie). *Soit $(\epsilon_n)_n$ une suite telle que $n\epsilon_n^2 \rightarrow \infty$ quand $n \rightarrow \infty$. Supposons qu'il existe des constantes $C_1, C_2 > 0$, des ensembles mesurables $\Theta_n \subset \Theta$ tels que :*

(a) *Condition de Kullback-Leibler :*

$$\Pi(B_n^{KL}(\theta_0, \epsilon_n)) \geq e^{-C_1 n \epsilon_n^2}.$$

(b) *Condition de "sieve" : $\Pi(\Theta_n^c) \leq e^{-(C_1+4)n\epsilon_n^2}$,*

(c) *Condition de tests : il existe une distance d_n sur Θ et des constantes $a, K > 0$ satisfaisant l'hypothèse de tests locaux (3.7) pour tout $\theta_1 \in \Theta_n$ tel que $d_n(\theta_0, \theta_1) > \epsilon_n$,*

(d) *Condition d'entropie : $\log N(\epsilon_n, \Theta_n, d_n) \leq C_2 n \epsilon_n^2$.*

Alors, il existe une constante $M = M(a, K, C_1, C_2)$ assez grande telle que la distribution a posteriori se concentre au taux ϵ_n autour de θ_0 :

$$\mathbb{E}_{\theta_0}^{(n)} \left[\Pi \left(\theta : d_n(\theta, \theta_0) > M\epsilon_n \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Ainsi, ces conditions assurent que l'information extraite des données est suffisante pour surpasser les informations a priori et obtenir une distribution a posteriori avec de bonnes propriétés de concentration. Elles permettent d'extraire des conditions suffisantes sur le prior pour avoir la contraction. Dans la section 4.2.1 du chapitre suivant, nous examinerons les modèles où cette propriété de contraction de la distribution a posteriori a été démontrée sous priors *spike-and-slab* dans le contexte de la grande dimension.

3.2.3 Consistance en sélection

La propriété de concentration de la posterior autour de la véritable valeur du paramètre θ_0 est donc un critère minimal pour valider une méthode d'estimation bayésienne et permet notamment de trouver des contraintes suffisantes sur le prior pour assurer la contraction. Cependant, dans un contexte de grande dimension et de sélection de variables, notre objectif est également d'estimer correctement le support du vecteur de régression. En d'autres termes, nous visons non seulement à obtenir des estimations des paramètres proches de leurs vraies valeurs, mais aussi à identifier correctement les coefficients nuls en les estimant effectivement à zéro.

3.2.3.1 Définition

Considérons toujours le modèle bayésien (3.3)-(3.4) de paramètre $\theta \in \Theta$ avec l'hypothèse fréquentiste que les données observées $Y^{(n)}$ sont générées selon la vraie valeur du paramètre θ_0 . Nous supposons de plus ici que $\Theta = \mathbb{R}^d$, $d \gg n$, avec $\theta = (\beta, \theta')$, où $\beta \in \mathbb{R}^p$ est un vecteur de régression de support S_β sur lequel nous voulons effectuer une sélection de variables avec $p \gg n$, et $\theta' \in \mathbb{R}^{d-p}$ est l'ensemble des autres paramètres à estimer. Par exemple, pour le modèle de régression linéaire (3.1), θ' correspond simplement à la variance résiduelle σ^2 , tandis que pour le modèle à effets mixtes (2.1)-(2.2), θ' correspond à (ψ, σ^2, Γ) . On note β_0 la vraie valeur du paramètre de régression (au sens fréquentiste) et S_0 son support.

Définition 5 (Consistance en sélection). *La distribution a posteriori satisfait la condition de consistance en sélection si :*

$$\sup_{\theta_0 \in \Theta_0} \mathbb{E}_{\theta_0}^{(n)} \left[\Pi(\beta : S_\beta \neq S_0 | Y^{(n)}) \right] \xrightarrow{n \rightarrow \infty} 0,$$

où Θ_0 est un sous-ensemble de Θ (éventuellement égal à Θ).

Intuitivement, cette propriété établit que le vrai support S_0 est retrouvé avec probabilité tendant vers 1. La consistance en sélection de la posterior implique en particulier que le modèle ayant la plus grande masse a posteriori est consistant au sens fréquentiste du terme (Castillo et al., 2015).

La consistance en sélection est un critère plus fin que la contraction de la distribution a posteriori. En effet, si certaines composantes de β_0 sont très proches de zéro, il devient difficile pour toute méthode de sélection de variables de les identifier comme non-nulles. Cela peut entraîner la sélection de faux positifs par la distribution a posteriori. Malgré cela, l'estimation des paramètres peut rester précise car les coefficients de ces fausses découvertes ont des valeurs négligeables. Pour éviter ces situations extrêmes qui compromettent la sélection correcte des variables, une condition de type "beta-min" (Bühlmann and Van De Geer, 2011) est nécessaire pour garantir que les vrais coefficients non-nuls de β_0 aient une valeur minimale suffisante (Castillo et al., 2015; Jiang and Sun, 2019; Jeong and Ghosal, 2021b). Par ailleurs, pour obtenir ce résultat de consistance, des hypothèses plus fortes sur le prior vont être également nécessaires, notamment pour rendre parcimonieuse la distribution a posteriori induite par ce prior.

3.2.3.2 Schéma classique de preuve

Depuis les travaux fondateurs de Castillo et al. (2015) sur la régression linéaire en grande dimension sous prior parcimonieux de type *spike-and-slab* discret, les propriétés asymptotiques

fréquentistes de la régression bayésienne parcimonieuse ont été explorées dans divers contextes, en suivant généralement le schéma de preuve proposé par [Castillo et al. \(2015\)](#). Ce schéma est détaillé ci-après.

1. **Dimension effective** : Afin d'atteindre la consistance en sélection pour la posterior, une première étape consiste à étudier une propriété de dimensionnalité effective liée à la taille du support de β . Plus précisément, il s'agit de démontrer que la distribution a posteriori tend à se concentrer sur les supports de taille relativement petite, pas beaucoup plus grande que la taille du vrai support :

$$\sup_{\theta_0 \in \Theta_0} \mathbb{E}_{\theta_0} \left[\Pi \left(\beta : |S_\beta| > K |S_0| \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0,$$

pour une certaine constante $K > 0$. Cette propriété repose sur le lemme 1 qui fournit une borne inférieure pour le dénominateur de la distribution a posteriori avec une probabilité tendant vers 1.

2. **Contraction de la posterior** : Le résultat précédent permettant de se concentrer sur les supports de petite taille, l'objectif ensuite est de dériver un taux de contraction de la distribution a posteriori en utilisant la théorie générale de [Ghosal and Van der Vaart \(2017\)](#) présentée dans la section 3.2.2.3. Pour cela, des hypothèses dites de compatibilité doivent être faites sur la matrice de design. En effet, dans le contexte de la régression linéaire (3.1) en grande dimension par exemple, le paramètre β ne peut pas être estimé sans conditions sur la matrice de design X . Plus précisément, comme β est supposé parcimonieux, une hypothèse d'inversibilité "locale" de la matrice de Gram $X^\top X$ est suffisante pour l'estimation : c'est l'objectif des hypothèses de compatibilité. Le lecteur peut se référer à [Castillo et al. \(2015\)](#) pour plus de précisions. Généralement, la contraction de la posterior est d'abord démontrée avec d la distance de Hellinger quadratique moyenne définie par :

$$\frac{1}{n} \sum_{i=1}^n d_H^2(p_{\theta_0,i}, p_{\theta,i}), \quad \text{avec } d_H^2(p, q) = \int \left(\sqrt{p(x)} - \sqrt{q(x)} \right)^2 dx,$$

en notant $P_\theta^{(n)}$ le produit des mesures $P_{\theta,1}, \dots, P_{\theta,n}$ de densités respectives $p_{\theta,1}, \dots, p_{\theta,n}$. Ceci permet ensuite d'en déduire la contraction de la posterior pour des distances plus interprétables telles que les normes ℓ_1 ou ℓ_2 dans le cas de densités gaussiennes. Cependant, dans le cas de densités gaussiennes multivariées avec une moyenne spécifique à l'individu et une matrice de covariance inconnue, établir la proximité en terme de distance de Hellinger quadratique moyenne ne permet pas de garantir la proximité en terme de distance euclidienne pour les paramètres de moyenne de ces densités. Pour surmonter ce problème, la solution proposée est une utilisation directe de la divergence moyenne de Rényi d'ordre $1/2$:

$$R_n(p_{\theta_0}^{(n)}, p_\theta^{(n)}) = -\frac{1}{n} \sum_{i=1}^n \log \left(\int \sqrt{p_{\theta_0,i}(x) p_{\theta,i}(x)} dx \right).$$

On trouve des exemples d'application de cette théorie dans les travaux de [Ning et al. \(2020\)](#) et de [Jeong and Ghosal \(2021b\)](#) par exemple.

3. **Approximation distributionnelle** : Une fois qu'on a obtenu la contraction de la distribution a posteriori autour de la vraie valeur du paramètre, l'étape clé suivante, introduite pour la première fois par [Castillo et al. \(2015\)](#), consiste à quantifier l'incertitude par le biais de la dispersion de la posterior. Cette quantification est généralement obtenue en utilisant une approximation distributionnelle de la posterior. Comme dans un théorème de Bernstein-von Mises (BvM), la distribution a posteriori du paramètre de régression est approchée par une distribution relativement plus simple, mais contrairement à un théorème BvM traditionnel en grande dimension ([Ghosal, 1999](#)), la distribution approximative est un mélange de gaussiennes multivariées au lieu d'une seule ([Castillo et al., 2015](#); [Ning et al., 2020](#); [Jeong and Ghosal, 2021b](#)) :

$$\sup_{\theta_0 \in \Theta_0} \mathbb{E}_{\theta_0} \left[\left\| \Pi \left(\beta \in \cdot | Y^{(n)} \right) - \Pi^\infty \left(\beta \in \cdot | Y^{(n)} \right) \right\|_{TV} \right] \xrightarrow{n \rightarrow \infty} 0,$$

où Π^∞ est un mélange aléatoire, sur chaque support possible, de distributions gaussiennes multivariées. Cependant, il convient de noter que la consistance en sélection de la posterior peut, dans certains cas, être obtenue sans approximation distributionnelle. Par exemple, dans une régression linéaire gaussienne parcimonieuse, [Narisetty and He \(2014\)](#) et [Song and Liang \(2023\)](#) ont exprimé directement la vraisemblance marginale du support du vecteur de régression en intégrant la vraisemblance complète en β et σ^2 à l'aide du noyau inverse gamma. Toutefois, cette approche ne peut pas s'étendre facilement à un cas où le modèle contient une matrice de covariance de dimension supérieure à 2.

4. **Contrôle de la sur-sélection** : Cette approximation de la forme de la distribution a posteriori permet d'obtenir le résultat suivant qui montre que la distribution a posteriori ne mettra pas de masse sur des supports contenant strictement le vrai modèle :

$$\sup_{\theta_0 \in \Theta_0} \mathbb{E}_{\theta_0} \left[\Pi \left(\beta : S_\beta \supset S_0, S_\beta \neq S_0 \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Pour démontrer ce théorème, la condition "beta-min" n'est même pas nécessaire. En effet, il suffit que le prior choisi induise suffisamment de parcimonie et favorise ainsi la sélection d'un support de petite taille, de manière à ne pas ajouter de coordonnées inutiles lorsque toutes les coordonnées réellement non-nulles sont présentes ([Castillo et al., 2015](#)).

5. **Consistance en sélection** : La propriété précédente n'exclut pas le cas où la distribution a posteriori charge des supports qui ne contiennent qu'un sous-ensemble du vrai modèle S_0 et éventuellement d'autres coordonnées (des faux positifs). Ainsi, la consistance en sélection découle de la combinaison entre la condition "beta-min" et la propriété précédente ([Castillo et al., 2015](#); [Ning et al., 2020](#); [Jeong and Ghosal, 2021a](#)). Ainsi, ce dernier résultat permet d'améliorer l'approximation de la distribution a posteriori puisqu'elle peut alors être approchée par une seule composante du mélange, celle associée au support S_0 , ce qui donne un théorème de BvM ([Castillo et al., 2015](#); [Ning et al., 2020](#); [Jeong and Ghosal, 2021b](#)).

CHAPITRE 4

Positionnement et contributions de la thèse

La problématique principale abordée dans cette thèse est la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes (NLMEM, *Non-Linear Mixed-Effects Models*). Plusieurs procédures ont déjà été mises au point dans le cadre spécifique des modèles linéaires à effets mixtes, mais les contributions proposées lorsque le modèle est non-linéaire sont rares et ne reposent sur aucun fondement théorique. Dans cette thèse, nous développons une procédure de sélection de covariables en grande dimension dans les modèles non-linéaires à effets mixtes, en étudiant à la fois son implémentation pratique et ses propriétés théoriques.

Dans ce chapitre, nous présentons les apports des chapitres 5 et 6. Plus précisément, la section 4.1 présente un état de l’art des méthodes de sélection de variables dans les modèles à effets mixtes, mettant en lumière un manque de procédures computationnellement efficaces en grande dimension, notamment dans les modèles non-linéaires à effets mixtes. Cette section introduit également les contributions du chapitre 5, qui incluent un développement méthodologique dans ces modèles, des études de simulation et une application sur des données réelles. Ensuite, la section 4.2 examine les propriétés asymptotiques fréquentistes des priors *spike-and-slab*, en soulignant l’absence de résultats dans les modèles non-linéaires. Elle présente ensuite les contributions du chapitre 6, qui développe des taux de contraction a posteriori pour un tel prior dans le cas d’un modèle non-linéaire à effets mixtes.

Table des matières

4.1	Procédures de sélection de variables en grande dimension dans les NLMEM	45
4.1.1	État de l’art des méthodes de sélection de variables en grande dimension dans les modèles à effets mixtes	45
4.1.1.1	Sélection de variables dans les modèles linéaires à effets mixtes	45
4.1.1.2	Sélection de variables dans les modèles non-linéaires à effets mixtes	46
4.1.2	Contribution du Chapitre 5 : Méthode de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes	47
4.2	Taux de contraction a posteriori en grande dimension sous prior <i>spike-and-slab</i>	49
4.2.1	État de l’art des propriétés asymptotiques fréquentistes des priors <i>spike-and-slab</i>	49
4.2.1.1	En régression linéaire	49
4.2.1.2	Dans des modèles plus complexes	50
4.2.2	Contribution du Chapitre 6 : Taux de contraction a posteriori en grande dimension dans un NLMEM	51

4.1 Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes

4.1.1 État de l’art des méthodes de sélection de variables en grande dimension dans les modèles à effets mixtes

Ces dernières années ont vu l’émergence de contributions variées sur la sélection de covariables en grande dimension dans les modèles à effets mixtes décrits par les équations suivantes, autrement appelée sélection des effets fixes :

$$\begin{cases} y_{ij} = f(\varphi_i, \psi, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i = X_i \beta + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \end{cases} \quad \begin{matrix} (4.1a) \\ (4.1b) \end{matrix}$$

avec $\varphi_i \in \mathbb{R}^q$, $X_i \in \mathbb{R}^{q \times p}$, $\beta \in \mathbb{R}^p$. Il est important de noter qu’identifier les covariables influentes dans ces modèles est difficile, car la sélection concerne des variables latentes du modèle : les paramètres individuels φ_i . Cela revient à identifier les composantes non-nulles de β . Les outils proposés diffèrent considérablement selon que la fonction de régression f est linéaire ou non linéaire par rapport aux paramètres individuels. Plus précisément, dans le cas linéaire, il est possible de développer des critères dont le calcul et/ou l’étude théorique impliquent des quantités explicites, ce qui est rarement le cas lorsque le modèle est non linéaire.

4.1.1.1 Sélection de variables dans les modèles linéaires à effets mixtes

Dans les modèles linéaires à effets mixtes, la majorité des approches de sélection de covariables en grande dimension s’appuient sur l’application de méthodes régularisées, la plupart d’entre elles bénéficiant de résultats théoriques garantissant les bonnes propriétés des méthodes proposées. On peut citer par exemple les travaux de [Schelldorfer et al. \(2011\)](#) qui proposent une estimation par maximum de vraisemblance avec une pénalité de type ℓ_1 et démontrent, dans le contexte de grande dimension, la consistance de leur estimateur ainsi qu’un résultat de type oracle non-asymptotique. Ils développent un algorithme de descente de gradient par coordonnées et démontrent sa convergence vers un point stationnaire du problème d’optimisation. [Fan and Li \(2012\)](#) proposent une procédure de sélection de variables en deux étapes basée sur une méthode de régularisation, permettant la sélection des effets fixes et des effets aléatoires. La sélection des effets aléatoires consistent à identifier les zéros dans la matrice de covariance des effets aléatoires Γ . Notamment, ils surmontent la difficulté associée à la matrice de covariance inconnue des effets aléatoires en la remplaçant par une matrice de substitution dans la vraisemblance pénalisée, et sous certaines conditions sur cette matrice de substitution, ils démontrent que leur procédure satisfait la consistance en sélection lorsque le nombre d’effets fixes peut croître de manière exponentielle avec la taille de l’échantillon. Plus récemment, [Li et al. \(2018\)](#) proposent une approche de double pénalisation, contrôlant à la fois la parcimonie dans le vecteur de régression et dans la matrice de covariance des effets aléatoires, afin d’estimer et de sélectionner simultanément les effets fixes et aléatoires.

Dans un contexte bayésien, [Frühwirth-Schnatter and Tüchler \(2008\)](#) développent une approche en grande dimension pour sélectionner les effets fixes et aléatoires basée sur la décomposition de Cholesky de la matrice de covariance des effets aléatoires ([Chen and Dunson, 2003](#)).

Ils utilisent alors un prior *spike-and-slab* sur les effets fixes et sur les éléments de la décomposition de Cholesky, et démontrent numériquement que l'efficacité de leur méthode en grande dimension. Plus récemment, [Heuclin et al. \(2023\)](#) proposent une approche qui combine la prior Horseshoe pour les effets fixes avec une version "pliée" du Horseshoe (ou *folded-Horseshoe*) pour les effets aléatoires, permettant ainsi la sélection simultanée des effets fixes et aléatoires. Le prior *folded-Horseshoe*, introduit dans cet article, consiste à utiliser la version pliée de la distribution gaussienne dans l'équation (3.2), c'est-à-dire la loi de $|X|$ pour X suivant une loi gaussienne. En se basant sur deux applications réelles, ils démontrent que leur prior semble insensible aux problèmes de grande dimension et conduit à une estimation non biaisée.

4.1.1.2 Sélection de variables dans les modèles non-linéaires à effets mixtes

En revanche, dans le cadre plus général des modèles non-linéaires à effets mixtes où la fonction de régression f est non-linéaire en φ_i , les résultats en grande dimension sont, à notre connaissance, peu nombreux et les seuls travaux publiés concernent des aspects computationnels. [Bertrand and Balding \(2013\)](#) comparent une approche progressive (*stepwise approach*) utilisant une estimation empirique de Bayes et des approches de régression pénalisées telles que les pénalités de Ridge, Lasso et HyperLasso pour sélectionner les covariables dans des études pharmacocinétiques. Ils concluent que l'approche pénalisée est plus efficace computationnellement et statistiquement. [Bertrand et al. \(2015\)](#) ont alors ensuite proposé une version pénalisée de l'algorithme SAEM, dans laquelle l'étape de maximisation correspond à un problème de moindres carrés pénalisé par une norme ℓ_1 , mais sans calibrer les paramètres de pénalité. [Ollier et al. \(2016\)](#) proposent aussi une extension de l'algorithme SAEM dans laquelle ils introduisent une pénalité Fused-Lasso pour estimer conjointement les modèles non-linéaires à effets mixtes sur plusieurs groupes et détecter les différences pertinentes entre les effets fixes et les variances des effets aléatoires. Dans l'idée de réduire les coûts computationnels, [Ollier \(2022\)](#) propose un algorithme de gradient proximal pour calculer un estimateur de type Lasso afin d'effectuer une sélection conjointe des covariables et des paramètres de corrélation entre les effets aléatoires. Ces deux dernières approches n'ont pas été évaluées dans le cadre de la grande dimension, où le nombre de covariables dépasse largement le nombre d'individus. En effet, leur méthode a uniquement été testée dans des cas de faible dimension, où le nombre de paramètres à estimer est déjà important dans leur démarche de sélection conjointe. Les codes disponibles pour les méthodes susmentionnées ne sont pas génériques et leur réglage est difficile. Dans des contextes de faible dimension, certains travaux se réfèrent à des critères de sélection de modèles tels que le BIC par exemple ([Delattre and Poursat, 2020](#)). Cependant, cette approche ne peut être généralisée au contexte de grande dimension, car le nombre de modèles candidats croît de manière exponentielle avec le nombre de covariables.

La solution la plus naïve pour un praticien confronté à un problème de grande dimension est donc d'adopter une approche en deux étapes, *i.e.* premièrement d'ajuster des modèles non linéaires indépendants à chaque individu i pour estimer φ_i , puis d'effectuer la sélection des variables en utilisant les paramètres individuels estimés (voir section 2.2.1.1 du chapitre 2). Cette deuxième étape repose sur une sélection de variables en grande dimension dans un modèle de régression linéaire, et on peut par exemple utiliser une pénalité Lasso. Comme discuté dans la section 2.2.1.1, ce type d'approche souffre de sa sensibilité aux estimations des paramètres individuels. [Prague and Lavielle \(2022\)](#) proposent également un algorithme, appelé SAMBA (*Stochastic Approximation for Model Building Algorithm*), permettant notamment de sélectionner

les covariables dans les modèles non-linéaires à effets mixtes mais cette sélection est basée sur une approche en deux temps et repose sur la minimisation d'un critère d'information comme AIC ou BIC, qui sont connus pour leur faible efficacité en grande dimension. Ils se comparent aux approches pas à pas classiques, comme *Stepwise Covariate Modeling* (SCM, [Jonsson and Karlsson \(1998\)](#)) ou *COnditional Sampling use for Stepwise Approach based on Correlation tests* (COSSAC, [Ayrat et al. \(2021\)](#)) par exemple, qui consistent à débiter avec le modèle contenant soit toutes les covariables, soit aucune, puis à chaque étape, des variables sont ajoutées ou éliminées du modèle en fonction d'un critère de sélection dérivé de la vraisemblance. Ce type de méthodes n'est pas efficace en grande dimension car l'espace modèle à tester est trop grand. De ce point de vue, SAMBA est plus efficace car ils se ramènent à un modèle linéaire. Dans la discussion de leur article, ils proposent de coupler SAMBA avec une procédure Lasso, ce qui pourrait grandement améliorer les performances de SAMBA en grande dimension, mais ce travail n'est pas encore publié à notre connaissance. Cependant, l'utilisation de ce type d'approche en deux temps entraîne une perte d'information sur l'incertitude des estimations des effets individuels. De plus, ces méthodes ne tirent pas parti de la structure de population des modèles à effets mixtes, qui permet que les individus avec des données manquantes puissent bénéficier des informations des individus entièrement observés.

Les approches bayésiennes pour la sélection des variables en grande dimension ont été relativement peu étudiées dans le contexte des modèles non-linéaires à effets mixtes. Leur expansion s'est principalement concentrée sur des modèles statistiques classiques comme la régression linéaire ou le modèle linéaire généralisé, où la sélection bayésienne des variables a connu un développement intense ces dernières années ([Castillo and van der Vaart, 2012](#); [Narisetty and He, 2014](#); [Jeong and Ghosal, 2021b](#)). Ceci sera plus largement discuté pour le cas du prior *spike-and-slab* dans la section 4.2.1. Ces méthodes bayésiennes encouragent la parcimonie du vecteur de régression en utilisant une variété de priors (voir par exemple [Tadesse and Vannucci \(2021\)](#) et les références qui y figurent), qui peuvent avoir de meilleures propriétés que le prior double exponentiel associé à la pénalité Lasso. Récemment [Lee \(2022\)](#) a proposé une vue d'ensemble des modèles bayésiens non-linéaires à effets mixtes. Il aborde notamment les méthodes d'inférence bayésienne, les options de priors et les méthodes de sélection de modèles dans ce contexte. Cependant, ces approches bayésiennes sont basées sur des méthodes MCMC qui sont rarement utilisables pour la sélection de variables en grande dimension à cause du coût computationnel.

4.1.2 Contribution du Chapitre 5 : Méthode de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes

À notre connaissance, il n'existe pas de méthode efficace, à la fois en termes de performances de sélection et d'implémentation algorithmique, pour la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Dans le cadre de cette thèse, nous avons donc développé une nouvelle méthode de sélection répondant à ces exigences, présentée en détail dans le chapitre 5 et résumée ci-après.

Plus précisément, considérons le modèle non-linéaire à effets mixtes suivant : pour tout $1 \leq i \leq n$ et $1 \leq j \leq n_i$, la réponse y_{ij} de l'individu i dans la condition t_{ij} est modélisée de la façon suivante :

$$\begin{cases} y_{ij} = g(\varphi_i, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \end{cases} \quad (4.2a)$$

$$\begin{cases} \varphi_i = \mu + \beta^\top V_i + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \end{cases} \quad (4.2b)$$

où la fonction g est non-linéaire en $\varphi_i \in \mathbb{R}^q$. L'équation (4.2b) est formulée de manière à considérer le même ensemble de p covariables pour les q paramètres individuels mais dont l'effet de ces covariables peut varier d'un paramètre individuel à l'autre grâce à la matrice d'effets fixes $\beta = (\beta_{\ell m})_{1 \leq \ell \leq p; 1 \leq m \leq q}$ de dimension $p \times q$. Notre objectif est donc d'identifier le support de β , défini par S_0 :

$$S_0 = \left\{ (\ell, m) \in \{1, \dots, p\} \times \{1, \dots, q\} \mid \beta_{\ell m}^0 \neq 0 \right\},$$

où β^0 désigne la vraie valeur de la matrice d'effets fixes. Notons $\theta = (\mu, \beta, \Gamma, \sigma^2)$ le paramètre de population à estimer. La difficulté ici est que la vraisemblance n'a pas de forme explicite du fait de la non-linéarité du modèle en les variables latentes $(\varphi_i)_i$. L'inférence est donc difficile car les estimateurs classiques n'ont pas de forme explicite.

Dans ce travail, nous avons abordé la sélection de variables par une approche bayésienne combinant l'utilisation d'une version multivariée du prior *spike-and-slab* gaussien pour β et l'algorithme MCMC-SAEM. Dans ce prior, c'est la variance de la distribution *spike*, notée ν_0 , qui joue un rôle de pénalisation. En effet, bien que le prior *spike-and-slab* gaussien ne permette pas une estimation parcimonieuse de β puisque $\nu_0 > 0$, cette valeur de ν_0 a un impact direct sur le nombre de covariables sélectionnées selon un critère de seuil détaillé ci-après. L'originalité de l'approche proposée réside dans la combinaison de cette régularisation bayésienne avec des algorithmes d'inférence traditionnellement utilisés en statistique fréquentiste. En effet, dans le cadre bayésien, les méthodes classiques d'inférence fixent arbitrairement la valeur de ν_0 , basée sur un seuil choisi par le praticien pour définir un niveau d'influence "négligeable" des covariables, puis procèdent à l'échantillonnage dans la distribution a posteriori via un algorithme MCMC. Toutefois, l'échantillonnage MCMC est très coûteux et son utilisation devient très vite prohibitive en grande dimension. Pour surmonter cette limitation, nous avons opté pour l'estimation du maximum a posteriori, ce qui nous permet de reformuler notre problème d'inférence comme un problème d'optimisation plutôt que d'échantillonnage. Ainsi, nous pouvons utiliser l'algorithme d'optimisation MCMC-SAEM, beaucoup plus rapide (environ 20 fois plus) que l'échantillonnage MCMC classique. La rapidité de notre algorithme nous permet d'explorer différents niveaux de parcimonie dans β en faisant varier la valeur de ν_0 . Notre procédure se compose donc de deux étapes, basée sur l'approche fréquentiste de sélection de modèles.

1. La première étape de la procédure consiste à réduire le nombre de modèles candidats. Pour ce faire, de façon similaire au Lasso, on explore une grille de valeurs du paramètre de pénalisation ν_0 , notée Δ , afin d'obtenir une collection de modèles prometteurs. Pour chaque valeur de $\nu_0 \in \Delta$, le paramètre de population est estimé par maximum a posteriori, noté $\hat{\theta}$, grâce à un algorithme MCMC-SAEM, que nous avons implémenté spécifiquement pour notre modèle. Ensuite, en appliquant un critère de seuil sur la probabilité a posteriori d'inclusion de chaque covariable sachant $\hat{\theta}$, nous identifions un sous-ensemble de covariables pertinentes $\hat{S}_{\nu_0}$.
2. Ensuite, la deuxième étape consiste à utiliser un critère d'information pour sélectionner le "meilleur" support parmi ceux conservés à l'étape précédente $(\hat{S}_{\nu_0})_{\nu_0 \in \Delta}$. Un choix efficace en grande dimension est de sélectionner $\hat{S}_{\hat{\nu}_0}$ pour $\hat{\nu}_0$ minimisant le critère eBIC (voir section 3.1.1.2 du chapitre 3). Un outil graphique est proposé pour visualiser les résultats de sélection tout au long de la grille, avec la valeur du critère eBIC associée (voir Figure 11 du chapitre 5).

Il est important de noter que le choix du prior gaussien *spike-and-slab*, plutôt qu'une distribution de Dirac en zéro comme distribution *spike* notamment, présente des avantages tant sur le plan algorithmique qu'applicatif. Premièrement, ce prior permet de maintenir le modèle dans la famille exponentielle courbe, ce qui facilite l'implémentation de l'algorithme SAEM. De plus, ce prior continu reflète plus fidèlement la réalité dans notre contexte biologique de sélection de marqueurs génétiques. En effet, il est raisonnable de supposer que de nombreux gènes ont un effet faible mais non nul sur un caractère d'intérêt. Cependant, notre objectif principal reste l'identification des gènes ayant un impact significatif sur ce caractère, justifiant l'utilisation de ce prior plus flexible.

Par ailleurs, cette procédure montre de très bonnes performances de sélection en simulation. Elle semble notamment éviter les faux positifs, ce qui est un atout majeur pour de nombreuses applications, notamment en biologie, où les faux positifs doivent être évités en raison du coût des expériences. Nous avons aussi montré l'efficacité de la méthode sur un jeu de données réelles pour l'identification de marqueurs génétiques impliqués dans le processus de sénescence des feuilles de blés.

4.2 Taux de contraction a posteriori en grande dimension sous prior *spike-and-slab*

À notre connaissance, l'étude théorique de la sélection de variables en grande dimension dans les modèles à effets mixtes n'est explorée dans la littérature que dans le cas linéaire. L'absence de résultats théoriques dans le cas non-linéaire peut refléter les difficultés et les complexités inhérentes à l'extension de ces analyses dans ce contexte. Dans cette thèse, nous nous concentrons sur l'étude théorique de l'utilisation du prior *spike-and-slab* pour la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes.

4.2.1 État de l'art des propriétés asymptotiques fréquentistes des priors *spike-and-slab*

Dans cette section, nous passons en revue la littérature portant sur les propriétés asymptotiques fréquentistes des divers priors *spike-and-slab* présentés dans la section 3.1.2.2 du chapitre 3, dans différentes configurations de modèles en grande dimension. Il est cependant important de noter que récemment, ces propriétés ont également été étudiées pour des priors de rétrécissement, comme ceux exposés dans la section 3.1.2.3 du chapitre 3, par Song and Liang (2023). Ils démontrent que, dans le contexte de la régression linéaire bayésienne en grande dimension, les priors de rétrécissement caractérisés par une densité présentant un pic dominant autour de zéro et une queue lourde plate, présentent des propriétés théoriques aussi solides que celles des priors *spike-and-slab*.

4.2.1.1 En régression linéaire

Dans un premier temps, Narisetty and He (2014) ont établi un résultat de consistance en sélection dans un modèle de régression linéaire gaussienne en utilisant un prior *spike-and-slab* continu gaussien sur le paramètre de régression β et un prior inverse gamma pour la variance résiduelle σ^2 . Comme mentionné précédemment dans la section 3.2.3.2 du chapitre 3, ils ont adopté une

approche qui consiste à intégrer la vraisemblance jointe en β et σ^2 à l'aide du noyau inverse gamma. Cependant, leur approche ne peut pas directement s'étendre à des modèles contenant une matrice de covariance de dimension supérieure à 2.

Ensuite, [Castillo et al. \(2015\)](#) ont développé le schéma de preuve présenté dans la section 3.2.3.2 du chapitre 3, et ont obtenu un résultat de consistance en sélection dans un modèle de régression linéaire gaussienne (3.1) en utilisant un prior *spike-and-slab* discret pour le vecteur de régression avec une loi Laplace comme distribution *slab*, et en supposant la variance résiduelle connue. Par la suite, [Ročková and George \(2018\)](#) ont étudié un prior *spike-and-slab* Lasso, où les distributions *spike* et *slab* suivent des lois de Laplace, pour la régression linéaire gaussienne en grande dimension avec variance résiduelle connue. Dans ce cadre, ils ont obtenu des taux de contraction de la distribution a posteriori similaires à ceux de [Castillo et al. \(2015\)](#), en imposant notamment que le paramètre de la distribution *spike* ne croisse pas plus rapidement que p^d , pour $d \geq 2$. Leur travail étend en quelque sorte les recherches de [Castillo et al. \(2015\)](#), car ils ont considéré une loi de Laplace comme distribution *spike* plutôt qu'une distribution de Dirac en zéro. Cependant, ils se sont limités à démontrer des résultats de contraction sans prouver la consistance en sélection.

[Jiang and Sun \(2019\)](#) ont proposé des conditions générales sous lesquelles un prior *spike-and-slab* générique peut atteindre des taux de contraction de la distribution a posteriori presque optimaux et la consistance en sélection dans un modèle de régression linéaire gaussienne en grande dimension avec variance inconnue. Pour obtenir ces résultats, ils supposent notamment une propriété de dégénérescence vers le Dirac en zéro sur leur distribution *spike* et une condition qui permet d'éviter une finesse excessive de la distribution *slab* autour des vrais coefficients de régression. Les auteurs démontrent également que leurs résultats incluent ceux de [Castillo et al. \(2015\)](#) et [Narisetty and He \(2014\)](#).

D'autres recherches basées sur des approches bayésiennes empiriques, où le prior est estimé à partir des données, sont également disponibles dans la littérature. Pour plus de détails, le lecteur peut se référer à [Martin et al. \(2017\)](#); [Chae et al. \(2019\)](#); [Belitser and Ghosal \(2020\)](#); [Tang and Martin \(2023\)](#).

4.2.1.2 Dans des modèles plus complexes

Ces propriétés asymptotiques fréquentistes du *spike-and-slab* ont également été récemment examinées dans des modèles plus complexes que la régression linéaire gaussienne. Par exemple, [Ning et al. \(2020\)](#) ont étendu les travaux de [Castillo et al. \(2015\)](#) à la régression linéaire multivariée en grande dimension. Plus précisément, ils ont obtenu un résultat de consistance en sélection sous une version multivariée du prior *spike-and-slab* discret de [Castillo et al. \(2015\)](#), adaptée à la parcimonie par groupe de variables. Comme discuté dans la section 3.2.3.2 du chapitre 3, [Ning et al. \(2020\)](#) ont proposé d'utiliser la divergence de Rényi moyenne d'ordre $1/2$, au lieu de la distance de Hellinger, pour obtenir des taux de contraction de la distribution a posteriori en présence d'une matrice de covariance résiduelle inconnue. En exploitant la propriété de diagonalisation de la matrice de covariance, ils ont assigné un prior inverse gaussien indépendant à chaque valeur propre de cette matrice. De plus, contrairement à [Castillo et al. \(2015\)](#), pour satisfaire la condition de Kullback-Leibler du Théorème 1, [Ning et al. \(2020\)](#) supposent que les vraies valeurs des paramètres sont restreintes à certaines régions. Ensuite, en s'appuyant sur les travaux de [Ning et al. \(2020\)](#), [Shen and Deshpande \(2022\)](#) ont obtenu des taux de contraction a posteriori pour le *spike-and-slab* Lasso multivarié en grande dimension.

Jeong and Ghosal (2021b) se sont quant à eux concentrés sur un modèle de régression linéaire avec paramètres de nuisances en grande dimension. Plus précisément, ils ont considéré le modèle suivant :

$$Y_i = X_i\theta + \xi_{\eta,i} + \varepsilon_i, \quad \varepsilon_i \stackrel{\text{ind}}{\sim} \mathcal{N}_{m_i}(0, \Delta_{\eta,i}), \quad i = 1, \dots, n,$$

où $Y_i \in \mathbb{R}^{m_i}$ est le vecteur réponse, $X_i \in \mathbb{R}^{m_i \times p}$ les covariables, $\theta \in \mathbb{R}^p$ vecteur de régression avec $p > n$, η est un paramètre de nuisance de dimension possiblement infinie, $\xi_{\eta,i} \in \mathbb{R}^{m_i}$ et $\Delta_{\eta,i} \in \mathbb{R}^{m_i \times m_i}$, $m_i \geq 1$. Ce type de modèle inclut par exemple les modèles linéaires à effets mixtes pour : $\xi_{\eta,i} = 0$, $\Delta_{\eta,i} = \sigma^2 I_{m_i} + Z_i \Gamma Z_i^\top$, et $\eta = \Gamma$ la matrice de covariance des effets aléatoires. Un prior *spike-and-slab* discret comme celui de Castillo et al. (2015) est utilisé comme prior pour les coefficients de régression, et un prior général est utilisé pour les paramètres de nuisance. Ce dernier est soumis à diverses hypothèses permettant d'atteindre les propriétés asymptotiques désirées. À noter également que l'utilisation de la loi Laplace pour la distribution *slab* n'est pas nécessaire ici, contrairement à Castillo et al. (2015), puisque une hypothèse de restriction sur le vrai signal est imposée. D'après Jeong and Ghosal (2021b), d'autres distributions *slab* devraient donc également fonctionner, comme une densité gaussienne par exemple, mais avec des restrictions plus fortes sur le signal réel. Ainsi, sous des hypothèses sur les vraies valeurs des paramètres et sur les priors notamment, Jeong and Ghosal (2021b) obtiennent la consistance en sélection de la posterior dans ce modèle générique. En particulier, la consistance en sélection dans les modèles linéaires à effets mixtes sous le prior *spike-and-slab* discret et le prior *Inverse-Wishart* sur la matrice de covariance des effets aléatoires est démontrée sous l'asymptotique $|S_0|^5 \log^3(p)/n \xrightarrow{n \rightarrow \infty} 0$ (voir Théorème 10 de Jeong and Ghosal (2021b)).

Jeong and Ghosal (2021a) ont étudié le modèle linéaire généralisé en grande dimension sous le prior *spike-and-slab* discret. En plus de la distribution a posteriori habituelle, ils ont évalué la posterior fractionnaire, qui est obtenue en appliquant le théorème de Bayes avec une puissance fractionnaire de la vraisemblance :

$$\Pi_\alpha(B|Y^{(n)}) = \frac{\int_B \Lambda_n^\alpha(\theta) d\Pi(\theta)}{\int_\Theta \Lambda_n^\alpha(\theta) d\Pi(\theta)}, \quad \forall B \in \mathcal{B},$$

où $\Lambda_n(\theta) = p_\theta^{(n)}/p_{\theta_0}^{(n)}(Y^{(n)})$ et $\alpha \in (0, 1]$. Cette définition se réduit à la distribution a posteriori habituelle si $\alpha = 1$. Cette posterior fractionnaire permet d'uniformiser la contraction de la distribution a posteriori sur un sous-ensemble plus large de l'espace des paramètres. Ainsi, sous ce modèle et diverses hypothèses sur les priors et les vraies valeurs des paramètres, Jeong and Ghosal (2021a) obtiennent des taux de contraction a posteriori similaires à ceux de Castillo et al. (2015) pour la régression linéaire classique.

4.2.2 Contribution du Chapitre 6 : Taux de contraction a posteriori dans un modèle non-linéaire à effets mixtes de grande dimension

Comme discuté dans la section précédente, des travaux récents ont mis en évidence un intérêt croissant pour l'étude des propriétés asymptotiques fréquentistes des procédures bayésiennes appliquées aux modèles linéaires en grande dimension avec des contraintes de parcimonie. Cependant, à notre connaissance, il existe une lacune dans la littérature concernant des résultats

théoriques similaires pour les modèles non-linéaires. Le chapitre 6 comble cette lacune en apportant une nouvelle contribution centrée spécifiquement sur un modèle non-linéaire à effets mixtes sous prior *spike-and-slab* discret. Plus précisément, nous considérons le modèle suivant :

$$Y_i = f_i(X_i\beta) + Z_i\xi_i + \varepsilon_i, \quad \varepsilon_i \stackrel{\text{ind.}}{\sim} \mathcal{N}_{n_i}(0, \sigma^2 I_{n_i}), \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \quad i = 1, \dots, n. \quad (4.3)$$

Dans l'équation ci-dessus, n est le nombre d'individus, $Y_i \in \mathbb{R}^{n_i}$ et n_i sont respectivement le vecteur et le nombre d'observations pour le sujet i , $\xi_i \in \mathbb{R}^q$ est un vecteur composé de q effets aléatoires, $\varepsilon_i \in \mathbb{R}^{n_i}$ est le terme d'erreur. $X_i \in \mathbb{R}^{q \times p}$ et $Z_i \in \mathbb{R}^{n_i \times q}$ sont des matrices de design composées de variables explicatives qui relient, respectivement, les observations aux effets fixes $\beta \in \mathbb{R}^p$ et aux effets aléatoires. Pour établir un parallèle avec la littérature classique sur les modèles à effets mixtes pour les données répétées, nous définissons

$$f_i(x) = (f(x, t_{i1}), f(x, t_{i2}), \dots, f(x, t_{in_i}))^\top,$$

où t_{ij} représente le j -ième temps d'observation pour l'individu i , et $f : \mathbb{R}^q \times \mathbb{R} \rightarrow \mathbb{R}$ représente une fonction de régression choisie pour capturer efficacement le phénomène observé. Bien que le choix de $f_i(X_i\beta) = X_i\beta$ donne le modèle linéaire standard à effets mixtes, il est courant dans de nombreuses applications de choisir une fonction non-linéaire f . En outre, lorsque f est non-linéaire, le modèle précédent est également appelé modèle marginal non-linéaire à effets mixtes (voir section 2.2.1.2 du chapitre 2). Le terme "marginal" provient du fait que, contrairement à de nombreux autres modèles à effets mixtes non-linéaires comme décrits par les équations (4.1), l'espérance et la variance des observations possèdent toutes deux une expression explicite, ce qui facilite les développements théoriques. La distribution de Y_i est en effet entièrement caractérisée par :

$$Y_i \sim \mathcal{N}(f_i(X_i\beta), \Delta_{\Gamma,i}), \text{ où } \Delta_{\Gamma,i} = Z_i\Gamma Z_i^\top + \sigma^2 I_{n_i}.$$

Cependant, il convient de noter qu'en utilisant la technique de linéarisation FOA décrite dans la section 2.2.1.2 du chapitre 2, le modèle non-linéaire à effets mixtes général (4.1) peut être approché par le modèle marginal non-linéaire à effets mixtes suivant :

$$Y_i = f_i(X_i\beta) + Z_i(\beta)\xi_i + \varepsilon_i, \quad (4.4)$$

où, contrairement au modèle (4.3) considéré dans ce travail, $Z_i(\beta) = \partial f_i(X_i\beta)/\partial \varphi_i$ dépend du paramètre inconnu β . Malgré cela, et à condition de disposer d'un estimateur fortement consistant de β , on peut traiter le modèle (4.4) dans le même cadre que le modèle (4.3) en fixant $Z_i(\beta)$ à son estimation obtenue en substituant β par son estimateur consistant (Vonesh and Carter, 1992). De cette manière, une méthode d'estimation développée pour le modèle (4.3) pourrait "approcher" les estimations du modèle non-linéaire à effets mixtes général (4.1). L'extension des résultats du chapitre 6 au modèle non-linéaire à effets mixtes général sera plus largement discuté dans les perspectives de cette thèse, chapitre 8.

Ainsi, on se concentre dans ce travail sur l'étude du modèle (4.3). Dans ce modèle, la variance résiduelle σ^2 est supposée connue, tandis que la matrice de covariance des effets aléatoires Γ et le vecteur de régression β sont inconnus et doivent être estimés. Dans ce contexte, notre objectif principal est de sélectionner judicieusement les covariables dans X_i permettant de comprendre la variabilité inter-individuelle, ce qui revient à identifier les composantes non-nulles dans β dans un cadre de grande dimension où $p \gg n$. Comme détaillé dans le schéma de preuve de la

consistance en sélection dans la section 3.2.3 du chapitre 3, l'objectif premier est donc d'estimer $(\beta, \Gamma) \in \mathcal{B} \times \mathcal{H}$ et d'obtenir des résultats de contraction de la distribution a posteriori. Notons (β_0, Γ_0) la vraie valeur des paramètres, et s_0 la taille du support de β_0 . Ainsi, dans le chapitre 6, nous établissons un ensemble de priors appropriés pour atteindre cet objectif. Plus précisément, nous utilisons un prior *slope-and-slab* discret de distribution slab Laplace pour le vecteur de régression β , comme introduit dans Castillo et al. (2015), et un prior *Inverse-Wishart* sur la matrice de covariance des effets aléatoires Γ .

Pour obtenir ces résultats de contraction, nous nous appuyons sur la théorie générale de Ghosal et al. (2000); Ghosal and Van Der Vaart (2007) et détaillée dans la section 3.2.2 du chapitre 3 (voir Théorème 2). Des hypothèses classiquement utilisées dans la littérature sont aussi imposées ici : notamment l'appartenance des vrais paramètres à un espace borné (Ning et al., 2020), notés $\beta_0 \in \mathcal{B}_0$ et $\Gamma_0 \in \mathcal{H}_0$, la décroissance exponentielle du prior de sélection de modèle (Castillo et al., 2015), le contrôle du paramètre de régularisation de la loi Laplace (Castillo et al., 2015) ou encore le contrôle des valeurs propres des matrices de design Z_i (Jeong and Ghosal, 2021b). La difficulté principale de la preuve consiste à contrôler la non-linéarité de la fonction de régression f . Plus précisément, et pour assurer des taux de contraction similaires à ceux observés dans le cas linéaire, nous faisons l'hypothèse que la fonction de régression f est Lipschitzienne. Ainsi, sous le régime asymptotique que le nombre de covariables p croît de manière presque exponentielle avec le nombre d'individus n et que le nombre total d'observations est dominé par une puissance de p , nous obtenons les taux de contraction suivant pour Γ et le terme de prédiction $f_i(X_i\beta)$ (voir Théorème 3 du chapitre 6).

Théorème 3 (Taux de contraction pour Γ et $f_i(X_i\beta)$). *Sous les hypothèses précédentes, il existe des constantes $C_3, C_4, C_5 > 0$ telles que :*

$$\sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\Gamma : \|\Gamma - \Gamma_0\|_F > C_4 \sqrt{\frac{s_0 \log(p)}{n}} \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0,$$

$$\sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \sqrt{\frac{1}{n} \sum_{i=1}^n \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2} > C_5 \sqrt{\frac{s_0 \log(p)}{n}} \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Par ailleurs, pour obtenir des taux de contraction similaires à ceux observés dans le cas linéaire pour la récupération du vecteur de régression, il est nécessaire de formuler des hypothèses supplémentaires. Ces hypothèses permettront d'assurer l'identifiabilité du vecteur β . Premièrement, une condition additionnelle sur la fonction de régression f permettant de garantir l'identifiabilité de $X_i\beta$ à partir des évaluations de f est nécessaire. Plus précisément, nous supposons que pour tout $i \in \{1, \dots, n\}$ et pour tout $t \in \mathbb{R}$, la vraie valeur $X_i\beta_0$ ne doit pas se situer dans un voisinage des points critiques de $f(\cdot, t)$. Ensuite, nous adoptons les conditions de compatibilité couramment utilisées pour assurer une hypothèse d'inversibilité "locale" de la matrice de Gram $X^\top X$ (Castillo et al., 2015), ce qui permet ainsi l'identifiabilité de β . En particulier, la non-linéarité du modèle nous conduit à contraindre un peu plus l'espace auquel appartient le vrai paramètre β_0 , nous notons cet espace $\bar{\mathcal{B}}_0$. Ainsi, sous ces hypothèses supplémentaires, nous obtenons des taux de contraction pour $X\beta$ et β semblables au cas linéaire (voir Théorème 4 du chapitre 6).

Théorème 4 (Taux de contraction pour $X\beta$ et β). *Sous les hypothèses précédentes, il existe des constantes $C_6, C_7, C_8 > 0$ telles que :*

$$\begin{aligned} \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|X(\beta - \beta_0)\|_2 > C_6 \sqrt{s_0 \log(p)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|\beta - \beta_0\|_2 > C_7 \frac{\sqrt{s_0 \log(p)}}{\|X\|_* \phi_2((C_1 + 1)s_0)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|\beta - \beta_0\|_1 > C_8 \frac{s_0 \sqrt{\log(p)}}{\|X\|_* \phi_1((C_1 + 1)s_0)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \end{aligned}$$

où $\|X\|_* = \max_j \|X_{\cdot j}\|_2$ pour $X_{\cdot j}$ la j -ème colonne de X , et pour $s > 0$:

$$\phi_2(s) = \inf_{\beta: 1 \leq s_\beta \leq s} \frac{\|X\beta\|_2}{\|X\|_* \|\beta\|_2}, \quad \phi_1(s) = \inf_{\beta: 1 \leq s_\beta \leq s} \frac{\|X\beta\|_2 \sqrt{s_\beta}}{\|X\|_* \|\beta\|_1}.$$

Ces taux coïncident avec ceux obtenus par Jeong and Ghosal (2021a) et Jeong and Ghosal (2021b), respectivement pour le modèle linéaire généralisé et pour le modèle linéaire à effets mixtes. Ces résultats de contraction obtenus avec le prior *slope-and-slab* discret constituent les premières garanties théoriques pour l'estimation en grande dimension dans un modèle non-linéaire à effets mixtes.

Partie II

Travaux et résultats

Sélection de covariables bayésienne en grande dimension dans les modèles non-linéaires à effets mixtes à l'aide de l'algorithme SAEM

Publication

Ce chapitre présente le travail de l'article suivant :

Naveau M., Kon Kam King G., Rincent R., Sansonnet L., and Delattre M., **Bayesian high-dimensional covariate selection in non-linear mixed-effects models using the SAEM algorithm**. *Statistics and Computing*, 2023, 34 (1), pp.53. [10.1007/s11222-023-10367-4](https://doi.org/10.1007/s11222-023-10367-4).

Pour faciliter la compréhension de la méthode, un exemple simple est fourni à la fin de ce chapitre. Cet exemple est également accessible avec les scripts R associés à cet article à l'adresse suivante : https://github.com/Marion-Naveau/Supp_Information_SAEMVS

Abstract

High-dimensional variable selection, with many more covariates than observations, is widely documented in standard regression models, but there are still few tools to address it in non-linear mixed-effects models where data are collected repeatedly on several individuals. In this work, variable selection is approached from a Bayesian perspective and a selection procedure is proposed, combining the use of a spike-and-slab prior and the Stochastic Approximation version of the Expectation Maximisation (SAEM) algorithm. Similarly to Lasso regression, the set of relevant covariates is selected by exploring a grid of values for the penalisation parameter. The SAEM approach is much faster than a classical MCMC (Markov chain Monte Carlo) algorithm and our method shows very good selection performances on simulated data. Its flexibility is demonstrated by implementing it for a variety of non-linear mixed effects models. The usefulness of the proposed method is illustrated on a problem of genetic markers identification, relevant for genomic-assisted selection in plant breeding.

Keywords: High-dimension, Non-linear mixed-effects models, SAEM algorithm, Spike-and-slab prior, Variable selection

Acknowledgement: The authors thank the Stat4Plant project ANR-20-CE45-0012 for funding. The authors are grateful to the INRAE MIGALE bioinformatics facility (MIGALE, INRAE, 2020. Migale bioinformatics Facility, doi: 10.15454/1.5572390655343293E12) for providing computing and storage resources. The authors are grateful to Emmanuel Heumez (INRAE Estrées-Mons, UE GCIE) for managing the trial in Estrées-Mons, and to Jacques Le Gouis for managing the PIA BreedWheat. The data-set was obtained thanks to the support of the PIA (Investment for the Future Program) Breedwheat (ANR-10-BTBR-03).

Contents

5.1	Introduction	59
5.2	Model description	61
5.2.1	Non-linear mixed-effects model and notations	61
5.2.2	Prior specification	63
5.3	Maximum <i>a posteriori</i> inference and thresholding	64
5.3.1	General description of SAEM algorithm	65
5.3.2	Central decomposition of the Q quantity in spike-and-slab non-linear mixed-effects models	67
5.3.3	Estimator thresholding	69
5.4	Covariate selection procedure	69
5.5	Numerical experiments	71
5.5.1	Comparison with a two-step approach	71
5.5.2	Impacts of the different parameters and collinearity between covariates	73
5.5.3	Comparison with an MCMC implementation	78
5.6	Application to plant senescence genetic marker identification	81
5.6.1	Data description and pre-processing	81
5.6.2	Modelling	82
5.6.3	Results and discussion	82
5.7	Conclusion and perspectives	84
5.8	Appendices	85
5.8.1	Algorithms: synthesis and implementation details	85
5.8.2	Results for scenarios with correlated covariates	91
5.8.3	Further details on real data	93
5.9	Application on a detailed example	95
5.9.1	Convergence of the MCMC-SAEM algorithm	95
5.9.2	Spike-and-slab regularisation plot and model selection	96

5.1 Introduction

Mixed-effects models have been introduced to analyse observations collected repeatedly on several individuals in a population of interest (Lavielle, 2014; Pinheiro and Bates, 2000). This type of data is particularly common in the fields of pharmacokinetics or when modelling biological growth for example, where data is customarily analysed using a non-linear model whose coefficients often

have a biological interpretation. In this case, the intrinsic variability of the data captured by the model parameters is then attributable to different sources (intra-individual, inter-individual, and residual) whose consideration is essential to characterise without bias the biological mechanisms behind the observations. Mixed-effects models allow the study of the responses of individuals with the same overall behaviour but with individual variations characterised by random individual parameters that are not observed. Thus, mixed-effects models are latent variable models. Parameter inference is therefore difficult because the likelihood and classical estimators do not have an explicit form. A widely used solution is to use an EM (Expectation-Maximisation) algorithm, or any variant, to compute the maximum likelihood estimator or the maximum *a posteriori* estimator in a Bayesian framework (Dempster et al., 1977).

Moreover, the description of inter-individual variability may involve a number of covariates much larger than the number of individuals. In this high-dimensional context, it is often desirable to be able to focus on the few most relevant covariates through a variable selection procedure. However, in mixed-effects models, identifying the influential covariates is difficult, as the selection concerns latent variables in the model. Recent years have seen the emergence of varied contributions on high-dimensional covariate selection in mixed-effects models. The proposed tools are very different according to whether the regression function is linear or non-linear with respect to the individual parameters. More precisely, the linear case allows the development of criteria whose calculation and/or theoretical study involve explicit quantities, which is rarely true when the model is non-linear. In linear mixed-effects models, many rely on the use of regularised methods (see Schellendorfer et al. (2011) and Fan and Li (2012) for example) and most of them include theoretical consistency results that guarantee the good properties of the proposed methods. In contrast, in the more general framework of non-linear mixed-effects models (NLMEM), there are few results and the only published works concern computational aspects. Bertrand and Balding (2013) compare a stepwise approach using an empirical Bayes estimate and penalised regression approaches like Ridge, Lasso and HyperLasso penalties to select covariates in pharmacokinetics studies, but without calibrating the penalty parameters. Ollier (2022) proposes a proximal gradient algorithm for computing a Lasso estimator in order to perform joint selection of covariates and correlation parameters between random effects, this method being more suitable in low rather than large dimensional covariate settings. The codes available for the above-mentioned methods are not generic and tuning them is difficult. So the easiest practical solution for a practitioner faced with a high-dimensional problem is to adopt a two-stage approach, *i.e.* to fit independent non-linear models to each individual and perform variable selection using the estimated parameters, losing the uncertainty on these estimates and the beneficial shrinkage property of mixed-effect models.

Bayesian approaches to variable selection have not received much attention in the NLMEM context. The focus for their development has been classical statistical models like linear regression or generalised linear model, for which Bayesian variable selection has been intensively developed in recent years. These methods encourage sparsity in the regression vector by using a variety of priors (see for example Tadesse and Vannucci (2021) and the references therein), which may have better properties than the double-exponential prior associated with the Lasso penalty. Very recently, Lee (2022) proposed an overview of the formulation, interpretation and implementation of Bayesian non-linear mixed-effects models. In particular, he discussed Bayesian inference methods, priors options, and model selection methods in this context. However, these Bayesian approaches are based on Markov Chain Monte-Carlo (MCMC) methods which seldom

scale well enough to be usable for high-dimensional variable selection. The main objective of this paper is to propose a fast Bayesian spike-and-slab approach that can be used to identify the relevant covariates in a non-linear mixed-effects model in a high-dimensional context. More precisely, we extend the EMVS (EM Variable Selection) approach of Ročková and George (2014) to the NLMEM setting. Like EMVS, the proposed approach involves two major steps. The first step is, for different values of the spike hyperparameter, to select a local version of the median probability model (Barbieri and Berger, 2004) using the Stochastic Approximation version of the EM algorithm (SAEM, see Delyon et al. (1999) and Kuhn and Lavielle (2004)). The second step consists in selecting the "best" model among those kept after the first step, using an extension of the BIC criterion (Chen and Chen, 2008). An important difference with Ročková and George (2014) is that our approach is applied to NLMEM and not to classical linear regression models. Due to the model non-linearity and to the latent nature of the model random effects, the central so-called Q -quantity of the EM algorithm often does not have a closed-form expression and posterior distributions are difficult to compute. To overcome these issues, we propose an inference method using the SAEM algorithm rather than simply the EM algorithm as in Ročková and George (2014). Another important difference is that optimal model selection among the sub-models obtained in the first step does not require the calculation of the marginal posterior of the models for a spike parameter being equal to 0, as in Ročková and George (2014), but only of the log-likelihood of the NLMEM taken at the maximum likelihood estimator.

This paper is organised as follows. Section 5.2 describes the non-linear mixed-effects model to introduce the notation, summarises the main objective of our procedure, and defines and motivates the hierarchical prior formulations. Section 5.3 details the key tools for our approach: the SAEM algorithm, to compute the maximum *a posteriori* estimator of the model parameters and a thresholding rule to select a local version of the median probability model and put some coefficients of the regression vector to zero. Next, Section 5.4 describes the variable selection procedure. Section 5.5 evaluates the selection performance of our method through an intensive simulation study and presents comparisons with existing methods. To illustrate the high-dimensional variable selection method proposed in this paper, Section 5.6 presents an application on a real data-set composed of European elite winter wheat varieties. Finally, Section 5.7 concludes with a summary discussion and prospects for future research. More details about algorithms and real data are postponed in the appendices.

5.2 Model description

5.2.1 Non-linear mixed-effects model and notations

The formalism used in this paper is that of Lavielle (2014) and Pinheiro and Bates (2000). Let n be the number of individuals and n_i the number of observations for individual i . Let $n_{tot} = \sum_{i=1}^n n_i$ the total number of observations. Let $X^\top$ denote the transpose of a vector or matrix X . Consider the following non-linear mixed-effects model: for all $1 \leq i \leq n$ and $1 \leq j \leq n_i$,

$$\begin{cases} y_{ij} = g(\varphi_i, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \end{cases} \quad (5.1a)$$

$$\begin{cases} \varphi_i = \mu + \beta^\top V_i + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma). \end{cases} \quad (5.1b)$$

This model is described in two levels. First, at the individual level, Equation (5.1a) describes the intra-individual variability, where the observations y_{ij} in $\mathbb{R}$ represent the response of individual i at time t_{ij} . It is assumed that all individuals follow the same known functional form g which depends non-linearly on an individual parameter φ_i which is q -dimensional. Thus, this function governs intra-individual behaviour. The value of q is closely linked to the choice of the mechanistic function g . It is therefore known, and usually small. The variance $\sigma^2 > 0$ of the Gaussian measurement noise is assumed unknown. Then, at the population level, Equation (5.1b) describes the inter-individual variability. For all $i \in \{1, \dots, n\}$, the q individual parameters $\varphi_i = (\varphi_{im})_{1 \leq m \leq q} \in \mathbb{R}^q$ are modelled as a multivariate Gaussian random variable whose mean is specified as the sum of an intercept μ in $\mathbb{R}^q$ and a linear combination of known covariates measured on individual i and contained in the vector $V_i = (V_{i1}, \dots, V_{ip})^\top \in \mathbb{R}^p$. The term "covariates" refers to explanatory variables which may be relevant to explain inter-individual variability. This term is used to distinguish them from other explanatory variables such as the time variable for example. The number of covariates is denoted by p , $\beta = (\beta_{\ell m})_{1 \leq \ell \leq p; 1 \leq m \leq q} \in \mathcal{M}_{p \times q}$ is an unknown fixed effects matrix, and the inter-individual variance-covariance matrix Γ is assumed unknown. Thus, the inter-individual variability is separated into two parts: on the one hand, β models the variability that can be explained by the covariates V_i , and on the other hand, ξ_i represents the part of variation that is not explained by the measured covariates. In the following, $y_i = (y_{ij})_{1 \leq j \leq n_i}$, $\mathbf{y} = (y_i)_{1 \leq i \leq n}$, $\varphi = (\varphi_{im})_{1 \leq i \leq n; 1 \leq m \leq q} \in \mathcal{M}_{n \times q}$ and $V = (V_{i\ell})_{1 \leq i \leq n; 1 \leq \ell \leq p} \in \mathcal{M}_{n \times p}$ respectively denote the vector of observations for individual i , the vector of all observations, the matrix of all individual parameters, and the matrix of all covariates. Let us also note $\theta = (\mu, \beta, \Gamma, \sigma^2)$ the unknown parameter, also-called the population parameter.

The goal of the present work is to identify the relevant covariates, *i.e.* those that best explain the variability between individuals. This can be framed as identifying the non-zero elements in β . Indeed, for each $(\ell, m) \in \{1, \dots, p\} \times \{1, \dots, q\}$, coefficient $\beta_{\ell m}$ describes the influence of covariate ℓ on the individual values of the m -th parameter of the non-linear curves $(\varphi_{im})_{1 \leq i \leq n}$. More precisely, $\beta_{\ell m} = 0$ means that covariate ℓ has no effect on $(\varphi_{im})_{1 \leq i \leq n}$ whereas $\beta_{\ell m} \neq 0$ means that covariate ℓ gives some information on these parameters. Identifying the relevant covariates for all individual parameters amounts to selecting the support of β , defined by S^* :

$$S^* = \left\{ (\ell, m) \in \{1, \dots, p\} \times \{1, \dots, q\} \mid \beta_{\ell m}^* \neq 0 \right\},$$

where β^* is the true fixed effects matrix. To solve this problem in a high-dimensional context, that is when $p \gg n$, it is natural to assume that each row of β^* is sparse, which means that many $\beta_{\ell m}^*$ are zero. An important point here is that model (5.1) is a model with incomplete-data. Indeed, although the first layer (5.1a) is observed, it is not the case for the individual parameters φ . The main difficulty here is that variable selection concerns latent variables of the model.

Remark 1. *As in the application presented in Section 5.6, we could perform variable selection on $q' < q$ components of φ_i , with the choice of the q' relevant components to be considered for variable selection driven by the initial biological question without changing the methodology presented in what follows.*

5.2.2 Prior specification

To solve this variable selection problem, it is convenient to adopt a Bayesian approach. The purpose of this section is to describe the prior distribution on $\theta = (\mu, \beta, \Gamma, \sigma^2)$. First, in order to find the non-zero coefficients of β , a spike-and-slab mixture prior (George and McCulloch, 1993, 1997; Ročková and George, 2014) is considered in a multivariate setting. To facilitate the formulation of this prior, binary latent variables $\delta = (\delta_{\ell m})_{1 \leq \ell \leq p; 1 \leq m \leq q}$ are introduced, such as:

$$\forall 1 \leq \ell \leq p, \forall 1 \leq m \leq q, \delta_{\ell m} = \begin{cases} 1 & \text{if } (\ell, m) \text{ is to be included in model } S^*, \\ 0 & \text{otherwise.} \end{cases} \quad (5.2)$$

Thus, $\delta_{\ell m} = 1$ indicates that the covariate ℓ provides information on the individual parameter m . In other words, δ characterises the support of β . Note that contrary to many multivariate variable selection methods such as the one implemented in **glmnet** (Friedman et al., 2010), the proposed procedure is not limited to selecting the same set of covariates for all dimensions m , i.e it is not required that $\forall 1 \leq m \leq q, \delta_{\ell m} = \delta_\ell$. The support S^* can therefore be reformulated as follows:

$$S^* = \left\{ (\ell, m) \in \{1, \dots, p\} \times \{1, \dots, q\} \mid \delta_{\ell m}^* = 1 \right\}, \quad (5.3)$$

where δ^* denotes the true support. Then, one would like to find $\hat{\delta}$ that maximises the posterior probability $\pi(\delta | \mathbf{y})$, which corresponds to the most promising support.

The prior formulations proposed here are based on the non-conjugate version of the hierarchical priors of George and McCulloch (1997), summarised as follows:

$$\pi(\beta_{\ell m} | \delta_{\ell m}) = \mathcal{N}(0, (1 - \delta_{\ell m})\nu_0 + \delta_{\ell m}\nu_1), \quad 1 \leq \ell \leq p, 1 \leq m \leq q, 0 < \nu_0 < \nu_1, \quad (5.4a)$$

$$\pi(\mu) = \mathcal{N}_q(0, \sigma_\mu^2 I), \quad \text{with } \sigma_\mu^2 > 0, \quad (5.4b)$$

$$\pi(\sigma^2) = \mathcal{IG}\left(\frac{\nu_\sigma}{2}, \frac{\nu_\sigma \lambda_\sigma}{2}\right), \quad \text{with } \nu_\sigma, \lambda_\sigma > 0, \quad (5.4c)$$

$$\pi(\Gamma) = \mathcal{IW}(\Sigma_\Gamma, d), \quad \text{with } \Sigma_\Gamma \in \mathcal{S}_q^{++}, d > 0, \quad (5.4d)$$

$$\pi(\delta_{\ell m} | \alpha_m) = \alpha_m^{\delta_{\ell m}} (1 - \alpha_m)^{1 - \delta_{\ell m}}, \quad \text{with } \alpha_m \in [0, 1], 1 \leq \ell \leq p, 1 \leq m \leq q, \quad (5.4e)$$

$$\pi(\alpha_m) = \text{Beta}(a_m, b_m), \quad \text{with } a_m, b_m > 0, 1 \leq m \leq q. \quad (5.4f)$$

The key prior distribution used for variable selection in this method is the spike-and-slab Gaussian mixture prior (5.4a) on β . In this prior, ν_0 and ν_1 are parameters controlling the penalisation inducing sparsity in the columns of β . More precisely, $(\beta_{\ell m})_{1 \leq \ell \leq p; 1 \leq m \leq q}$ are independent conditionally on δ , with $\pi(\beta_{\ell m} | \delta_{\ell m} = 0) = \mathcal{N}(0, \nu_0)$ and $\pi(\beta_{\ell m} | \delta_{\ell m} = 1) = \mathcal{N}(0, \nu_1)$. The general recommendation for this type of prior is to set ν_0 small to encourage the exclusion of insignificant effects, and ν_1 large enough to accommodate all plausible β values (see George and McCulloch, 1997). Indeed, when $\delta_{\ell m} = 0$, the prior constrains $\beta_{\ell m}$ to very small values which implies that covariate ℓ has no impact in the individual parameter m in the model. Thus, through the values $\delta_{\ell m}$, the spike-and-slab prior makes it possible to distinguish the selected covariates from the rest.

Note that, since φ is unobserved, one cannot simply centre the variable on which the selection is made as is usually the case in more standard models, and so the inclusion of an intercept μ is necessary. Thus, a vaguely informative Gaussian prior (5.4b) is used for μ , with σ_μ^2 large

enough. This choice of prior has the advantage of simplifying the calculations for parameter inference, thanks to a useful reformulation $\tilde{\beta} = \begin{pmatrix} \mu^\top \\ \beta \end{pmatrix} \in \mathcal{M}_{(p+1) \times q}$ and, for all $1 \leq i \leq n$, $\tilde{V}_i = (\tilde{V}_{i\ell'})_{1 \leq \ell' \leq p+1} = (1, V_i)^\top \in \mathbb{R}^{p+1}$, so that $\mu + \beta^\top V_i = \tilde{\beta}^\top \tilde{V}_i$. Let $\tilde{V} = (\tilde{V}_{i\ell'})_{1 \leq i \leq n; 1 \leq \ell' \leq p+1}$ such as $\varphi = \tilde{V}\tilde{\beta} + \xi$, where $\xi = \begin{pmatrix} \xi_1^\top \\ \vdots \\ \xi_n^\top \end{pmatrix} \in \mathcal{M}_{n \times q}$. Then, by introducing $\tilde{\delta} = (\tilde{\delta}_{\ell'm})_{1 \leq \ell' \leq p+1; 1 \leq m \leq q}$ such as:

$$\forall 1 \leq \ell' \leq p+1, \forall 1 \leq m \leq q, \tilde{\delta}_{\ell'm} = \begin{cases} 1 & \text{if } \ell' = 1, \\ \delta_{\ell'm} & \text{where } \ell = \ell' - 1, \text{ if } \ell' > 1, \end{cases}$$

to force the inclusion of the intercept in the model, Equations (5.4a) and (5.4b) can be rewritten as: for $1 \leq \ell' \leq p+1, 1 \leq m \leq q$,

$$\pi(\tilde{\beta}_{\ell'm} | \tilde{\delta}_{\ell'm}) = \mathcal{N}(0, (1 - \tilde{\delta}_{\ell'm})\nu_0 + \tilde{\delta}_{\ell'm}(\mathbb{1}_{\ell' > 1}\nu_1 + \mathbb{1}_{\ell'=1}\sigma_\mu^2)) \quad (5.5)$$

For the variance parameter σ^2 , an inverse-gamma prior is chosen (5.4c), which prohibits negative values. One possibility is to set $\nu_\sigma, \lambda_\sigma$ equal to 1 for example, to make it relatively non-influential. For the inter-individual variance-covariance matrix Γ , the inverse-Wishart prior is chosen (5.4d), which is the multivariate extension of the inverse-Gamma density. The hyperparameter Σ_Γ is a positive matrix which can be specified as $\Sigma_\Gamma = kI_q$ with k chosen of comparable size to the maximum likely variance of $(\varphi_{im})_{1 \leq i \leq m}$ for all $m \in \{1, \dots, q\}$. The hyperparameter d , which is a degree of freedom, can be chosen as the smallest integer value ensuring the existence of $\mathbb{E}[\Gamma]$, that is $q + 2$.

Following Ročková and George (2014), the i.i.d. Bernoulli prior (5.4e) is used for the inclusion variable δ , where the hyperparameters $(\alpha_m)_m$ can be seen as the proportion of relevant covariates for each individual parameter, and a Beta distribution prior (5.4f) is chosen on each α_m , for $m \in \{1, \dots, q\}$. To encourage sparsity in the model, Castillo and van der Vaart (2012) suggest choosing a_m small and b_m large, for example $a_m = 1$ and $b_m = p$ for all m . In the following, α, a and b respectively denote the vectors $(\alpha_m)_{1 \leq m \leq q}$, $(a_m)_{1 \leq m \leq q}$ and $(b_m)_{1 \leq m \leq q}$. See Ročková and George (2014), Lique et al. (2017), Deshpande et al. (2019), for more details on these choices of priors and the choice of hyperparameters values.

5.3 Maximum a posteriori inference and thresholding

The purpose of this section is to discuss the estimation of $\Theta = (\theta, \alpha) = (\tilde{\beta}, \Gamma, \sigma^2, \alpha)$ in model (5.1) - (5.4). Recall that $\Xi = (\nu_0, \nu_1, \sigma_\mu^2, \nu_\sigma, \lambda_\sigma, \Sigma_\Gamma, d, a, b)$ are fixed hyperparameters. In the following, model (5.1) - (5.4) is called SSNLME (Spike-and-Slab Non-Linear Mixed-Effects) model. Note that φ , which is not observed, could be considered as a parameter to be estimated, included in Θ . However, to design a scalable inference scheme, we consider it as a latent variable which we marginalise out of the posterior. This enables us to use an EM-type approach which is considerably faster than a full MCMC approach (see details in Subsection 5.5.3).

The EM-type approach proposed in Section 5.4 requires to compute the maximum a posteriori (MAP) estimator for Θ :

$$\hat{\Theta}^{MAP} = \underset{\Theta \in \Lambda}{\operatorname{argmax}} \pi(\Theta | \mathbf{y}), \text{ with } \pi(\Theta | \mathbf{y}) = \frac{p_\Theta(\mathbf{y})\pi(\Theta)}{\int_\Lambda p_\Theta(\mathbf{y})\pi(\Theta)d\Theta}, \quad (5.6)$$

where $p_{\Theta}(\mathbf{y})$ and $\pi(\Theta)$ respectively denote the probability density of $\mathbf{y}$ conditionally to Θ , and the prior density of Θ , and Λ denotes the parameter space. However, since the individual parameters φ are marginalised out, the $p_{\Theta}(\mathbf{y})$ distribution is not explicit. Denoting $Z = (\varphi, \delta) \in \mathcal{Z}$ the latent variables, the marginalised posterior distribution $\pi(\Theta|\mathbf{y})$ takes the form:

$$\pi(\Theta|\mathbf{y}) = \int_{\mathcal{Z}} \pi(\Theta, Z|\mathbf{y}) dZ, \text{ with } \pi(\Theta, Z|\mathbf{y}) = \frac{p(\mathbf{y}|\Theta, Z)p(\Theta, Z)}{\int_{\mathcal{Z}} \int_{\Lambda} p(\mathbf{y}|\Theta, Z)p(\Theta, Z) d\Theta dZ},$$

where $p(\mathbf{y}|\Theta, Z)$ and $p(\Theta, Z)$ respectively designate the probability density of $\mathbf{y}$ conditionally to (Θ, Z) , and the joint distribution of Θ and Z .

Targeting only the maximum *a posteriori* replaces a sampling problem by an optimisation problem, which turns out to be much more scalable than exploring the full posterior. Equation (5.6) is an optimisation problem in an incomplete data model, which is gainfully tackled using the Stochastic Approximation version of the EM algorithm (SAEM, Delyon et al. (1999)). Often in the literature, the EM algorithm and its extensions are presented in the frequentist framework for the calculation of the maximum likelihood estimator (MLE). Nevertheless, these algorithms are also very well adapted to the computation of the MAP estimator (Dempster et al., 1977).

5.3.1 General description of SAEM algorithm

In this subsection, we consider the general framework of an incomplete data model with observations $\mathbf{y}$ and latent variables Z that characterise the distribution of observations. It is assumed that the density of the complete data $(\mathbf{y}, Z)$ is parameterised by Θ , which is unknown and associated with a prior $\pi(\Theta)$. The EM algorithm is iterative and allows to build a sequence $(\Theta^{(k)})_k$ of parameter estimates, which under certain regularity conditions converges to a local maximum of the observed posterior distribution

$$\pi(\Theta|\mathbf{y}) = \int \pi(\Theta, Z|\mathbf{y}) dZ,$$

(see Delyon et al. (1999) for more details). However, this integral is generally intractable and the idea is to maximise it by iteratively maximising an easier quantity:

$$Q(\Theta|\Theta') = \mathbb{E}_{Z|(\mathbf{y}, \Theta')} [\log(\pi(\Theta, Z|\mathbf{y})) | \mathbf{y}, \Theta'],$$

the conditional expectation of the complete log-posterior $\log(\pi(\Theta, Z|\mathbf{y}))$ given the observations $\mathbf{y}$ and the current value of the parameter estimates Θ' . However, the quantity $Q(\Theta|\Theta')$ does not always have a closed form. This is especially the case in non-linear mixed-effects models like SSNLME model. However, even though this expectation cannot be computed in closed-form, it can be approximated by simulation. One solution proposed by Wei and Tanner (1990) is to replace the E-step, *i.e.* the computation of the Q quantity, by a Monte-Carlo approximation based on a large number of independent simulations of the latent variables Z , called the MCEM algorithm. However, this large number of simulations is computationally expensive. The SAEM algorithm is an other alternative that replaces the E-step by a stochastic approximation based on a single simulation of the latent variables, which considerably reduces the computational cost (Delyon et al., 1999). More precisely, the E-step of the EM algorithm is replaced by two steps: a simulation step (S-step) and a stochastic approximation step (SA-step). Then the k -th iteration of the SAEM algorithm proceeds as follows:

1. **S-step (Simulation):** simulate a realisation $Z^{(k)}$ of the latent variables according to the conditional distribution $\pi(Z|\mathbf{y}, \Theta^{(k)})$.
2. **SA-step (Stochastic Approximation):** update the approximation $Q_{k+1}(\Theta)$ of $Q(\Theta|\Theta^{(k)})$ by a stochastic approximation method, according to:

$$Q_{k+1}(\Theta) = Q_k(\Theta) + \gamma_k(\log \pi(\Theta, Z^{(k)}|\mathbf{y}) - Q_k(\Theta)),$$

where $(\gamma_k)_k$ is a sequence of step sizes decreasing towards 0 such that $\forall k, \gamma_k \in [0, 1]$, $\sum_k \gamma_k = \infty$ and $\sum_k \gamma_k^2 < \infty$.

3. **M-step (Maximisation):** update the parameter value by computing:

$$\Theta^{(k+1)} = \operatorname{argmax}_{\Theta \in \Lambda} Q_{k+1}(\Theta).$$

Remark 2. If the model belongs to the curved exponential family, that is the complete log-posterior can be written as:

$$\log(\pi(\Theta, Z|\mathbf{y})) = -\psi(\Theta) + \left\langle S(\mathbf{y}, Z), \phi(\Theta) \right\rangle,$$

where ψ and ϕ denote two functions of Θ , with $\langle \cdot, \cdot \rangle$ denoting the scalar product, and $S(\mathbf{y}, Z)$ the minimal sufficient statistics of the model, then,

$$Q(\Theta|\Theta^{(k)}) = -\psi(\Theta) + \left\langle \mathbb{E}_{Z|(\mathbf{y}, \Theta^{(k)})}[S(\mathbf{y}, Z)|\mathbf{y}, \Theta^{(k)}], \phi(\Theta) \right\rangle.$$

It is therefore sufficient to focus on the minimal sufficient statistics instead of $Q(\Theta|\Theta^{(k)})$ itself. More precisely, the SA-step and M-step of the SAEM algorithm are replaced by:

- **SA-step:** update S_{k+1} , the stochastic approximation of $\mathbb{E}_{Z|(\mathbf{y}, \Theta^{(k)})}[S(\mathbf{y}, Z)|\mathbf{y}, \Theta^{(k)}]$, according to:

$$S_{k+1} = S_k + \gamma_k(S(\mathbf{y}, Z^{(k)}) - S_k).$$

- **M-step:** update the parameter value by computing:

$$\Theta^{(k+1)} = \operatorname{argmax}_{\Theta \in \Lambda} \left\{ -\psi(\Theta) + \langle S_{k+1}, \phi(\Theta) \rangle \right\}.$$

Note that theoretical convergence results of the SAEM algorithm are provided in [Delyon et al. \(1999\)](#) under the assumption that the model belongs to the curved exponential family. The SSNLME model belongs to this curved exponential family.

Note that the simulation step is not always directly feasible. This is particularly true in non-linear mixed-effects models since the conditional distribution of the latent variables knowing the observations and the current value of the parameters is known only to a nearest multiplicative constant. [Kuhn and Lavielle \(2004\)](#) proposed an alternative which consists in coupling the SAEM method with an MCMC procedure. Interestingly, convergence of the MCMC part at each S-step is not necessary and only a few MCMC iterations are required in practice (see [Kuhn and Lavielle \(2004\)](#) and [Kuhn and Lavielle \(2005\)](#) for practical and theoretical considerations about this algorithm). In the following, this extension is called MCMC-SAEM.

5.3.2 Central decomposition of the Q quantity in spike-and-slab non-linear mixed-effects models

In the following, the notations from Section 5.2 are used again, and $\|\cdot\|$ denotes the Euclidean norm on $\mathbb{R}^n$. The SSNLME model (5.1) - (5.4) is a particular latent variables model with $\mathbf{y} = (y_{ij})_{i,j}$ and $Z = (\boldsymbol{\varphi}, \boldsymbol{\delta})$. The aim here is to decompose the Q quantity of the SAEM algorithm in the particular case of the SSNLME model, allowing to describe an algorithm for computing the MAP estimator of Θ in the following subsection.

First, note that, by using the tower property of conditional expectation, the quantity $Q(\Theta|\Theta^{(k)})$ in model (5.1) - (5.4) is written as:

$$\begin{aligned} Q(\Theta|\Theta^{(k)}) &= \mathbb{E}_{(\boldsymbol{\varphi}, \boldsymbol{\delta})|(\mathbf{y}, \Theta^{(k)})} [\log(\pi(\Theta, \boldsymbol{\varphi}, \boldsymbol{\delta}|\mathbf{y})) | \mathbf{y}, \Theta^{(k)}] \\ &= \mathbb{E}_{\boldsymbol{\varphi}|(\mathbf{y}, \Theta^{(k)})} \left[\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)}) \middle| \mathbf{y}, \Theta^{(k)} \right], \end{aligned}$$

where

$$\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)}) = \mathbb{E}_{\boldsymbol{\delta}|(\boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)})} [\log(\pi(\Theta, \boldsymbol{\varphi}, \boldsymbol{\delta}|\mathbf{y})) | \boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)}]. \quad (5.7)$$

It is interesting to write $Q(\Theta|\Theta^{(k)})$ like this because $\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)})$ has a closed form. Indeed, Proposition 1 shows that $\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)})$ can be decomposed such as:

$$\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)}) = C + \tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) + \tilde{Q}_2(\alpha, \Theta^{(k)}), \quad (5.8)$$

where $\tilde{Q}_1$ and $\tilde{Q}_2$ have a closed form given in Proposition 1. Thus, the separability of (5.8) into two distinct functions, $\tilde{Q}_1$ which depends on $(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)})$ and $\tilde{Q}_2$ on $(\alpha, \Theta^{(k)})$, allows to update the estimations of θ and α independently from one another. Moreover, since $\tilde{Q}_2$ does not depend on $\boldsymbol{\varphi}$, Proposition 1 allows to write that:

$$Q(\Theta|\Theta^{(k)}) = C + \mathbb{E}_{\boldsymbol{\varphi}|(\mathbf{y}, \Theta^{(k)})} \left[\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) \middle| \mathbf{y}, \Theta^{(k)} \right] + \tilde{Q}_2(\alpha, \Theta^{(k)}). \quad (5.9)$$

However, even if $\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)})$ has a closed form, this is not the case of $Q(\Theta|\Theta^{(k)})$ because the function g is non-linear with respect to φ_i , and so $\pi(\boldsymbol{\varphi}|\mathbf{y}, \Theta^{(k)})$ is only known to a nearest multiplicative constant. Thus, it is necessary to use a stochastic approximation method to approximate $\mathbb{E}_{\boldsymbol{\varphi}|(\mathbf{y}, \Theta^{(k)})} \left[\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) \middle| \mathbf{y}, \Theta^{(k)} \right]$ in Equation (5.9). The originality of the present extension of the MCMC-SAEM algorithm is that it combines an exact computation $\tilde{Q}_2(\alpha, \Theta^{(k)})$ and a stochastic approximation of $\mathbb{E}_{\boldsymbol{\varphi}|(\mathbf{y}, \Theta^{(k)})} \left[\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) \middle| \mathbf{y}, \Theta^{(k)} \right]$ instead of a stochastic approximation of the entire quantity $Q(\Theta|\Theta^{(k)})$. This results in the combination of an exact EM algorithm and of an MCMC-SAEM algorithm for the estimation of α and θ respectively.

Also, let us notice that $\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)})$ takes an exponential form. Thus, according to Remark 2, it suffices to approximate stochastically $\mathbb{E}_{\boldsymbol{\varphi}|(\mathbf{y}, \Theta^{(k)})} [S(\mathbf{y}, \boldsymbol{\varphi}) | \mathbf{y}, \Theta^{(k)}]$ at SA-step (see Appendix 5.8.1.1 for more details about $S(\mathbf{y}, \boldsymbol{\varphi})$ and this exponential form).

Algorithm 6 in Appendix 5.8.1.1 summarises the proposed extension of the MCMC-SAEM algorithm for computing the MAP estimator of Θ in the SSNLME model.

Proposition 1. Consider $\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)})$ defined by Equation (5.7) where $\Theta = (\tilde{\beta}, \Gamma, \sigma^2, \alpha)$. Then:

$$\tilde{Q}(\mathbf{y}, \boldsymbol{\varphi}, \Theta, \Theta^{(k)}) = C + \tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) + \tilde{Q}_2(\alpha, \Theta^{(k)}),$$

where C is a normalisation constant which does not depend on Θ , and with:

$$\begin{aligned} \tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) = & -\frac{1}{2\sigma^2} \sum_{i=1}^n \sum_{j=1}^{n_i} (y_{ij} - g(\varphi_i, t_{ij}))^2 - \frac{n_{tot} + \nu_\sigma + 2}{2} \log(\sigma^2) - \frac{\nu_\sigma \lambda_\sigma}{2\sigma^2} \\ & - \frac{1}{2} \text{Tr} \left((\boldsymbol{\varphi} - \tilde{V} \tilde{\beta})^\top (\boldsymbol{\varphi} - \tilde{V} \tilde{\beta}) \Gamma^{-1} \right) - \frac{1}{2} \sum_{m=1}^q \sum_{\ell'=1}^{p+1} \tilde{\beta}_{\ell'm}^2 \tilde{d}_{\ell'm}^*(\Theta^{(k)}) \\ & - \frac{n + d + q + 1}{2} \log(|\Gamma|) - \frac{1}{2} \text{Tr}(\Sigma_\Gamma \Gamma^{-1}) \end{aligned}$$

where $|A|$ and $\text{Tr}(A)$ respectively denote the determinant and the trace of a matrix A , and

$$\begin{aligned} \tilde{Q}_2(\alpha, \Theta^{(k)}) = & \sum_{m=1}^q \log \left(\sqrt{\frac{\nu_0}{\nu_1}} \frac{\alpha_m}{1 - \alpha_m} \right) \sum_{\ell=1}^p p_{\ell m}^*(\Theta^{(k)}) + \\ & (a_m - 1) \log(\alpha_m) + (p + b_m - 1) \log(1 - \alpha_m). \end{aligned}$$

Quantities $p_{\ell m}^*(\Theta^{(k)})$, $1 \leq \ell \leq p$, $1 \leq m \leq q$, and $\tilde{d}_{\ell'm}^*(\Theta^{(k)})$, $1 \leq \ell' \leq p+1$, $1 \leq m \leq q$, are defined as follows:

$$p_{\ell m}^*(\Theta^{(k)}) = \mathbb{E}[\delta_{\ell m} | \boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)}] = \frac{\alpha_m^{(k)} \phi_{\nu_1}(\beta_{\ell m}^{(k)})}{\alpha_m^{(k)} \phi_{\nu_1}(\beta_{\ell m}^{(k)}) + (1 - \alpha_m^{(k)}) \phi_{\nu_0}(\beta_{\ell m}^{(k)})} \quad (5.10)$$

where $\phi_\nu(\cdot)$ is the normal density with zero mean and variance ν , and

$$\begin{aligned} \tilde{d}_{\ell'm}^*(\Theta^{(k)}) = & \mathbb{E} \left[\frac{1}{(1 - \tilde{\delta}_{\ell'm}) \nu_0 + \tilde{\delta}_{\ell'm} (\mathbb{1}_{\ell' > 1} \nu_1 + \mathbb{1}_{\ell' = 1} \sigma_\mu^2)} \middle| \boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)} \right] \\ = & \frac{1}{\sigma_\mu^2} \mathbb{1}_{\ell' = 1} + \left(\frac{1 - p_{\ell m}^*(\Theta^{(k)})}{\nu_0} + \frac{p_{\ell m}^*(\Theta^{(k)})}{\nu_1} \right) \mathbb{1}_{\ell' \neq 1} \end{aligned} \quad (5.11)$$

where $\ell = \ell' - 1$.

Remark 3. Note that $\mathbb{E}[\delta_{\ell m} | \boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)}] = \mathbb{E}[\delta_{\ell m} | \Theta^{(k)}]$ because the posterior distribution of δ given $(\boldsymbol{\varphi}, \mathbf{y}, \Theta^{(k)})$ depends on $\mathbf{y}$ and $\boldsymbol{\varphi}$ only through the current estimates $\Theta^{(k)}$.

Remark 4. Note that for a linear mixed-effects model, that is when g is linear with respect to φ_i , a classical EM algorithm is applicable.

5.3.3 Estimator thresholding

As in Ročková and George (2014), after obtaining an estimator $\hat{\Theta}^{MAP}$, the support S^* , defined in Equation (5.3), can be naturally estimated as the most probable model conditionally on $\hat{\Theta}^{MAP}$. Indeed, for all $m \in \{1, \dots, q\}$ and for all $\ell \in \{1, \dots, p\}$, since $\pi(\delta_{\ell m} = 1 | \hat{\alpha}_m^{MAP}) = \hat{\alpha}_m^{MAP}$ and $\pi(\delta_{\ell m} = 0 | \hat{\alpha}_m^{MAP}) = 1 - \hat{\alpha}_m^{MAP}$, the *a posteriori* inclusion probability of the covariate ℓ knowing $\hat{\Theta}^{MAP}$ for the individual parameter m can be obtained as:

$$\mathbb{P}(\delta_{\ell m} = 1 | \mathbf{y}, \hat{\beta}_{\ell m}^{MAP}, \hat{\alpha}_m^{MAP}) = \frac{\pi_1(\hat{\beta}_{\ell m}^{MAP}) \hat{\alpha}_m^{MAP}}{\pi_1(\hat{\beta}_{\ell m}^{MAP}) \hat{\alpha}_m^{MAP} + \pi_0(\hat{\beta}_{\ell m}^{MAP}) (1 - \hat{\alpha}_m^{MAP})},$$

where $\pi_k(\hat{\beta}_{\ell m}^{MAP}) = \pi(\hat{\beta}_{\ell m}^{MAP} | \delta_{\ell m} = k)$ for $k \in \{0, 1\}$. Then, $\hat{\delta}$, which is the most probable δ knowing that $\Theta = \hat{\Theta}^{MAP}$, can be computed as follows:

$$\begin{aligned} \hat{\delta}_{\ell m} = 1 &\iff \mathbb{P}(\delta_{\ell m} = 1 | \mathbf{y}, \hat{\beta}_{\ell m}^{MAP}, \hat{\alpha}_m^{MAP}) \geq 0.5 \\ &\iff |\hat{\beta}_{\ell m}^{MAP}| \geq \sqrt{2 \frac{\nu_0 \nu_1}{\nu_1 - \nu_0} \log \left(\sqrt{\frac{\nu_1}{\nu_0}} \frac{1 - \hat{\alpha}_m^{MAP}}{\hat{\alpha}_m^{MAP}} \right)} = s_\beta(\nu_0, \nu_1, \hat{\alpha}_m^{MAP}). \end{aligned}$$

Note that this estimator can be seen as a local version of the median probability model of Barbieri and Berger (2004). Thus, the following subset of covariates for the individual parameter m is selected via a thresholding operation:

$$\hat{S} = \left\{ (\ell, m) \in \{1, \dots, p\} \times \{1, \dots, q\} \mid |\hat{\beta}_{\ell m}^{MAP}| \geq s_\beta(\nu_0, \nu_1, \hat{\alpha}_m^{MAP}) \right\}. \quad (5.12)$$

Remark 5. Note that threshold $s_\beta(\nu_0, \nu_1, \hat{\alpha}_m^{MAP})$ is the same for all the covariates but depends on the individual parameter m and on the values of the spike and slab hyperparameters ν_0 and ν_1 which act as tuning parameters for the penalty.

Remark 6. It is interesting to note that the thresholding rule is unchanged from the easier situation where the individual parameters φ_i 's would have been directly observed, which would have fit into the framework treated in Ročková and George (2014).

5.4 Covariate selection procedure

Similarly to Lasso regression (Tibshirani, 1996), it is interesting to exploit the flexibility of the spike-and-slab prior to study different levels of sparsity in the β 's columns, and thanks to the speed of the MCMC-SAEM algorithm, it is possible to explore a grid of values for the spike hyperparameter ν_0 rather than focusing on a single value. Indeed, mechanically, the higher ν_0 is, the fewer covariates are included in the estimated support of β 's columns. This is why it is more interesting to look at a grid of values and then use a model selection criterion to choose the optimal model. Let us denote Δ this grid, and $|\Delta|$ the number of grid points. Then, for all $\nu_0 \in \Delta$, the MCMC-SAEM algorithm is executed to obtain the MAP estimate of Θ , $\hat{\Theta}_{\nu_0}^{MAP}$, which is then used to determine a subset of relevant covariates for all individual parameters, $\hat{S}_{\nu_0}$, as explained by Equation (5.12) in Subsection 5.3.3. This first step reduces the total collection of 2^{pq} possible models to a smaller collection of $|\Delta| \ll 2^{pq}$ promising sub-models $(\hat{S}_{\nu_0})_{\nu_0 \in \Delta}$ with

high posterior probability. Next, a model selection criterion can be applied to choose the "best" model from this collection.

As explained in [Ročková and George \(2014\)](#), a possible criterion is to maximise, along the grid, the marginal posterior of δ under the prior with $\nu_0 = 0$. This corresponds to the so-called Dirac-and-slab prior, where the spike is a Dirac distribution ([Mitchell and Beauchamp, 1988](#)). However, in our case, it is not possible to have an explicit expression for this marginal and it is also difficult to obtain it numerically, so this criterion is not convenient.

However, as the collection of models has been reduced to a small sub-collection $(\hat{S}_{\nu_0})_{\nu_0 \in \Delta}$, that contains at most $|\Delta|$ models, an information criterion can be used to choose the final model. Thus, covariate selection would consist in choosing the "best" $\nu_0 \in \Delta$, that is noted $\hat{\nu}_0$, as:

$$\hat{\nu}_0 = \underset{\nu_0 \in \Delta}{\operatorname{argmin}} \left\{ \operatorname{crit}(\hat{S}_{\nu_0}) \right\}, \quad (5.13)$$

where

$$\operatorname{crit}(\hat{S}_{\nu_0}) = -2 \log \left(p(\mathbf{y}; \hat{\theta}_{\nu_0}) \right) + \operatorname{pen}(\nu_0), \quad (5.14)$$

with:

- $\log(p(\mathbf{y}; \theta))$ the log-likelihood of model (5.1),
- $\hat{\theta}_{\nu_0} = (\hat{\beta}_{\nu_0}, \hat{\Gamma}_{\nu_0}, \hat{\sigma}_{\nu_0}^2)$ a point estimator of the parameter $\theta = (\tilde{\beta}, \Gamma, \sigma^2)$ in sub-model $\hat{S}_{\nu_0}$,
- $\operatorname{pen}(\nu_0)$ a penalty function which penalizes the complexity of $\hat{S}_{\nu_0}$.

There are many ways to define the penalty in the criterion (5.14), e.g. AIC ([Akaike, 1973](#)), BIC ([Schwarz, 1978](#); [Delattre et al., 2014](#)), DIC ([Spiegelhalter et al., 2002](#)), etc. Here, a pragmatic and effective choice is to use the eBIC (extended Bayesian Information Criterion, [Chen and Chen \(2008\)](#)) which is tailored to the high-dimensional setting. Indeed, eBIC's penalty has the following form:

$$\operatorname{pen}_{\text{eBIC}}(\nu_0) = |\hat{S}_{\nu_0}| \times \log(n) + 2 \log \left(\binom{pq}{|\hat{S}_{\nu_0}|} \right), \quad (5.15)$$

where $|\hat{S}_{\nu_0}|$ is the size of this support, which allows to take into account that the number of possible models with $r \leq pq$ covariates increases quickly as r increases. The eBIC uses $\hat{\theta}_{\nu_0}$ the maximum likelihood estimate (MLE). Note that this MLE and log-likelihood that are required to compute criteria with a penalty of the form (5.15) do not have an explicit form here because the individual parameters are latent and the function g is non-linear with respect to φ_i . They are calculated using an MCMC-SAEM algorithm and importance sampling techniques respectively (see e.g. [Kuhn and Lavielle \(2005\)](#) and [Lavielle \(2014\)](#) for details). Note that some Bayesian model selection criteria could also be appropriate instead of eBIC.

The proposed variable selection procedure can be summarised as in Algorithm 7 in Appendix 5.8.1.2. In the following, this procedure is called SAEMVS (SAEM Variable Selection). A detailed example of the application of this procedure on a toy example is presented in a supporting web material given in the section "Data and code availability".

Remark 7. Note that for Algorithm 7 it is possible to parallelise the computations along the grid because the outputs of the algorithm for two given values of $\nu_0 \in \Delta$ do not depend on each other.

5.5 Numerical experiments

This section studies the performance of the global selection procedure, named SAEMVS, in the non-linear mixed effects model with spike-and-slab prior, called the SSNLME model. As a reminder, this procedure is summarised in Algorithm 7 in Appendix 5.8.1.2, and uses an MCMC-SAEM algorithm for inference, which is detailed in Algorithm 6 in Appendix 5.8.1.1.

The numerical study is divided into three parts. The first part is a comparison with strategies that can be easily implemented from existing methods. It aims essentially at demonstrating the interest of carrying out the selection of covariates from the data of all the individuals simultaneously thanks to the mixed effects model. The second part studies in great detail the influences of the number of subjects, the number of covariates, the signal-to-variability ratio and the collinearity between covariates on the performance of SAEMVS. To both show the flexibility of our approach and simplify the presentation of this study, this second part is conducted on another non-linear mixed effects model, this time with a one-dimensional random effect. In the third part, a comparison of SAEMVS in terms of computation time with an MCMC implementation is presented, to quantify precisely the speed improvement afforded by the SAEM algorithm.

5.5.1 Comparison with a two-step approach

The proposed approach is compared to a standard solution readily available to a practitioner using existing tools. The goal being to study the impact of covariates on specific parameters of the non-linear models, a manageable approach would be to proceed in two steps and fit independently the non-linear model to each individual, then perform variable selection using as dependent variables the estimated parameters. This second step can be carried out using for instance the popular **glmnet** package, which allows multivariate response variable selection. This strategy is expected to work fine in data-rich scenarios when each parameter can be estimated very precisely, but it loses the uncertainty on the estimated parameters and the shrinkage property of the mixed-effect model. We are not aware of alternative solutions with freely available code that a practitioner could use for a high-dimensional problem as discussed in the introduction. The interest of SAEMVS, which considers the mixed effect model and fully embraces the uncertainty in the estimations, is shown in the context of repeated measures. We show that SAEMVS performs better than the two-steps approach on two grounds:

- when the estimation of the individual parameters is more challenging, for instance in the classical scenario of lost to follow-up patients where some observation times may be missing,
- when different sets of covariates have an impact on the different dimensions of the response.

5.5.1.1 Model and simulation design

The following model, commonly used in pharmacokinetics, is considered:

$$\begin{cases} y_{ij} = \frac{D\varphi_{i1}}{\mathcal{V}}\varphi_{i1} - \varphi_{i2} \left(e^{-\frac{\varphi_{i2}}{\mathcal{V}} t_{ij}} - e^{-\varphi_{i1} t_{ij}} \right) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i = \mu + \boldsymbol{\beta}^\top V_i + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \end{cases} \quad (5.16)$$

where $\varphi_i = (\varphi_{i1}, \varphi_{i2})^\top$, namely $q = 2$. The constants D and $\mathcal{V}$ are set to 100 and 30 respectively. The aim is to obtain a set of active covariates for the two individual parameters. For this study, the parameters are set to: $n = 200$ individuals, $p = 500$ covariates, $\sigma^2 = 10^{-3}$, $\Gamma = \begin{pmatrix} 0.2 & 0.05 \\ 0.05 & 0.1 \end{pmatrix}$ so that the correlation between the two individual parameters is of the order of 0.35, $\mu = (6, 8)^\top$, and $\beta = \begin{pmatrix} 3 & 2 & 1 & 0 & 0 & 0 & \dots & 0 \\ 0 & 0 & 3 & 2 & 1 & 0 & \dots & 0 \end{pmatrix}^\top$. The individual covariates $V_i \in \mathbb{R}^p$, $1 \leq i \leq n$, are simulated independently according to a binomial distribution with a success probability of 0.2. Note that the support of the β is not the same for each dimension. It seems unrealistic that the covariates impacting φ_{i1} would be exactly the same as those impacting φ_{i2} . The covariates are standardised. Thus, 100 data-sets are simulated according to these parameter values where the number of observations by individual is fixed to $n_1 = \dots = n_n = 12$ with the following observation time points: $(t_{i1}, \dots, t_{i12}) = (0.05, 0.15, 0.25, 0.4, 0.5, 0.8, 1, 2, 7, 12, 24, 40)$.

This comparison was carried out on different scenarios, one corresponding to a data-rich situation where the two-step approach is expected to work relatively well, and a more challenging scenario where it is expected that there will be a difference between the two-step method and SAEMVS. The two scenarios correspond to different observation periods for each individual:

1. **Complete data-set.** This is the baseline scenario where all individuals are observed during the entire experiment.
2. **Partial observations.** For each $p_{\text{partial}} \in \{0.1, 0.2, 0.3, 0.4\}$, the other scenarios correspond to the case where $N_1 = p_{\text{partial}}n$ individuals are assumed to be no longer part of the experiment after the 3rd observation time, that is: for each previously simulated data-set, only the first 3 observation times are kept for the first N_1 individuals, and all observation times for the remaining $N_2 = n - N_1$ individuals.

5.5.1.2 Competing methods

For this comparison study, SAEMVS is compared to two other procedures. For both of these procedures, the first step consists in estimating the φ_i 's individual-by-individual with the method of least squares thanks to the **nlm** R function (Non-Linear Minimization, see [Schnabel et al. \(1985\)](#)). Then, the Lasso method from the **glmnet** R package ([Friedman et al., 2010](#)) is applied in its multivariate and univariate versions on the second level of the model (5.16) using the estimated φ_i 's:

- **Multivariate setting.** To take into account the correlation between the two individual parameters, the multi-response Gaussian family is used for the **glmnet** function on the estimated φ_i 's. Note that this function uses a group Lasso penalty, forcing the two individual parameters to have the same support. This method is called "mgaussian" in the following.
- **Univariate setting.** To circumvent this constraint, we also compare SAEMVS to the case where selection is made on the two individual parameters independently. For this, the **glmnet** function with the Gaussian family is used for each of the individual parameters separately on the estimated φ_{i1} 's and φ_{i2} 's. This method is called "gaussian" in the following.

For both of these methods, a cross-validation procedure is performed using the `cv.glmnet` function and the largest λ at which the mean squared error (MSE) is within one standard error of the smallest MSE is chosen as recommended in [Hastie et al. \(2009\)](#). The algorithmic settings of SAEMVS are given in Appendix 5.8.1.3 with an example of convergence graphs of the MCMC-SAEM algorithm.

5.5.1.3 Results

First, Table 5.1 compares the mean estimation errors for the two individual parameters using the first step of the two-step approach. It should be noted that 3 observation times appear sufficient to estimate the first parameter accurately, but insufficient for the second, as can be seen with the increasing estimation error for decreasing amount of data in Table 5.1.

Table 5.1: Comparison of the mean estimation errors for the first individual parameter (MEE1) and the second (MEE2) calculated on all individuals over the 100 data-sets using the first step of the two-step approach.

p_{partial}	0	0.1	0.2	0.3	0.4
MEE1 ^a	0.088	0.10	0.10	0.11	0.12
MEE2 ^b	0.12	0.62	1.14	1.68	2.20

^aMEE1 is the mean of the difference in absolute value between the true φ_{i1} and its estimate over all the individuals and the 100 data-sets.

^bMEE2 is the mean of the difference in absolute value between the true φ_{i2} and its estimate over all the individuals and the 100 data-sets.

Thus, it is expected that the number of partially observed individuals will have a negative impact on the estimated support of the second individual parameter. Indeed, this is what is observed in Figure 5.1 for gaussian and mgaussian methods. On this figure, for the first individual parameter (graph a), we can see that gaussian and mgaussian methods select a model that almost always includes the true support (striped bars). However, mgaussian method almost never selects the true model and the gaussian method only in one case out of 2. In contrast, SAEMVS selects exactly the right model (unpatterned bars) in a large majority of cases (about 90%) with no false positives. Note that in many applications, especially in biology, false positives are to be avoided due to the cost of the experiments, and therefore SAEMVS seems to be more efficient in this regard. For the second individual parameter (graph b), the methods gaussian and mgaussian suffer greatly from the increase in the number of partially observed individuals, whereas the mixed effect model structure in SAEMVS shows greater robustness thanks to the classical pooling of information phenomenon, whereby individuals with missing data can benefit from the remaining fully observed individuals.

5.5.2 Impacts of the different parameters and collinearity between covariates

5.5.2.1 Model

Now, the performance of SAEMVS is studied under a variety of scenarios. This exploration is carried out on a different non-linear model, highlighting our approach's flexibility. For ease

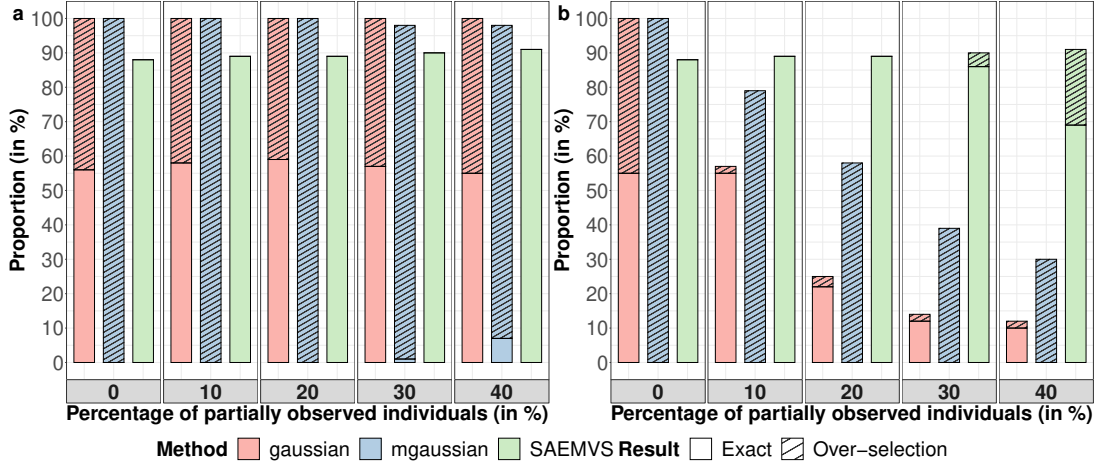


Fig. 5.1: Proportion of data-sets on which the three methods (in colour) select the correct model ("Exact", unpatterned bars), or a model that strictly includes the correct model ("Over-selection", striped bars) for the first individual parameter (a) and the second individual parameter (b), and different percentage of partially observed individuals (on the x -axis).

of presentation, we consider variable selection in a one-dimensional setting, *i.e.* $q = 1$. The other parameters are assumed to be shared among individuals and they are estimated jointly. A data-set is simulated according to a logistic growth model such as:

$$\begin{cases} y_{ij} = \frac{\psi_1}{1 + \exp\left(-\frac{t_{ij} - \varphi_i}{\psi_2}\right)} + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i = \mu + \beta^\top V_i + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \Gamma^2), \end{cases} \quad (5.17)$$

where $\psi = (\psi_1, \psi_2)$ is seen as unknown fixed effects, $\varphi_i \in \mathbb{R}$, $\mu \in \mathbb{R}$, $\beta \in \mathbb{R}^p$, and $\Gamma^2 > 0$. Thus, parameter ψ must also be estimated and therefore the population parameter is $\theta = (\mu, \beta, \psi, \Gamma^2, \sigma^2)$. The procedure presented earlier in this paper can easily be extended to this case.

Indeed, as the function g here is not separable into φ_i and ψ , the model does not belong to the curved exponential family since it is not possible to write $\tilde{Q}_1$ with an exponential form. As a result, the expression for the maximum argument in ψ at M-step is not explicit. One solution would be to do numerical optimisation in ψ . However, for ease of implementation of the MCMC-SAEM algorithm, following the idea of [Kuhn and Lavielle \(2005\)](#), an extended model belonging to the curved exponential family is used to estimate the parameters by considering:

$$\begin{cases} y_{ij} \stackrel{\text{ind.}}{\sim} \mathcal{N}(g(\varphi_i, \psi, t_{ij}), \sigma^2), \\ \varphi_i \stackrel{\text{ind.}}{\sim} \mathcal{N}(\mu + \beta^\top V_i, \Gamma^2), \\ \psi \sim \mathcal{N}(\eta, \Omega), \end{cases} \quad (5.18)$$

with φ and ψ independent, $\Omega = \text{diag}(\omega_1^2, \omega_2^2)$ known and $\theta^{ext} = (\mu, \beta, \eta, \sigma^2, \Gamma^2)$ the new population parameter to be estimated. The estimation of η is then used as an estimation of ψ . As

previously, the indicators $\delta = (\delta_\ell)_{1 \leq \ell \leq p}$ (Equation (5.2) with $q = 1$) are introduced, and we consider the same priors as in (5.4) for $(\mu, \beta, \sigma^2, \Gamma^2, \delta, \alpha)$ but in their one-dimensional version:

$$\begin{aligned} \pi(\beta|\delta) &= \mathcal{N}_p(0, D_\delta), \text{ with } D_\delta = \text{diag}((1 - \delta)\nu_0 + \delta\nu_1), 0 < \nu_0 < \nu_1, \\ \pi(\mu) &= \mathcal{N}(0, \sigma_\mu^2), \text{ with } \sigma_\mu^2 > 0, \\ \pi(\sigma^2) &= \mathcal{IG}\left(\frac{\nu_\sigma}{2}, \frac{\nu_\sigma \lambda_\sigma}{2}\right), \text{ with } \nu_\sigma, \lambda_\sigma > 0, \\ \pi(\Gamma^2) &= \mathcal{IG}\left(\frac{\nu_\Gamma}{2}, \frac{\nu_\Gamma \lambda_\Gamma}{2}\right), \text{ with } \nu_\Gamma, \lambda_\Gamma > 0, \\ \pi(\delta|\alpha) &= \alpha^{|\delta|}(1 - \alpha)^{p - |\delta|}, \text{ with } \alpha \in [0, 1] \text{ and } |\delta| = \sum_{\ell=1}^p \delta_\ell, \\ \pi(\alpha) &= \text{Beta}(a, b), \text{ with } a, b > 0. \end{aligned} \tag{5.19}$$

For η , the following prior is chosen: for $r \in \{1, 2\}$, $\pi(\eta_r) = \mathcal{N}(0, \rho_r^2)$, with $\rho_r^2 > 0$ known. This amounts to randomising hyperparameters of the prior on ψ , implying a less informative prior than if η were fixed. Here, $\Theta = (\theta^{ext}, \alpha)$ is the population parameter and $Z = (\varphi, \psi, \delta)$ are the latent variables. The steps of the MCMC-SAEM algorithm can be adapted to this model. The estimation method is unchanged for parameters $(\mu, \beta, \sigma^2, \Gamma^2, \alpha)$ because the quantity $\tilde{Q}_1$ is separable into $(\mu, \beta, \sigma^2, \Gamma^2, \alpha)$ and η . The main difference is that there is another latent variable ψ , which must also be simulated at the S-step. The thresholding procedure is not modified. See Appendix 5.8.1.4 for more details about this extension of SAEMVS.

Remark 8. *In order to limit the estimation error between the initial model and this extended model, the value of the covariance matrix is adapted during the iterations. Inspired by the results of Allasonnière and Debaveleere (2021) for the case of the computation of the MLE, the following process is chosen: start with a fairly large initial value $\Omega^{(0)} = \text{diag}(\omega_1^{2(0)}, \omega_2^{2(0)})$ for a certain number κ of iterations, then multiply it by $0 < \tau < 1$, and iterate this process every κ iterations. Starting from a large initial value, the value of Ω remains large enough during the first iterations to allow a rather fast convergence speed, then it is slowly decreased towards 0, while remaining always strictly positive, to limit the estimation error between the initial model and the extended model.*

5.5.2.2 Simulation design

For this simulation study, individual profiles are simulated according to model (5.17) by considering $n_1 = \dots = n_n = 10$ observations per individual and regular observation time points such that $t_{ij} = t_j = 150 + (j - 1)\frac{3000 - 150}{J - 1}$, $\sigma^2 = 30$, $\psi_1 = 200$, $\psi_2 = 300$, $\mu = 1200$, $\beta = (100, 50, 20, 0, \dots, 0)^\top$. Thus, only the first three covariates are assumed to be influential and their respective intensities are contrasted. The individual covariates $V_i \in \mathbb{R}^p$, $1 \leq i \leq n$, are simulated independently according to a centred multivariate Gaussian distribution with covariance matrix $\Sigma \in \mathcal{M}_p(\mathbb{R})$. To test the sensitivity of SAEMVS to the correlation that may exist between covariates, different scenarios are tested corresponding to different structures for matrix Σ . Different values of n (number of subjects) and p (number of covariates) are used according to the scenario. Several values of Γ^2 (variance of the random effects) are also used in order to evaluate the performances of SAEMVS in different "signal-to-noise ratio" situations.

- **Scenario with uncorrelated covariates.** This is the baseline scenario where optimal performance of Algorithm 7 is expected. This corresponds to $\Sigma = I_p$, where I_p is the

identity matrix of size p . The following values for n , p and Γ^2 are used: $n \in \{100, 200\}$, $p \in \{500, 2000, 5000\}$ and $\Gamma^2 \in \{200, 1000, 2000\}$.

• **Scenarios with correlations between covariates.**

1. The first scenario leaves the three influential covariates uncorrelated with all other covariates whereas the non-influential covariates are correlated with each other. An autoregressive correlation structure is considered between the non-influential covariates. This corresponds to $\Sigma = \left(\begin{array}{c|c} I_3 & 0_{3,p-3} \\ \hline 0_{p-3,3} & (\rho_\Sigma^{|i-j|})_{i,j \in \{4, \dots, p\}} \end{array} \right)$, with $|\rho_\Sigma| < 1$.
2. In the second scenario, the third influential covariate is assumed to be correlated to every non-influential covariate according to an autoregressive correlation structure. This corresponds to $\Sigma = \left(\begin{array}{c|c} I_3 & A \\ \hline A^\top & I_{p-3} \end{array} \right)$, with $A = \begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ & & (\rho_\Sigma^{|3-j|})_{j \in \{4, \dots, p\}} & \end{pmatrix}$, $|\rho_\Sigma| < 1$.
3. The third scenario considers correlations between the sole influential covariates. Again, an autoregressive correlation structure is used. This corresponds to $\Sigma = \left(\begin{array}{c|c} (\rho_\Sigma^{|i-j|})_{i,j \in \{1, \dots, 3\}} & 0_{3,p-3} \\ \hline 0_{p-3,3} & I_{p-3} \end{array} \right)$, $|\rho_\Sigma| < 1$.
4. In the fourth scenario, an autoregressive correlation structure is used between the covariates without making any distinction between the influential covariates and the non-influential covariates. This corresponds to $\Sigma = (\rho_\Sigma^{|i-j|})_{i,j \in \{1, \dots, p\}}$, $|\rho_\Sigma| < 1$.

To study the impact of correlations between covariates according to the four scenarios above, the following values for n , p , Γ^2 and ρ_Σ are used: $n = 200$, $p \in \{500, 2000, 5000\}$, $\Gamma^2 \in \{200, 2000\}$ and $\rho_\Sigma \in \{0.3, 0.6\}$.

For each of the five scenarios described above and each combination (n, p, Γ^2) or $(n, p, \Gamma^2, \rho_\Sigma)$, 100 different data-sets are simulated and the support of β is estimated by applying Algorithm 7 on each data-set. The algorithmic settings are given in Appendix 5.8.1.5. Note that, in order to be able to compare covariates that do not have the same order of magnitude, the covariates are standardised. The performances in terms of exact selection of the true influential covariates, over-selection and under-selection are examined (see Subsection 5.5.2.3).

5.5.2.3 Results

Scenario with uncorrelated covariates: The results are presented in Table 5.2 and Figure 5.2. SAEMVS selects exactly the right model in a large majority of cases for a sufficiently large number of individuals n . When n increases, the results improve, which suggests a consistency property in selection. With n and p fixed, the more the inter-individual variance Γ^2 is important, the more the results degrade. Indeed, as Γ^2 increases, the "signal-to-variability" ratio decreases, leading to difficulties in detecting the third covariate associated with the lowest non-zero coefficient in β . It could also be noted that with n and Γ^2 fixed, the results deteriorate when p increases but the effect of p seems weak when n is large. In addition, when SAEMVS

fails, it is most often because it under-selects, that is, it selects fewer variables than there are. Indeed, average sensitivity values are lower than those for specificity. It seems that SAEMVS tends to avoid false positives, even though this may result in not having selected all the truly influential covariates.

Table 5.2: Uncorrelated covariates. For each of the quantifiers (Sensitivity, Specificity, Accuracy), the mean over the 100 data sets is shown, with the empirical standard deviation divided by $\sqrt{100}$ in brackets.

Γ^2	p	$n = 100$			$n = 200$		
		Se ^a	Sp ^b	Ac ^c	Se ^a	Sp ^b	Ac ^c
200	500	0.883 (0.0180)	1 (5e-05)	0.999 (0.0001)	0.973 (0.0091)	1 (2e-05)	1 (6e-05)
200	2000	0.900 (0.0160)	1 (2e-05)	1 (3e-05)	0.987 (0.0066)	1 (2e-05)	1 (2e-05)
200	5000	0.870 (0.0189)	1 (6e-06)	1 (1e-05)	0.983 (0.0073)	1 (1e-05)	1 (1e-05)
1000	500	0.860 (0.0185)	1 (3e-05)	0.999 (0.0001)	0.977 (0.0086)	1 (3e-05)	1 (6e-05)
1000	2000	0.840 (0.0180)	1 (9e-06)	1 (3e-05)	0.977 (0.0086)	1 (1e-05)	1 (2e-05)
1000	5000	0.813 (0.0197)	1 (4e-06)	1 (1e-05)	0.963 (0.0105)	1 (6e-06)	1 (8e-06)
2000	500	0.810 (0.0185)	1 (3e-05)	0.999 (0.0001)	0.950 (0.0120)	1 (0)	1 (7e-05)
2000	2000	0.777 (0.0178)	1 (2e-05)	1 (3e-05)	0.937 (0.0131)	1 (1e-05)	1 (2e-05)
2000	5000	0.767 (0.0174)	1 (6e-06)	1 (1e-05)	0.930 (0.0137)	1 (4e-06)	1 (9e-06)

^aSensitivity is the proportion of true positives correctly identified.

^bSpecificity is the proportion of the true negatives correctly identified.

^cAccuracy is the proportion of true results, either true positive or true negative.

Scenarios with correlated covariates: The results for scenarios 1 and 4 with $\Gamma^2 = 200$ are presented in Figure 5.3 and Table 5.3 for the two values of ρ_Σ . The results for the other scenarios are given in Appendix 5.8.2. On these figures, one can compare the selection performance of SAEMVS in the different scenarios of correlations between covariates with the case without correlations (i.i.d scenario). First, for scenario 1, that is when the non-active covariates are correlated, quite similar performances to the i.i.d scenario are observed, but with more over-selection. Indeed, as a consequence of the correlation between irrelevant covariates, the latter tend to be selected more often and in small groups. Then, scenario 4 corresponds to a full correlation matrix between all covariates. Note that the correlation matrix chosen for this scenario assumes a fairly strong correlation between the three true covariates. Thus, as these active covariates explain the response variable in a similar way, this scenario leads to much under-selection compared to the i.i.d case. This can be seen in the sensitivity values. Moreover, it over-selects more than scenario i.i.d because of the correlations between the true and false covariates. However, by looking at the

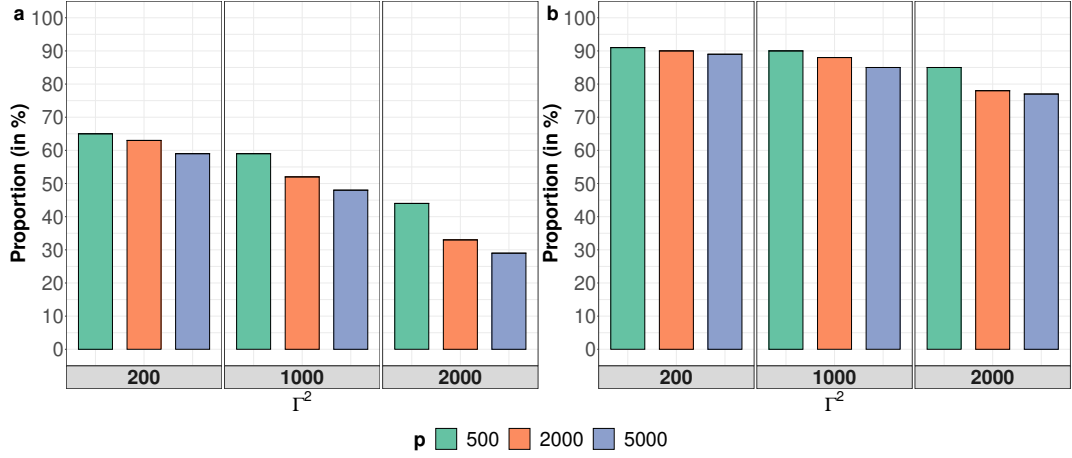


Fig. 5.2: Uncorrelated covariates. Proportion of data-sets on which SAEMVS selects the correct model for $n = 100$ (a) and $n = 200$ (b), and different values of p and Γ^2 .

specificity values, we can see that even in cases where SAEMVS over-selects, the false positive rate (which is equivalent to $(1 - \text{specificity})$) remains very close to 0.

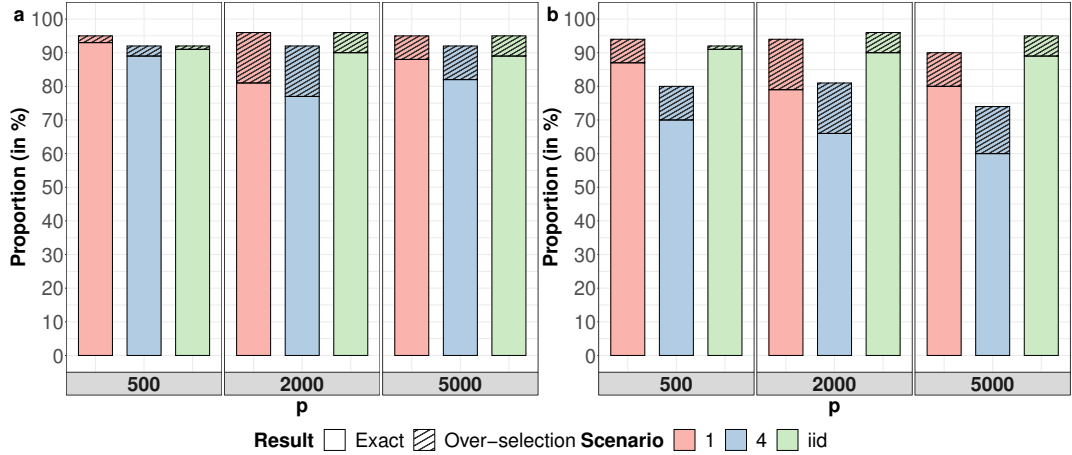


Fig. 5.3: Correlated covariates for $\Gamma^2 = 200$. Proportion of data-sets on which SAEMVS selects the correct model ("Exact", unpatterned bars) or a model that strictly includes the correct model ("Over-selection", striped bars) for $\rho_\Sigma = 0.3$ (a), $\rho_\Sigma = 0.6$ (b), and different values of p . Scenario "i.i.d" corresponds to the case where the covariates are not correlated and is used as a reference.

5.5.3 Comparison with an MCMC implementation

It is reasonably straightforward to implement an MCMC algorithm for full posterior inference on the spike-and-slab variable selection for non-linear mixed-effects model. To understand precisely the added value of the SAEM algorithm, we compare the run time of a full MCMC approach and the MCMC-SAEM method proposed in this paper, and highlight the better scaling properties

Table 5.3: Correlated covariates for $\Gamma^2 = 200$. For each of the quantifiers (Sensitivity, Specificity, Accuracy), the mean over the 100 data sets is shown, with the empirical standard deviation divided by $\sqrt{100}$ in brackets.

Scenario	p	$\rho_\Sigma = 0.3$			$\rho_\Sigma = 0.6$		
		Se ^a	Sp ^b	Ac ^c	Se ^a	Sp ^b	Ac ^c
i.i.d	500	0.973 (0.0091)	1 (2e-05)	1 (6e-05)	0.973 (0.0091)	1 (2e-05)	1 (6e-05)
1	500	0.983 (0.0073)	1 (3e-05)	1 (5e-05)	0.980 (0.0080)	1 (6e-05)	1 (8e-05)
4	500	0.973 (0.0091)	1 (4e-05)	1 (7e-05)	0.933 (0.0134)	1 (0.0001)	0.999 (0.0002)
i.i.d	2000	0.987 (0.0066)	1 (2e-05)	1 (2e-05)	0.987 (0.0066)	1 (2e-05)	1 (2e-05)
1	2000	0.987 (0.0066)	1 (3e-05)	1 (3e-05)	0.980 (0.0080)	1 (3e-05)	1 (3e-05)
4	2000	0.973 (0.0091)	1 (2e-05)	1 (3e-05)	0.937 (0.0131)	1 (3e-05)	1 (4e-05)
i.i.d	5000	0.983 (0.0073)	1 (1e-05)	1 (1e-05)	0.983 (0.0073)	1 (1e-05)	1 (1e-05)
1	5000	0.983 (0.0073)	1 (1e-05)	1 (1e-05)	0.967 (0.0101)	1 (9e-06)	1 (1e-05)
4	5000	0.973 (0.0091)	1 (8e-06)	1 (9e-06)	0.913 (0.0147)	1 (1e-05)	1 (2e-05)

^aSensitivity is the proportion of true positives correctly identified.

^bSpecificity is the proportion of the true negatives correctly identified.

^cAccuracy is the proportion of true results, either true positive or true negative.

of the latter. To build the most informative comparison, the same model with a smooth spike is considered for both the MCMC and MCMC-SAEM approaches, remarking that spike-and-slab priors with a Dirac spike are known to pose challenges for MCMC (see [Bai et al., 2021](#)). For the MCMC algorithm, an efficient C++ implementation of a random walk adaptive MCMC is used through the Nimble software ([de Valpine et al., 2017](#)), which uses an adaptive scheme proposed in [Shaby and Wells \(2010\)](#). To make the comparison as fair as possible, we marginalise the sampler over the discrete inclusion variables δ , to mirror the marginalisation in (5.7). This was found to appreciably improve the mixing of the MCMC algorithm. It is possible to retrieve the δ variables from the posterior samples using their conditional posterior distribution.

Common data-sets are simulated according to model (5.17) with the following parameters: $n = 200$ individuals, $n_1 = \dots = n_n = 10$ observations per individual, $p \in \{500, 700, 1000, 1500, 2000, 2500\}$ covariates, $\sigma^2 = 30$, $\psi_1 = 200$, $\psi_2 = 300$, $\mu = 1200$, $\beta = (100, 50, 20, 0, \dots, 0)^\top$ and $\Gamma^2 = 200$. For $i \in \{1, \dots, n\}$ and $j \in \{1, \dots, 10\}$, $t_{ij} = t_j = 150 + (j - 1) \frac{3000 - 150}{J - 1}$. Covariates are simulated independently and identically distributed according to $\mathcal{N}(0, 1)$. The objective is to compare the time needed to estimate the parameters $\Theta = (\mu, \beta, \psi, \Gamma^2, \sigma^2, \alpha)$ between the MCMC-SAEM algorithm proposed in this arti-

cle (Algorithm 6 adapted to model (5.17), see Appendix 5.8.1.4) and the full MCMC procedure described above. As explained in Subsection 5.5.2.1, to estimate the parameters, we consider the extended model (5.18). The same model structure (5.18) and priors are used for both approaches. For $(\mu, \beta, \Gamma^2, \delta, \alpha)$ the priors are as in (5.19). For η , the prior of Subsection 5.5.2.1 is chosen: for $r \in \{1, 2\}$, $\pi(\eta_r) = \mathcal{N}(0, \rho_r^2)$, with $\rho_r^2 > 0$ known. To stabilise the MCMC procedure, the prior on σ^2 is modified to a uniform distribution on $[0, 200]$ for both methods. This has very little consequence for SAEMVS. Indeed, the only difference lies in the updating of σ^2 at the M-step of the MCMC-SAEM algorithm, which becomes:

$$\sigma^{2(k+1)} = \begin{cases} \frac{s_{1,k+1}}{n_{tot}} & \text{if } \frac{s_{1,k+1}}{n_{tot}} \leq 200 \\ 200 & \text{else.} \end{cases}$$

The two methods are both initialised with: $\forall \ell \in \{1, \dots, 10\}$, $\beta_\ell^{(0)} = 100$, $\forall \ell \in \{11, \dots, p\}$, $\beta_\ell^{(0)} = 1$, $\mu^{(0)} = 1400$, $\sigma^{2(0)} = 100$, $\alpha^{(0)} = 0.1$ and $\eta^{(0)} = (400, 400)^\top$. In practice, to avoid convergence toward a local maximum in the MCMC-SAEM algorithm, a simulated annealing version of SAEM (see Lavielle, 2014) is implemented. Thus, in SAEMVS, Γ^2 is initialised very large to explore the space during the first iterations, with $\Gamma^{2(0)} = 5000$. For the full MCMC procedure, a more plausible value of Γ^2 , $\Gamma^{2(0)} = 500$, is chosen as initialisation. The hyperparameters are set in the same way for both methods as well: $\nu_0 = 0.04$, $\nu_1 = 12000$, $\sigma_\mu = 3000$, $\nu_\Gamma = \lambda_\Gamma = 1$, $a = 1$, $b = p$, $\Omega = \text{diag}(20, 20)$ and $\rho_1^2 = \rho_2^2 = 1200$.

For ease of presentation, we compare the two approaches for a single value of ν_0 . It is standard practice to run an MCMC spike-and-slab model for a single value (see for instance George and McCulloch (1997) or Malsiner-Walli and Wagner (2018)). The MCMC algorithm was run for 3000 iterations, which was just enough to reach convergence (assessed by comparing multiple chains) for a variety of ν_0 and p values. The MCMC-SAEM algorithm was run for 500 iterations and showed appropriate convergence. This convergence was assessed qualitatively by looking at the difference between successive parameter estimates during iterations. Under these conditions, for all $p \in \{500, 700, 1000, 1500, 2000, 2500\}$, both methods were run for 50 different data-sets and the minimum time was kept for each method. The results obtained are shown in Figure 5.4. In this figure, computation times of the full MCMC procedure (in purple) and of MCMC-SAEM (in blue) are represented by the points for the different values of p . The lines represent the regression line associated with each method. Note that a $\log_{10}$ - $\log_{10}$ scale is used in this figure. This shows that both methods have an execution time that grows polynomially with p . Furthermore, the polynomial complexity of the two methods, *i.e.* the slope of the regression lines, is slightly lower for the MCMC-SAEM method. Thus, if we note respectively τ_{MCMC} and $\tau_{MCMC-SAEM}$ the execution time associated with each of the methods under the conditions previously described, empirically Figure 5.4 strongly suggests that $\frac{\tau_{MCMC}}{\tau_{MCMC-SAEM}} \approx 10^{0.7} p^{0.2}$.

To sum up, the MCMC-SAEM algorithm proposed in this paper appears $10^{0.7} p^{0.2}$ times faster than the classical MCMC procedure, *i.e.* between 17 and 24 times faster for p between 500 and 2500. In other words, SAEMVS allows to browse a grid of about 20 values of the penalisation parameter ν_0 while a classical MCMC only looked at one value of this parameter.

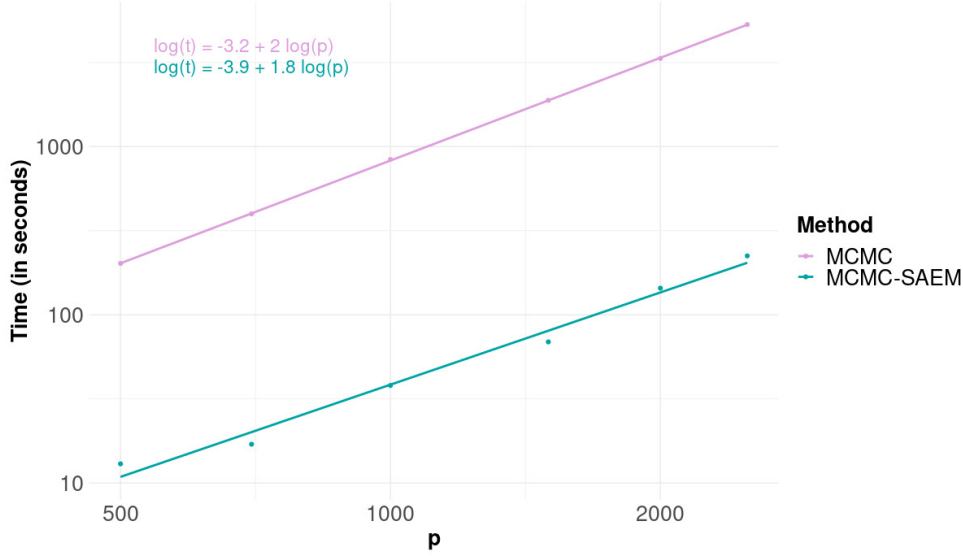


Fig. 5.4: Comparison of computation times between MCMC (in purple) and Algorithm 6 (MCMC-SAEM, in blue) inference methods in $\log_{10}$ - $\log_{10}$ scale.

5.6 Application to plant senescence genetic marker identification

In this section, SAEMVS is applied to a problem of marker-assisted selection to assist breeding of winter wheat. We are interested in identifying genetic markers that impact the senescence process, *i.e.* the ageing of these plants, which is under genetic determinism. Because heading date is related to senescence, the relevance of the selected senescence markers is then compared to known flowering genes as well as to markers associated to heading data obtained by applying an association mapping model (Yu et al., 2006) to these data. In the following, these markers are called heading QTLs. Note that the heading date can be seen as an easily measurable approximation of the flowering date.

5.6.1 Data description and pre-processing

The plant material used here has been previously described in Rincent et al. (2018, 2019) and Touzy et al. (2019). It is composed of $n = 220$ wheat varieties. Senescence was measured as the global proportion of senesced surface of the canopy. This proportion was observed on each variety $J = 18$ times over time.

For each of the n varieties, information from high-throughput genotyping is available on several tens of thousands of SNPs positioned along the entire genome (Rimbert et al., 2018) (see Appendix 5.8.3.1). These binary variables constitute the covariates to be selected in order to explain the differences between varieties in terms of senescence. There are structurally strong correlation and collinearity between these variables, which hampers variable selection, as pointed out in Section 5.5.2 and discussed in many references (see e.g. Malsiner-Walli and Wagner (2018) and Heuclin et al. (2020) in spike and slab models). To permit a comparison with the genetic markers associated with flowering in the context of collinearity, SAEMVS is applied chromosome by chromosome. Note, however, that variable selection on the whole genome (≈ 30000 markers)

would be computationally feasible in a matter of hours with strict pre-processing to address multicollinearity/near-multicollinearity. We still apply the following minimal pre-processing: if several SNPs in the same chromosome are strictly collinear for all varieties, only one of them is retained for analysis. This eliminates 916 SNPs across the genome. As a result, for a given chromosome C , the number p of covariates in V_i^C is always less than 2000, while remaining in a high-dimensional framework where $p \gg n$ (see Table 5.4 in Appendix 5.8.3.4 where the values of p are given for each chromosome). The wheat varieties in this data-set are structured into genetic groups, which may cause confusion between the effect of the subpopulation structure and that of SNPs. To control for subpopulation structure, we adapt our model to include covariates not subject to selection. We consider the first 5 principal components of a principal coordinates analysis performed on the available SNPs in the model. These 5 adjustment variables, capturing subpopulation structure, are denoted by $v_i \in \mathbb{R}^5$ for variety i and are guaranteed to be used in the regression.

5.6.2 Modelling

Some senescence curves are shown in Figure 9 in Appendix 5.8.3.2. We can see that a logistic growth model (Equation (5.17)) with a maximum value of 100% and inter-individual variability on the two other parameters is coherent with the shape of these curves. This is a simple model for the purposes of the example, but note that the variable selection method can accommodate other models like beta regression models with minimal changes. Here, we choose to analyse the effects of SNPs from each chromosome C on a single parameter of interest: the variability of characteristic times between varieties. Denoting y_{ij} the proportion of senesced surface of the plant of the variety $1 \leq i \leq n$ at time t_{ij} , this leads to the following model when focussing on chromosome C :

$$\begin{cases} y_{ij} = \frac{100}{1 + \exp\left(-\frac{t_{ij} - \varphi_i}{\psi_i}\right)} + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i = \mu + \lambda^\top v_i + \beta^\top V_i^C + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \Gamma^2), \\ \psi_i = \eta + \omega_i, & \omega_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \Omega^2). \end{cases}$$

The parameter to be estimated is therefore: $\theta = (\mu, \lambda, \beta, \eta, \sigma^2, \Gamma^2, \Omega^2)$. As previously, the indicators δ (Equation (5.2) with $q = 1$) are introduced and we force the inclusion of variables v_i in the model by a useful reformulation similar to what was done for the intercept in (5.5). We use the following priors for η and Ω^2 : $\pi(\eta) = \mathcal{N}(0, \sigma_\eta^2)$, with σ_η^2 known, and $\pi(\Omega^2) = \mathcal{IG}\left(\frac{\nu_\Omega}{2}, \frac{\nu_\Omega \lambda_\Omega}{2}\right)$, with $\nu_\Omega, \lambda_\Omega > 0$. The same priors as in (5.19) are used for the other parameters. Here, $\Theta = (\theta, \alpha)$ is the population parameter and $Z = (\varphi, \psi, \delta)$ are the latent variables, where $\varphi = (\varphi_i)_{1 \leq i \leq n}$ and $\psi = (\psi_i)_{1 \leq i \leq n}$. The SAEMVS procedure can be easily implemented with minor modifications to the version detailed in Appendix 5.8.1.4. The algorithmic settings are provided in Appendix 5.8.3.3.

5.6.3 Results and discussion

The results are shown in Figure 5.5 and further detailed in Table 5.4 in Appendix 5.8.3. For each chromosome, the SNPs selected by SAEMVS are compared to heading QTLs and major flowering

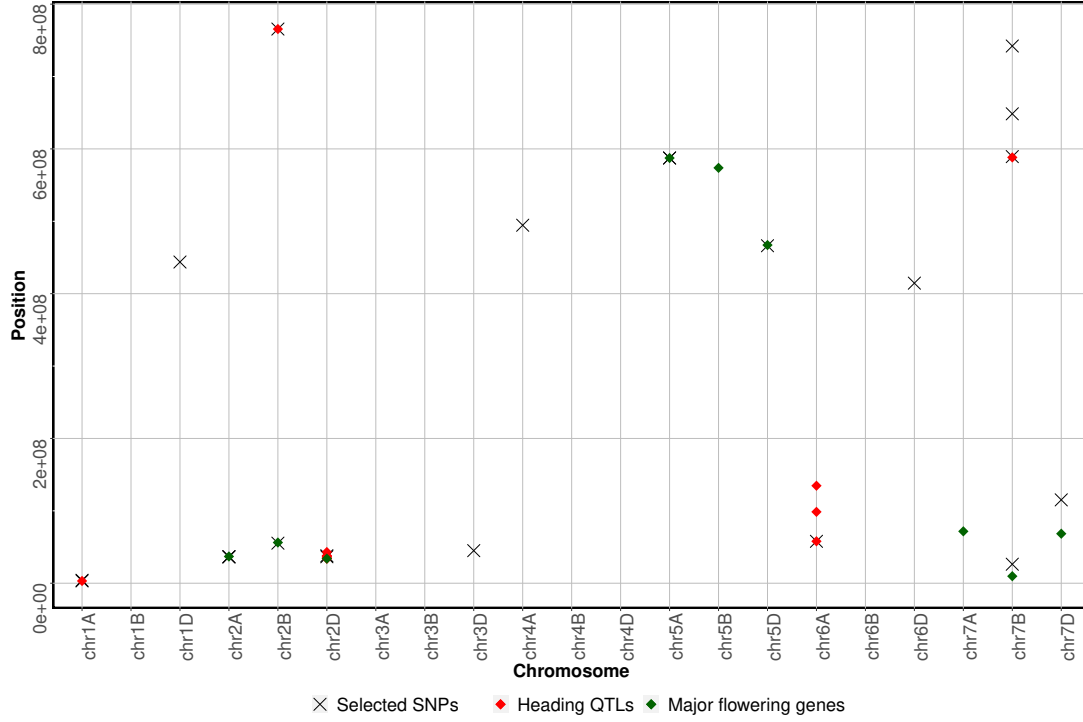


Fig. 5.5: Position on each chromosome of the markers selected by SAEMVS (in black cross), compared to heading QTLs (in red diamond) and major flowering genes (in green diamond).

genes that were identified in previous analysis (called Ppd and Vrn genes). Indeed, as mentioned above, as flowering and senescence are linked biological processes, we can expect to find SNPs that are close to these specific QTLs and genes on the genome. As an example, SAEMVS applied to chromosome 1A leads to the selection of two SNPs that are close to each other (*i.e.* within one mega-base of each other) but also close to a heading QTL on the chromosome. On a genome-wide basis, Figure 5.5 displays many colocalisations between flowering genes or heading QTLs and the SNPs selected by SAEMVS on the senescence data. Our procedure thus returns consistent results from a biological point of view. Conversely, SAEMVS also selects SNPs that are not close to zones associated with flowering on the genome, for example on chromosome 1D, 3D, 4A, 6D, 7B and 7D. The selected markers reveal potentially "stay-green" SNPs associated with late senescence independently of flowering time, which would be very interesting for plant breeders.

It is important to note that SAEMVS suffers from selection switch among markers that are highly correlated. This is illustrated by the chevron structure in Figure 5.6. On the left-hand plot, the red curves correspond to the selection threshold and the markers that are selected at least twice along the grid Δ of ν_0 values are denoted in different colours. In such situations, we expect any variable selection to face the same challenges, as one cannot hope to select markers truly associated with the phenomenon of interest. The usual approach in biology is to rather identify regions of the genome containing markers associated to the phenomenon of interest, usually via a post-processing of the variable selection results.

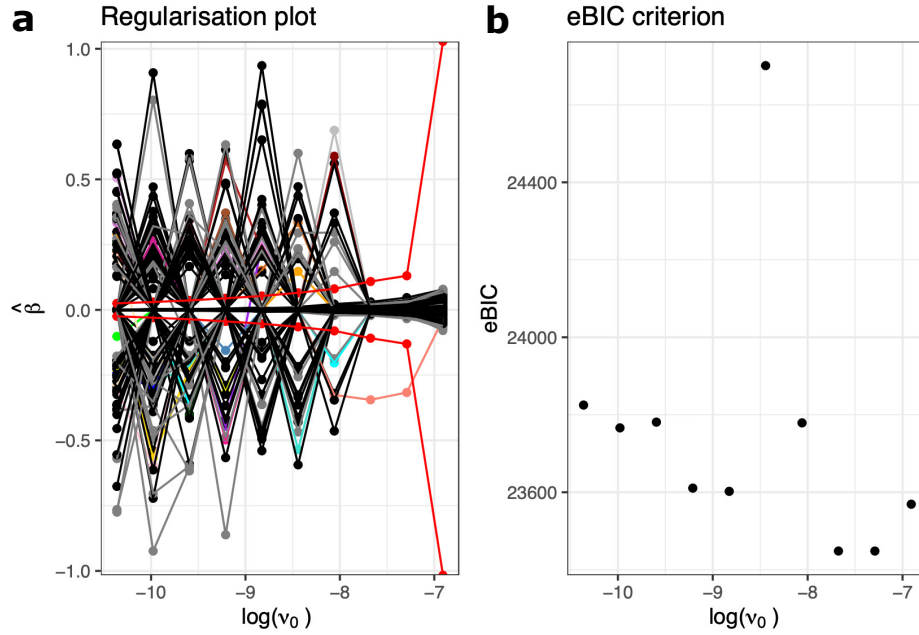


Fig. 5.6: Regularisation plot and eBIC criterion for chromosome 6A

5.7 Conclusion and perspectives

The main objective of this paper was to propose a new procedure for high-dimensional variable selection in non-linear mixed-effects models. In this work, variable selection was approached from a Bayesian perspective and a selection procedure combining the use of spike-and-slab Gaussian mixture prior and the SAEM algorithm was proposed. The spike-and-slab prior on the regression coefficients allows both the shrinkage towards zero of small non-significant coefficients through the spike distribution, while the largely uninformative slab distribution allows estimating influential covariates without bias from the penalisation. The speed of the SAEM algorithm allows exploring different levels of sparsity in the model through the variance of the spike distribution v_0 , which we observed to be beneficial in selecting sparse models for regression. Varying the level of sparsity provides a collection of good models among which we select the one minimising an eBIC criterion.

The SAEMVS method can do both one-dimensional and multidimensional variable selection, and it does not assume the same support for each dimension. It is very flexible and was illustrated on three different models. The proposed methodology showed very good selection performance on simulated data. Indeed, SAEMVS appears to select the right support in a large majority of cases. As expected, for different numbers of covariates p fixed, the right support is selected more often as the number of individuals n increases and the inter-individual variance Γ^2 decreases. Even more interesting, this method is much faster than an MCMC stochastic search alternative and can solve higher-dimensional variable selection problems. The application of SAEMVS on a real data-set shows, on the one hand, the flexibility of the procedure, and on the other hand, convincing results from a biological point of view despite strong correlations/multicollinearity between the covariates.

Moreover, it was observed that a reasonable correlation between covariates has little effect

on the selection performance of the proposed procedure. However, when the level of correlation becomes high, the performance decreases. This could be improved if structural information on the covariates were *a priori* known. Indeed, in this article, the i.i.d. Bernoulli prior on the indicators δ (5.4e), entails the assumption that each covariate has the same probability *a priori* of being included in the model. However, there are situations, such as genomic data, in which certain covariates are *a priori* more likely to be included together in the model. This *a priori* structural information on the covariates can be taken into account in SAEMVS by choosing a more flexible prior on δ . In Stingo et al. (2010) and Stingo and Vannucci (2011), authors propose the independent logistic regression prior or the Markov random field prior. This could also be considered in our methodology.

Another important remark is that, in this article, we considered a Gaussian distribution for $p(\mathbf{y}|\boldsymbol{\varphi}, \sigma^2)$ in model (5.1). It is possible to relax this assumption and consider larger distribution classes, such as discrete distributions like the Poisson distribution for example. The proposed methodology can therefore be applied in many contexts.

5.8 Appendices

5.8.1 Algorithms: synthesis and implementation details

5.8.1.1 MCMC-SAEM algorithm in SSNLME model

First, in Subsection 5.3.2, we notice that $\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)})$ takes an exponential form. More precisely, we have:

$$\tilde{Q}_1(\mathbf{y}, \boldsymbol{\varphi}, \theta, \Theta^{(k)}) = -\psi(\theta, \Theta^{(k)}) + \left\langle S(\mathbf{y}, \boldsymbol{\varphi}), \phi(\theta) \right\rangle, \quad (5.20)$$

with:

- $S(\mathbf{y}, \boldsymbol{\varphi}) = \left(\sum_{i,j} (y_{ij} - g(\varphi_i, t_{ij}))^2, \text{vec}(\boldsymbol{\varphi}^\top \boldsymbol{\varphi}), \text{vec}(\boldsymbol{\varphi}) \right)$
- $\phi(\theta) = \left(-\frac{1}{2\sigma^2}, -\frac{1}{2}\text{vec}(\Gamma^{-1}), \text{vec}(\tilde{V}\tilde{\beta}\Gamma^{-1}) \right)$
- $\psi(\theta, \Theta^{(k)}) = \frac{1}{2} \langle (\tilde{V}\tilde{\beta})^\top \tilde{V}\tilde{\beta}, \Gamma^{-1} \rangle + \frac{1}{2} \sum_{m=1}^q \sum_{\ell'=1}^{p+1} \tilde{\beta}_{\ell'm}^2 \tilde{d}_{\ell'm}^*(\Theta^{(k)}) + \frac{n_{tot} + \nu_\sigma + 2}{2} \log(\sigma^2) + \frac{n + d + q + 1}{2} \log(|\Gamma|) + \frac{\nu_\sigma \lambda_\sigma}{2\sigma^2} + \frac{1}{2} \text{Tr}(\Sigma_\Gamma \Gamma^{-1})$

where $\text{vec}(A)$ denotes the vectorisation of a matrix A . To simplify the formulas and using that for two matrix A and B , $\langle A, B \rangle := \text{Tr}(A^\top B) = \langle \text{vec}(A), \text{vec}(B) \rangle$, we denote by

$$(s_1(\mathbf{y}, \boldsymbol{\varphi}), s_2(\boldsymbol{\varphi}), s_3(\boldsymbol{\varphi})) = \left(\sum_{i,j} (y_{ij} - g(\varphi_i, t_{ij}))^2, \boldsymbol{\varphi}^\top \boldsymbol{\varphi}, \boldsymbol{\varphi} \right). \quad (5.21)$$

The use of the decomposition discussed in Subsection 5.3.2 leads to the following extension of the MCMC-SAEM algorithm, Algorithm 6, for computing the MAP estimator of Θ in the

SSNLME model (5.1) - (5.4), where Λ_θ denotes the parameter space restricted to θ , h is small (between 1 and 5), and K is usually in the order of a few hundred.

In practice, to allow more flexibility during the first iterations and thus to move away more quickly from the initial condition, it is usual to start the algorithm with n_{burnin} burn-in iterations, *i.e.* to use a step sizes sequence $(\gamma_k)_k$ of the form: $\gamma_k = 1$ for $0 \leq k \leq n_{\text{burnin}} - 1$ and $\gamma_k = 1/(k - n_{\text{burnin}} + 1)^\gamma$ for $n_{\text{burnin}} \leq k \leq K - 1$, where $\gamma \in]0.5, 1[$, $n_{\text{burnin}} < K$ with K the number of iterations of the SAEM algorithm (see Kuhn and Lavielle, 2005). Moreover, to avoid convergence toward a local maximum in the MCMC-SAEM algorithm, a simulated annealing version of SAEM (see Lavielle, 2014) is implemented.

Algorithm 6 MCMC-SAEM

Input: $K \in \mathbb{N}^*$, $\Theta^{(0)}$ initial parameter, hyperparameters vector $\Xi = (\nu_0, \nu_1, \sigma_\mu^2, \nu_\sigma, \lambda_\sigma, \Sigma_\Gamma, d, a, b)$, $S_0 = 0$ and $(\gamma_k)_k$ a step sizes sequence decreasing towards 0 such that $\forall k, \gamma_k \in [0, 1]$, $\sum_k \gamma_k = \infty$ and $\sum_k \gamma_k^2 < \infty$.

for $k = 0$ to $K - 1$ **do**

1. **S-Step:** simulate $\varphi^{(k)}$ using the result of h iterations of an MCMC procedure with $\pi(\varphi|\mathbf{y}, \Theta^{(k)})$ for target distribution.
2. **SA-Step:** for $u \in \{1, 2, 3\}$, compute $s_{u,k+1} = s_{u,k} + \gamma_k(s_u(\mathbf{y}, \varphi^{(k)}) - s_{u,k})$ with $s_u(\mathbf{y}, \varphi^{(k)})$ defined by (5.21).

3. **M-Step:** update $\theta^{(k+1)} = \underset{\theta \in \Lambda_\theta}{\operatorname{argmax}} \left\{ -\psi(\theta, \Theta^{(k)}) + \langle S_{k+1}, \phi(\theta) \rangle \right\}$ and

$\alpha^{(k+1)} = \underset{\alpha \in [0,1]^q}{\operatorname{argmax}} \tilde{Q}_2(\alpha, \Theta^{(k)})$, which reduces to the following explicit forms in model (5.1)-(5.4):

- $\operatorname{vec}(\tilde{\beta}^{(k+1)}) = \left(I_q \otimes \tilde{V}^\top \tilde{V} + (\Gamma^{(k)} \otimes I_{p+1}) \operatorname{diag}(\operatorname{vec}(\tilde{D}^*)) \right)^{-1} \operatorname{vec}(\tilde{V}^\top s_{3,k+1})$
- $\Gamma^{(k+1)} = \frac{\Sigma_\Gamma + s_{2,k+1} - (\tilde{V} \tilde{\beta}^{(k+1)})^\top s_{3,k+1} - s_{3,k+1}^\top \tilde{V} \tilde{\beta}^{(k+1)} + (\tilde{V} \tilde{\beta}^{(k+1)})^\top \tilde{V} \tilde{\beta}^{(k+1)}}{n + d + q + 1}$
- $\sigma^{2(k+1)} = \frac{\nu_\sigma \lambda_\sigma + s_{1,k+1}}{n_{\text{tot}} + \nu_\sigma + 2}$
- $\alpha_m^{(k+1)} = \frac{\sum_{\ell=1}^p p_{\ell m}^*(\Theta^{(k)}) + a_m - 1}{p + b_m + a_m - 2}$ for $1 \leq m \leq q$

where ψ and ϕ are defined by (5.20), and $\tilde{Q}_2(\alpha, \Theta^{(k)})$, $(p_{\ell m}^*(\Theta^{(k)}))_{1 \leq \ell \leq p; 1 \leq m \leq q}$ and $(\tilde{d}_{\ell' m}^*(\Theta^{(k)}))_{1 \leq \ell' \leq p+1; 1 \leq m \leq q}$ are defined in Proposition 1, with $\tilde{D}^* = (\tilde{d}_{\ell' m}^*(\Theta^{(k)}))_{1 \leq \ell' \leq p+1; 1 \leq m \leq q} \in \mathcal{M}_{(p+1) \times q}$.

end for

Output: $\hat{\Theta}^{MAP} = (\hat{\beta}^{MAP}, \hat{\Gamma}^{MAP}, \hat{\sigma}^{2MAP}, \hat{\alpha}^{MAP}) = (\tilde{\beta}^{(K)}, \Gamma^{(K)}, \sigma^{2(K)}, \alpha^{(K)})$.

5.8.1.2 Proposed variable selection procedure: SAEMVS

The proposed variable selection procedure SAEMVS can be summarised as in Algorithm 7.

Algorithm 7 SAEMVS procedure

Input: Δ a grid of ν_0 values, and all required arguments for MCMC-SAEM (Algorithm 6).

▷ **Reduce the model collection:**

for $\nu_0 \in \Delta$ **do**

1. Compute the MAP estimate $\hat{\Theta}_{\nu_0}^{MAP}$ by Algorithm 6.
2. Threshold the estimator $\hat{\beta}_{\nu_0}^{MAP}$ to define sub-model $\hat{S}_{\nu_0}$ according to Equation (5.12).

end for

▷ **Compute the eBIC criterion:**

for each unique sub-model among $(\hat{S}_{\nu_0})_{\nu_0 \in \Delta}$ **do**

1. Compute the MLE estimate $\hat{\theta}_{\nu_0}^{MLE}$ in sub-model $\hat{S}_{\nu_0}$ with an MCMC-SAEM algorithm.
2. Compute the log-likelihood $\log p(\mathbf{y}; \hat{\theta}_{\nu_0}^{MLE})$ with importance sampling techniques.
3. Compute the associated $\text{eBIC}(\hat{S}_{\nu_0})$ according to Equation (5.15).

end for

▷ **Identify the best level of sparsity:** compute $\hat{\nu}_0$ defined by Equation (5.13).

Output: $\hat{S}_{\hat{\nu}_0}$.

5.8.1.3 Algorithmic settings of SAEMVS for the comparison study

For the comparison study, the following settings are used for Algorithm 7.

- The hyperparameter values are set to $\nu_\sigma = \lambda_\sigma = 1$, $d = 4$, $\Sigma_\Gamma = 0.2I_2$, $a = (1, 1)^\top$, $b = (p, p)^\top$, $\sigma_\mu = 5$, $\nu_1 = 1000$, and the spike parameter ν_0 runs through a grid Δ defined as $\log_{10}(\Delta) = \left\{ -3 + k \times \frac{1}{3}, k \in \{0, \dots, 9\} \right\}$.
- The step sizes are defined with $\gamma = 2/3$, $n_{\text{burnin}} = 150$ and $K = 300$ as explained in Appendix 5.8.1.1.
- The MCMC-SAEM algorithm is initialised with: $\forall m \in \{1, 2\}, \forall \ell \in \{1, \dots, 10\}$ $\beta_{\ell m}^{(0)} = 1$, and $\forall m \in \{1, 2\}, \forall \ell \in \{11, \dots, p\}$ $\beta_{\ell m}^{(0)} = 0.1$, $\mu^{(0)} = (10, 10)^\top$, $\sigma^{2(0)} = 10^{-2}$, $\Gamma^{(0)} = \begin{pmatrix} 0.5 & 0.1 \\ 0.1 & 0.5 \end{pmatrix}$ and $\alpha^{(0)} = 0.5$. Note that different initialisations have been tested and have shown similar performances.

Figure 7 represents the convergence graphs of one run of the MCMC-SAEM algorithm for μ , some components of β , σ^2 , Γ and α . It is observed that the algorithm converges in a few iterations

for any parameter. Note that the parameters are all relatively correctly estimated, except for Γ but this was expected because of a over-fitting situation. Indeed, the underestimation of Γ can be explained by the fact that since $\nu_0 > 0$, none of the estimates of the coefficients of β is zero and therefore all the covariates are active in the model, which makes the variance estimation of the random effect tend towards 0.

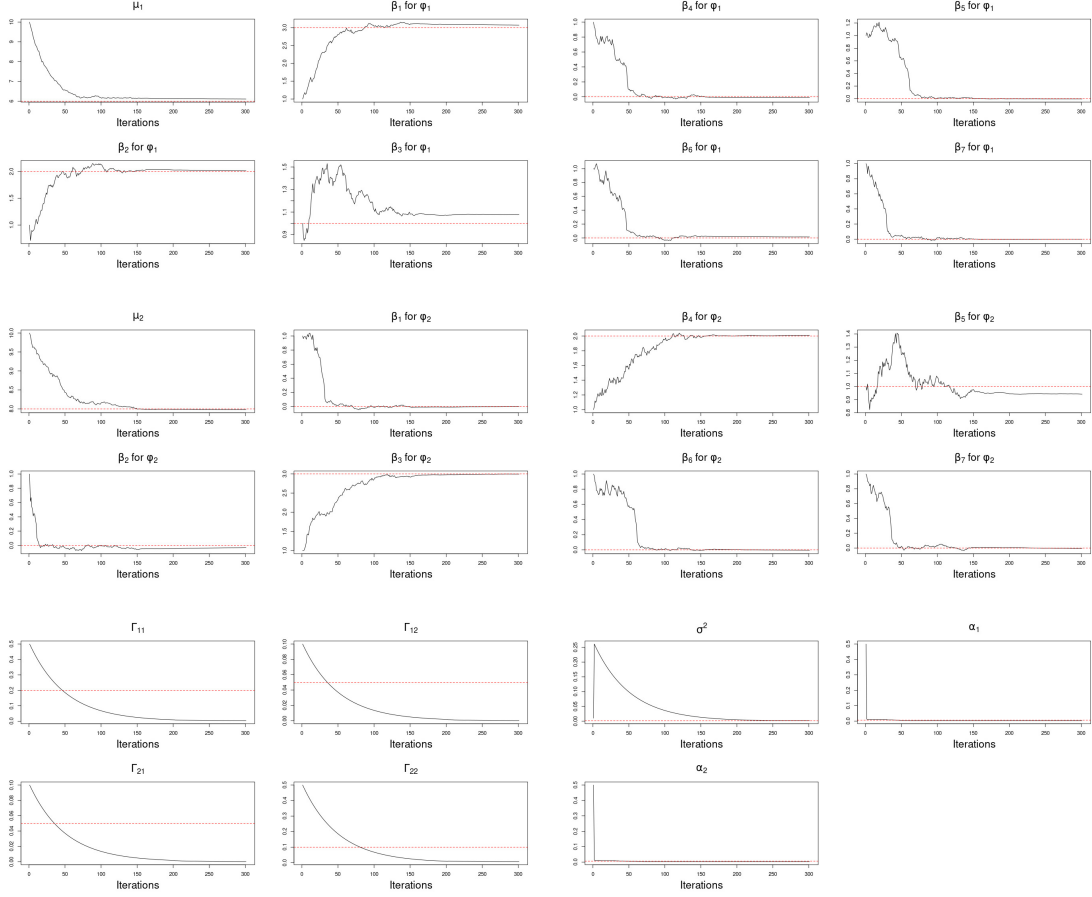


Fig. 7: Convergence graphs of the MCMC-SAEM algorithm for μ , some components of β , σ^2 , Γ and α on one simulated data-set, for $\nu_0 = 0.005$ and $\nu_1 = 1000$. The red dashed line corresponds to the true value of the considered parameter.

5.8.1.4 Extension in one-dimensional setting with estimation of fixed effects

As it is explained in Subsection 5.5.2.1, SAEMVS can easily be adapted to a model where fixed effects must be estimated. For this, the following general one-dimensional model is considered:

$$\begin{cases} y_{ij} & \overset{\text{ind.}}{\sim} \mathcal{N}(g(\varphi_i, \psi, t_{ij}), \sigma^2), \\ \varphi_i & \overset{\text{ind.}}{\sim} \mathcal{N}(\mu + \beta^\top V_i, \Gamma^2), \end{cases}$$

where $\psi \in \mathbb{R}^s$ are fixed effects to be estimated. Thus, by considering ψ as a latent variable with a normal distribution centred in η , an unknown parameter, and with a known covariance matrix Ω , and no longer as a parameter, it is possible to obtain an exponential form similar to Equation (5.20) for $\tilde{Q}_1$. The priors described in Subsection 5.5.2.1 in the particular case of the logistic growth model are used here, and we have $\Theta = (\mu, \beta, \eta, \sigma^2, \Gamma^2, \alpha)$ and $Z = (\varphi, \psi, \delta)$. The calculation of the quantity Q of the EM algorithm becomes:

$$\begin{aligned} Q(\Theta|\Theta^{(k)}) &= \mathbb{E}_{(\varphi, \psi, \delta)|(\mathbf{y}, \Theta^{(k)})} [\log(\pi(\Theta, \varphi, \psi, \delta|\mathbf{y}))|\mathbf{y}, \Theta^{(k)}] \\ &= \mathbb{E}_{(\varphi, \psi)|(\mathbf{y}, \Theta^{(k)})} \left[\tilde{Q}(\mathbf{y}, \varphi, \psi, \Theta, \Theta^{(k)}) \middle| \mathbf{y}, \Theta^{(k)} \right], \end{aligned}$$

where:

$$\begin{aligned} \tilde{Q}(\mathbf{y}, \varphi, \psi, \Theta, \Theta^{(k)}) &= \mathbb{E}_{\delta|(\varphi, \psi, \mathbf{y}, \Theta^{(k)})} [\log(\pi(\Theta, \varphi, \psi, \delta|\mathbf{y}))|\varphi, \psi, \mathbf{y}, \Theta^{(k)}] \\ &= C + \tilde{Q}_1(\mathbf{y}, \varphi, \psi, \theta, \Theta^{(k)}) + \tilde{Q}_2(\alpha, \Theta^{(k)}), \end{aligned}$$

with

$$\begin{aligned} \tilde{Q}_1(\mathbf{y}, \varphi, \psi, \theta, \Theta^{(k)}) &= -\frac{1}{2\sigma^2} \sum_{i,j} (y_{ij} - g(\varphi_i, \psi, t_{ij}))^2 - \frac{1}{2\Gamma^2} \|\varphi - \tilde{V}\tilde{\beta}\|^2 - \\ &\quad \frac{1}{2} \sum_{\ell'=1}^{p+1} \tilde{\beta}_{\ell'}^2 \tilde{d}_{\ell'}^*(\Theta^{(k)}) - \frac{n_{tot} + \nu_\sigma + 2}{2} \log(\sigma^2) - \frac{n + \nu_\Gamma + 2}{2} \log(\Gamma^2) - \\ &\quad \frac{\nu_\Gamma \lambda_\Gamma}{2\Gamma^2} - \frac{\nu_\sigma \lambda_\sigma}{2\sigma^2} - \sum_{r=1}^s \frac{(\psi_r - \eta_r)^2}{2\omega_r^2} - \sum_{r=1}^s \frac{\eta_r^2}{2\rho_r^2} \end{aligned}$$

and

$$\tilde{Q}_2(\alpha, \Theta^{(k)}) = \log \left(\sqrt{\frac{\nu_0}{\nu_1}} \frac{\alpha}{1 - \alpha} \right) \sum_{\ell=1}^p p_\ell^*(\Theta^{(k)}) + (a - 1) \log(\alpha) + (p + b - 1) \log(1 - \alpha).$$

$(p_\ell^*(\Theta^{(k)}))_{1 \leq \ell \leq p}$ and $(\tilde{d}_{\ell'}^*(\Theta^{(k)}))_{1 \leq \ell' \leq p+1}$ are defined in Proposition 1, Equations (5.10) and (5.11) with $q = 1$.

Note that $\tilde{Q}_1$ is still of the exponential form. Indeed,

$$\tilde{Q}_1(\mathbf{y}, \varphi, \psi, \theta, \Theta^{(k)}) = -\Psi(\theta, \Theta^{(k)}) + \left\langle S(\mathbf{y}, \varphi, \psi), \phi(\theta) \right\rangle \quad (5.22)$$

with:

- $S(\mathbf{y}, \varphi, \psi) = \left(\sum_{i,j} (y_{ij} - g(\varphi_i, \psi, t_{ij}))^2, \sum_{i=1}^n \varphi_i^2, \varphi, \psi^2, \psi \right)$
- $\phi(\theta) = \left(-\frac{1}{2\sigma^2}, -\frac{1}{2\Gamma^2}, \frac{\tilde{V}\tilde{\beta}}{\Gamma^2}, \left(-\frac{1}{2\omega_r^2} \right)_{1 \leq r \leq s}, \left(\frac{\eta_r}{\omega_r^2} \right)_{1 \leq r \leq s} \right)$

$$\bullet \Psi(\theta, \Theta^{(k)}) = \frac{\|\tilde{V}\tilde{\beta}\|^2}{2\Gamma^2} + \frac{1}{2} \sum_{\ell'=1}^{p+1} \tilde{\beta}_{\ell'}^2 \tilde{d}_{\ell'}^*(\Theta^{(k)}) + \frac{n_{tot} + \nu_\sigma + 2}{2} \log(\sigma^2) + \frac{n + \nu_\Gamma + 2}{2} \log(\Gamma^2) + \frac{\nu_\Gamma \lambda_\Gamma}{2\Gamma^2} + \frac{\nu_\sigma \lambda_\sigma}{2\sigma^2} + \sum_{r=1}^s \frac{\eta_r^2}{2\omega_r^2} + \sum_{r=1}^s \frac{\eta_r^2}{2\rho_r^2}$$

The k -th iteration of the MCMC-SAEM algorithm on this model is therefore:

1. **S-Step:** simulate $(\varphi^{(k)}, \psi^{(k)})$ using the result of some iterations of a Metropolis-Hastings within Gibbs algorithm with $\pi(\varphi, \psi | \mathbf{y}, \Theta^{(k)})$ for target distribution.
2. **SA-Step:** compute $S_{k+1} = S_k + \gamma_k (S(\mathbf{y}, \varphi^{(k)}, \psi^{(k)}) - S_k)$ with $S(\mathbf{y}, \varphi, \psi)$ defined by (5.22), where $S_k = (s_{1,k}, s_{2,k}, s_{3,k}, s_{4,k}, s_{5,k}) \in \mathbb{R} \times \mathbb{R} \times \mathbb{R}^n \times \mathbb{R}^q \times \mathbb{R}^q$.
3. **M-Step:** update $\theta^{(k+1)} = \underset{\theta \in \Lambda_\theta}{\operatorname{argmax}} \left\{ -\psi(\theta, \Theta^{(k)}) + \langle S_{k+1}, \phi(\theta) \rangle \right\}$ and $\alpha^{(k+1)} = \underset{\alpha \in [0,1]}{\operatorname{argmax}} \tilde{Q}_2(\alpha, \Theta^{(k)})$. More precisely,

$$\begin{aligned} \bullet \tilde{\beta}^{(k+1)} &= (\tilde{V}^\top \tilde{V} + \Gamma^{2(k)} \operatorname{diag}((\tilde{d}_{\ell'}^*(\Theta^{(k)}))_{1 \leq \ell' \leq p+1}))^{-1} \tilde{V}^\top s_{3,k+1}, \\ \bullet \Gamma^{2(k+1)} &= \frac{\|\tilde{V}\tilde{\beta}^{(k+1)}\|^2 + \nu_\Gamma \lambda_\Gamma + s_{2,k+1} - 2\langle s_{3,k+1}, \tilde{V}\tilde{\beta}^{(k+1)} \rangle}{n + \nu_\Gamma + 2}, \\ \bullet \sigma^{2(k+1)} &= \frac{\nu_\sigma \lambda_\sigma + s_{1,k+1}}{n_{tot} + \nu_\sigma + 2}, \\ \bullet \eta_r^{(k+1)} &= \frac{(s_{5,k+1})_r}{1 + \frac{\omega_r^2}{\rho_r^2}} \text{ for } 1 \leq r \leq s, \\ \bullet \alpha^{(k+1)} &= \frac{\sum_{\ell=1}^p p_\ell^*(\Theta^{(k)}) + a - 1}{p + b + a - 2}. \end{aligned}$$

As you can see, $\tilde{Q}_1$ is separable into $(\mu, \beta, \sigma^2, \Gamma^2, \alpha)$ and η , which means that the inference method used for the parameters $(\mu, \beta, \sigma^2, \Gamma^2, \alpha)$ is unchanged, *i.e.* the formulas to update these parameters in M-step are identical, it is only the way to simulate the sufficient statistics that has changed.

Thus, thanks to this algorithm, it is obtained an estimation $\hat{\theta}_{\nu_0}^{MAP}$, where $\hat{\theta}_{\nu_0}^{MAP} = (\hat{\mu}_{\nu_0}^{MAP}, \hat{\beta}_{\nu_0}^{MAP}, \hat{\eta}_{\nu_0}^{MAP}, \hat{\Gamma}_{\nu_0}^{2,MAP}, \hat{\sigma}_{\nu_0}^{2,MAP})$, and the estimation of η is used as an estimation of ψ . Then, to finish the model collection reduction step of SAEMVS, Algorithm 7, the estimator $\hat{\beta}_{\nu_0}^{MAP}$ is thresholded to obtain a promising sub-model $\hat{S}_{\nu_0}$ given by Equation (5.12). The selection threshold formula is unchanged because it only depends on the second layer of the model (5.17).

For the model selection step, to compute the eBIC criterion, it is also necessary to go through the extended model (5.18). Indeed, as described in Kuhn and Lavielle (2005), the MLE in the sub-model $\hat{S}_{\nu_0}$ is computed in the extended model by using an MCMC-SAEM algorithm and the estimation of η is used as an estimation of ψ . Then, the log-likelihood is approached by a Monte-Carlo method: for T large enough,

$$\log(p(\mathbf{y}; \hat{\theta}_{\nu_0}^{MLE})) \approx \sum_{i=1}^n \log \left(\frac{(2\pi \hat{\sigma}_{\nu_0}^{2,MLE})^{-n_i/2}}{T} \sum_{t=1}^T \exp \left(- \sum_{j=1}^{n_i} \frac{(y_{ij} - g(\varphi_i^{(t)}, \hat{\psi}_{\nu_0}^{MLE}, t_{ij}))^2}{2\hat{\sigma}_{\nu_0}^{2,MLE}} \right) \right)$$

where $p(\mathbf{y}; \theta)$ denotes the likelihood of model (5.17), and for all $i \in \{1, \dots, n\}$, $(\varphi_i^{(t)})_{t \in \{1, \dots, T\}}$ are simulated i.i.d. according to $p(\varphi_i; \hat{\theta}_{\nu_0}^{MLE}) = \mathcal{N}(\hat{\mu}_{\nu_0}^{MLE} + (\hat{\beta}_{\nu_0}^{MLE})^\top V_i, \hat{\Gamma}_{\nu_0}^{2MLE})$.

5.8.1.5 Algorithmic settings in the simulation study

For the simulation study, the following settings are used for Algorithm 7.

- The hyperparameter values are set to $\nu_\sigma = \lambda_\sigma = \nu_\Gamma = \lambda_\Gamma = 1$, $a = 1$, $b = p$, $\sigma_\mu = 3000$, $\rho_1^2 = \rho_2^2 = 1200$, $\nu_1 = 12000$, and the spike parameter ν_0 runs through a grid Δ defined as $\log_{10}(\Delta) = \left\{ -2 + k \times \frac{4}{19}, k \in \{0, \dots, 19\} \right\}$.
- The step sizes are defined with $\gamma = 2/3$, $n_{\text{burnin}} = 350$ and $K = 500$ as explained in Appendix 5.8.1.1.
- The MCMC-SAEM algorithm is initialised with: $\forall \ell \in \{1, \dots, 10\} \beta_\ell^{(0)} = 100$, $\forall \ell \in \{11, \dots, p\} \beta_\ell^{(0)} = 1$, $\mu^{(0)} = 1400$, $\sigma^{2(0)} = 100$, $\Gamma^{2(0)} = 5000$, $\alpha^{(0)} = 0.5$ and $\eta^{(0)} = (400, 400)^\top$. Note that different initialisations have been tested and have shown similar performances.
- At the beginning of the algorithm, $\Omega = \text{diag}(20, 20)$ and it is slowly reduced during the iterations as explained in Remark 8 with $\kappa = 40$ and $\tau = 0.9$.

5.8.2 Results for scenarios with correlated covariates

In this appendix, Figure 8, we give the results for scenarios not discussed in Section 5.5.2.3.

In scenario 2, it is assumed that the third relevant covariate is correlated to the non-active covariates. In this case, similar results to the i.i.d scenario are observed. Indeed, the selection performances of SAEMVS are only slightly affected by this scenario of correlations. This can be explained by the fact that, in this case, among the group of correlated covariates, the method will tend to select only one (or at least a very limited number of covariates among them): the most intense is chosen, *i.e.* the third true covariate. Next, scenario 3 describes correlations between the relevant covariates. Like the previous scenario, the procedure tends to select few covariates among the correlated covariates since they explain the response variable in a similar way. This also explains the degradation of the results when ρ_Σ increases.

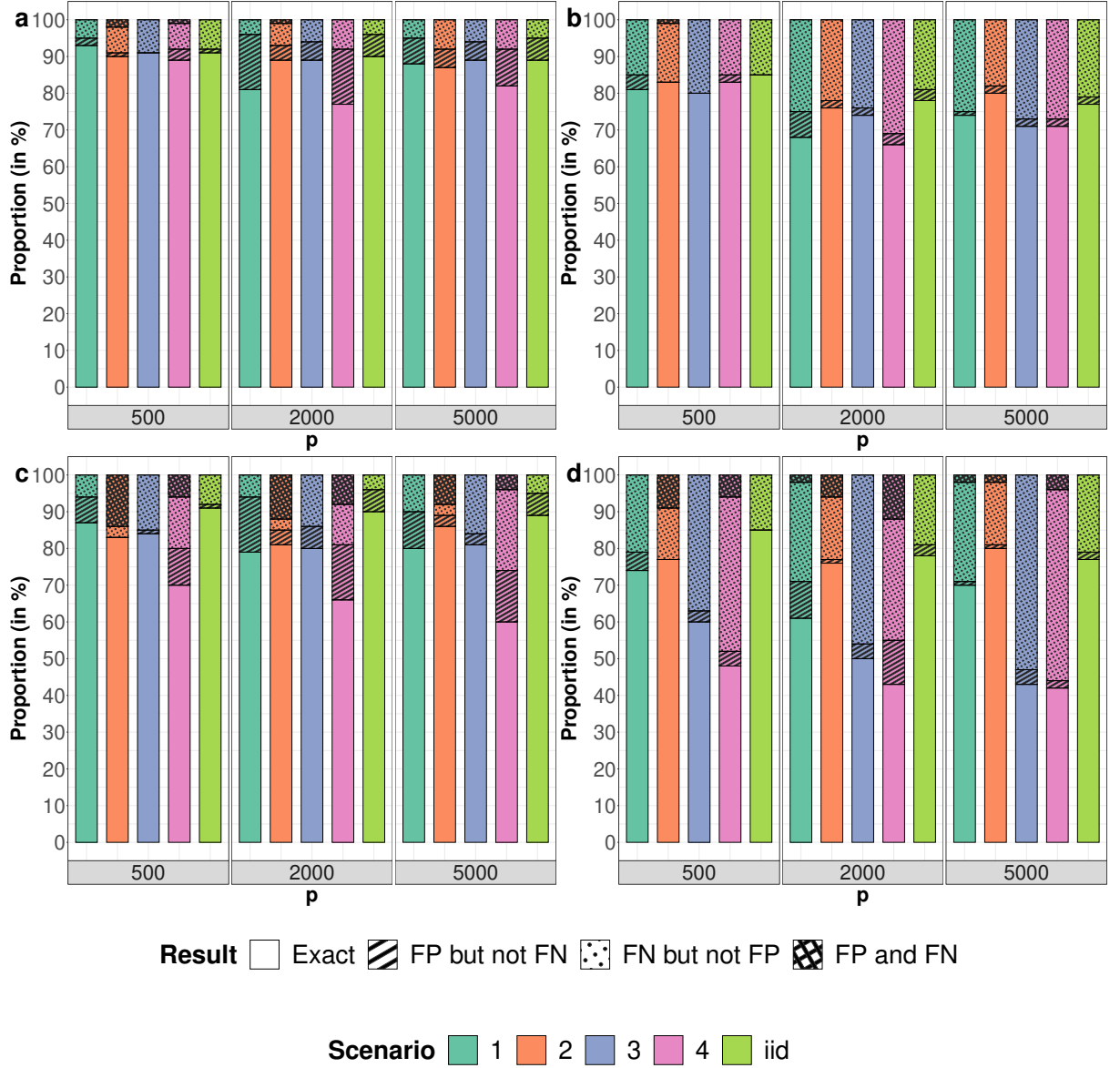


Fig. 8: Correlated covariates. Proportion of data-sets on which Algorithm 7 selects the correct model ("Exact", unpatterned bars), a model that contains false positives but not false negatives ("FP but not FN", striped bars), a model that contains false negatives but not false positives ("FN but not FP", dotted bars), or a model that contains both false positives and false negatives ("FP and FN", crosshatched bars) for $\rho_\Sigma = 0.3$ and $\Gamma^2 = 200$ (a), $\rho_\Sigma = 0.3$ and $\Gamma^2 = 2000$ (b), $\rho_\Sigma = 0.6$ and $\Gamma^2 = 200$ (c), and $\rho_\Sigma = 0.6$ and $\Gamma^2 = 2000$ (d), and different values of p . Scenario "i.i.d" corresponds to the case where the covariates are not correlated and is used as a reference.

5.8.3 Further details on real data

5.8.3.1 More details on the variety genotyping process

The varieties of the data-set used in this application was genotyped with the TaBW280K high-throughput genotyping array (Rimbert et al., 2018). This array was designed to cover both genic and intergenic regions of the three bread wheat subgenomes. Markers in strong Linkage Disequilibrium (LD) were filtered out using the pruning function of PLINK (Purcell et al., 2007) with a window of size 100 SNPs (Single Nucleotide Polymorphism, also called "molecular markers" in the following), a step of 5 SNPs and a LD threshold of 0.8, as proposed in Charmet et al. (2020). Missing values were imputed as the marker observed frequency, and then these imputed values are replaced by 0 or 1 using a threshold of 0.5. Monomorphic and unmapped markers were removed from the data-set. Eventually, we obtained $p = 26,189$ polymorphic high-resolution SNPs with a physical position on the v1 reference genome (International Wheat Genome Sequencing Consortium, 2018).

5.8.3.2 Representation of the data-set

Figure 9 shows part of the real data-set.

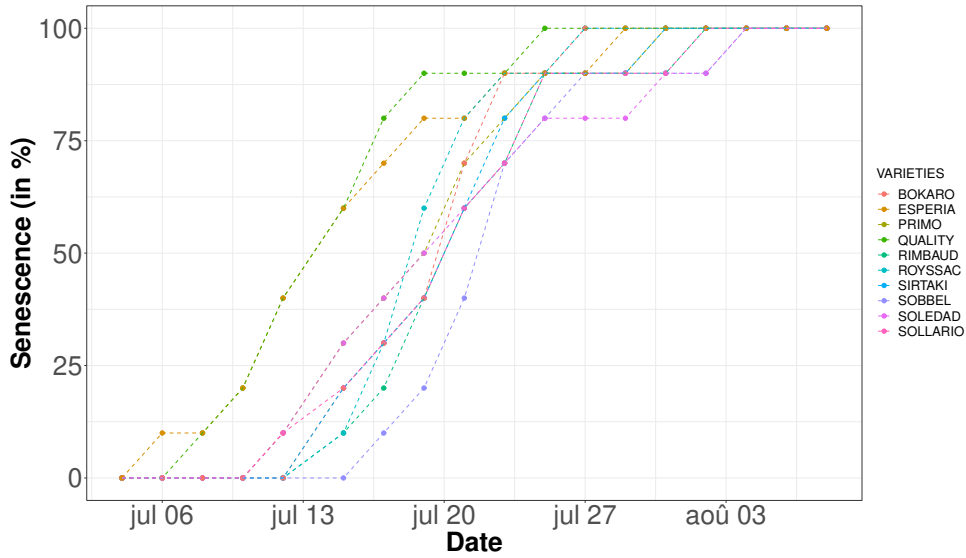


Fig. 9: Representation of the data-set for 11 different varieties

5.8.3.3 Settings

For application on real data, the following settings are used for Algorithm 7.

- The hyperparameter values are set to $\nu_\sigma = \lambda_\sigma = \nu_\Gamma = \lambda_\Gamma = \nu_\Omega = \lambda_\Omega = 1$, $a = 1$, $b = p$, $\sigma_\mu^2 = \sigma_\lambda^2 = \sigma_\eta^2 = 100$, $\nu_1 = 10$, and the spike parameter ν_0 runs through a grid Δ defined as $\log_{10}(\Delta) = \left\{ -4.5 + k \times \frac{1.5}{9}, k \in \{0, \dots, 9\} \right\}$.

- The step sizes are defined with $\gamma = 2/3$, $n_{\text{burnin}} = 250$ and $K = 400$ as explained in Appendix 5.8.1.1.
- The MCMC-SAEM algorithm is initialised with: $(\beta^{(0)})_\ell = 0.25$ if ℓ is a marker less than 1 mega base distance from a heading QTL or a major flowering gene, and otherwise $(\beta^{(0)})_\ell = 0.1$, $\mu^{(0)} = 20$, $\lambda^{(0)} = (0.25, \dots, 0.25)^\top$, $\sigma^{2(0)} = 80$, $\Gamma^{2(0)} = 50$, $\eta^{(0)} = 5$, $\Omega^{2(0)} = 50$, and $\alpha^{(0)} = 0.5$.

5.8.3.4 Table of results of the application on real data

Table 5.4 summarises for each chromosome, the number of covariates, the number of heading QTLs, the number of major flowering genes present in that chromosome, and the number of SNPs selected by SAEMVS.

Table 5.4: Summary of data and number of SNPs selected by SAEMVS for each chromosome

Chromosome	p	Number of heading QTLs	Number of flowering genes	Number of selected SNPs
1A	1473	1	0	2
1B	1604	0	0	0
1D	497	0	0	1
2A	1416	0	1	3
2B	1672	1	1	2
2D	696	7	1	2
3A	1477	0	0	0
3B	1888	0	0	0
3D	722	0	0	1
4A	1259	0	0	1
4B	961	0	0	0
4D	571	0	0	0
5A	1598	0	1	2
5B	1535	0	1	0
5D	772	0	1	1
6A	1119	3	0	1
6B	1317	0	0	0
6D	584	0	0	1
7A	1819	0	1	0
7B	1515	1	1	4
7D	778	0	1	1

5.9 Application on a detailed example

This section presents a toy example of applying the selection method proposed in this article on simulated data within a one-dimensional framework, where only a singular individual effect is present, that is $q = 1$. As a reminder, this procedure is summarised in Algorithm 7 in Appendix 5.8.1.2, and uses an MCMC-SAEM algorithm for inference, which is detailed in Algorithm 6 in Appendix 5.8.1.1.

A data-set is simulated according to a logistic growth model that is model (5.1) with:

$$g(\varphi_i, t_{ij}) = \frac{\psi_1}{1 + \exp\left(-\frac{t_{ij} - \varphi_i}{\psi_2}\right)},$$

where ψ_1 and ψ_2 are known constants. This is a common and realistic model used in many fields of life sciences, such as plant growth for example. Let us consider $n = 200$ individuals, $p = 500$ covariates and $n_1 = \dots = n_n = 10$ observations per individual. For all $i \in \{1, \dots, n\}$, for all $j \in \{1, \dots, n_i\}$, $t_{ij} = t_j = 150 + (j - 1)\frac{3000 - 150}{10 - 1}$. For each individual, the p covariates are simulated independently according to standard normal distributions $\mathcal{N}(0, 1)$. The parameter values are set to $\sigma^2 = 30$, $\psi_1 = 200$, $\psi_2 = 300$, $\mu = 1200$, $\beta = (100, 50, 20, 0, \dots, 0)^\top$ and $\Gamma^2 = 200$. Thus, only the first three covariates are influential, *i.e.* $S_\beta^* = \{1, 2, 3\}$.

5.9.1 Convergence of the MCMC-SAEM algorithm

Algorithm 6, is initialised here with: $\forall \ell \in \{1, \dots, 10\}$, $\beta_\ell^{(0)} = 100$, $\forall \ell \in \{11, \dots, p\}$, $\beta_\ell^{(0)} = 1$, $\mu^{(0)} = 1500$, $\sigma^{2(0)} = 100$, $\Gamma^{2(0)} = 5000$ and $\alpha^{(0)} = 0.5$. The hyperparameters are set to: $\nu_0 = 0.02$, $\nu_1 = 12000$, $\nu_\sigma = \lambda_\sigma = \nu_\Gamma = \lambda_\Gamma = 1$, $a = 1$, $b = p$ and $\sigma_\mu = 3000$.

Figure 10 represents the convergence graphs of one run of the MCMC-SAEM algorithm for μ , some components of β , σ^2 , Γ^2 and α . It is observed that the algorithm converges in a few iterations for any parameter. In this example, after 500 iterations, the algorithm returns $\hat{\beta}_1^{MAP} = 95.7$, $\hat{\beta}_2^{MAP} = 51.96$, $\hat{\beta}_3^{MAP} = 22.46$, $\hat{\mu}^{MAP} = 1202$, $\hat{\sigma}^{2MAP} = 33.86$, $\hat{\Gamma}^{2MAP} = 1.77$ and $\hat{\alpha}^{MAP} = 0.003$. Moreover, the estimates of the null fixed effects of the covariates are all less than 0.07 in absolute value. Note that the parameters are all relatively correctly estimated, except for Γ^2 but this was expected because of a over-fitting situation. Indeed, the underestimation of Γ^2 can be explained by the fact that since $\nu_0 > 0$, none of the estimates of the coefficients of β is zero and therefore all the covariates are active in the model, which makes the variance estimation of the random effect tend towards 0 in Equation (5.1b). This illustrates the need to threshold the estimators as described in Subsection 5.3.3.

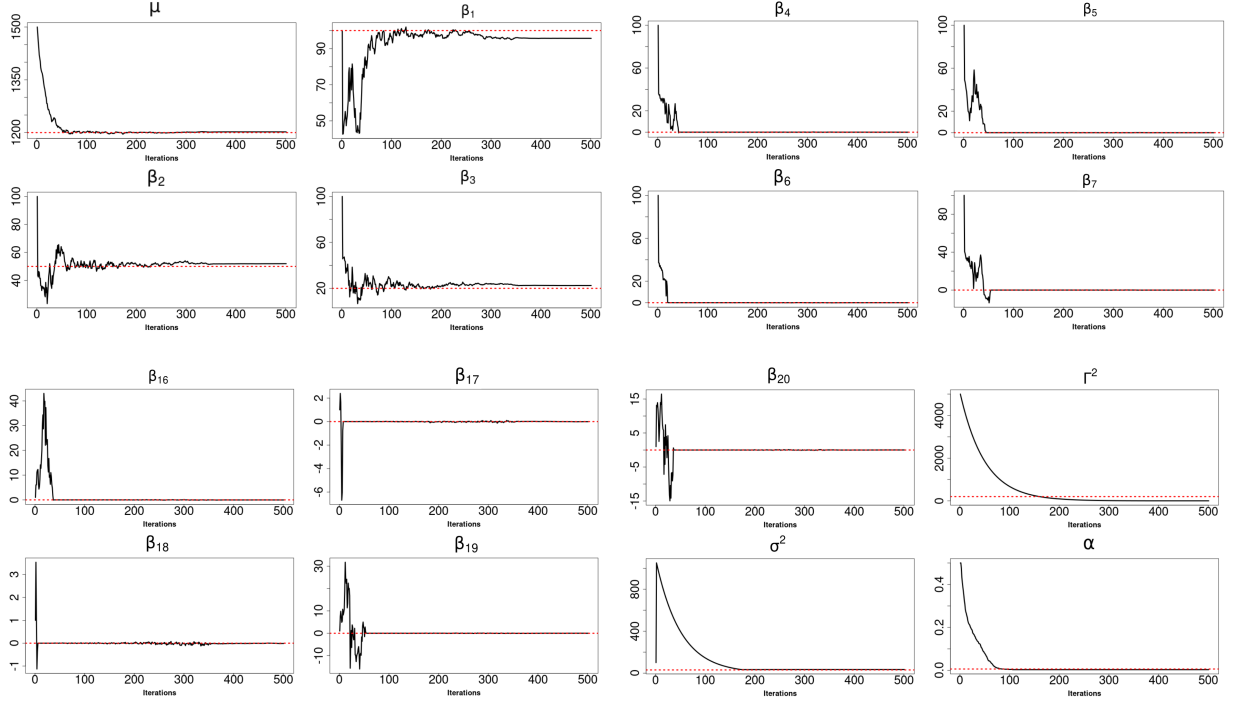


Fig. 10: Convergence graphs of the MCMC-SAEM algorithm for μ , some components of β , σ^2 , Γ^2 and α on one simulated data-set, for $\nu_0 = 0.02$ and $\nu_1 = 12000$. The red dashed line corresponds to the true value of the considered parameter.

5.9.2 Spike-and-slab regularisation plot and model selection

To illustrate how the variable selection works, the full procedure is applied on this simulated example on the grid of ν_0 values Δ such that

$$\log_{10}(\Delta) = \left\{ -2 + k \times \frac{4}{19}, k \in \{0, \dots, 19\} \right\}.$$

To have a visual representation of this procedure, a spike-and-slab regularisation plot inspired from Ročková and George (2014) is drawn. It represents the evolution of the β_ℓ estimates for all $\ell \in \{1, \dots, p\}$ along the grid of ν_0 , and the value of the selection threshold associated with $\nu_0 \in \Delta$. The regularisation plot on Figure 11(A) shows the MAP estimates of the covariate fixed effects vector β obtained for each $\nu_0 \in \Delta$. The blue lines are associated with the three relevant covariates, while the black lines are associated with the null fixed effects of the covariates. Moreover, the red lines correspond to the selection threshold of the covariates. Thus, for each $\nu_0 \in \Delta$, the selected covariates $\hat{S}_{\nu_0}$ are those associated with a $(\hat{\beta}_{\nu_0}^{MAP})_\ell$ located outside the two red curves in the regularisation plot. As expected, the larger ν_0 is, the smaller the support of the associated $\hat{\beta}_{\nu_0}^{MAP}$ is. Indeed, on the one hand, the selection threshold increases with ν_0 , and on the other hand, the larger ν_0 is, the more $(\hat{\beta}_{\nu_0}^{MAP})_\ell$'s are truncated in the spike distribution. This illustrates the interest of going through a grid rather than focusing on a single ν_0 value.

Figure 11(B) represents the value of the eBIC criterion for all ν_0 in Δ . As desired, it is minimal for the values of ν_0 for which exactly the right model is selected. The procedure returns the second value of $\nu_0 \in \Delta$, *i.e.* $\hat{\nu}_0 \approx 0.016$, and $\hat{S}_{\hat{\nu}_0} = \{1, 2, 3\}$. So, in this simulated example, our procedure returns exactly the right model, that is the one with only the first three covariates.

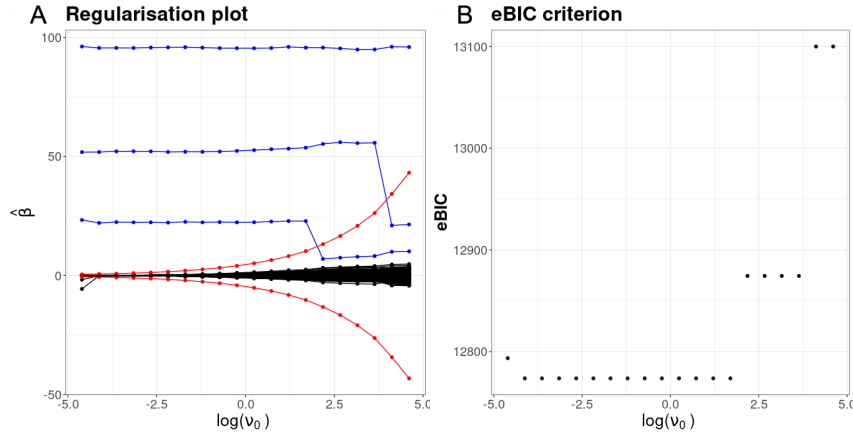


Fig. 11: Example of a regularisation plot (A) with eBIC criterion graph (B) for model selection. On (A), the blue lines are associated with the three true relevant covariates, while the black lines are associated with the null fixed effects of the covariates. The red lines correspond to the selection threshold of the covariates.

Taux de contraction a posteriori dans un modèle non-linéaire à effets mixtes de grande dimension

Publication

Ce chapitre présente le travail de la pré-publication suivante :

Naveau M., Delattre M., Sansonnet L., **Posterior contraction rates in a sparse non-linear mixed-effects model**. *Soumis*, 2024, arXiv:2405.01206.

Abstract

Recent works have shown an interest in investigating the frequentist asymptotic properties of Bayesian procedures for high-dimensional linear models under sparsity constraints. However, there exists a gap in the literature regarding analogous theoretical findings for non-linear models within the high-dimensional setting. The current study provides a novel contribution, focusing specifically on a non-linear mixed-effects model. In this model, the residual variance is assumed to be known, while the covariance matrix of the random effects and the regression vector are unknown and must be estimated. The prior distribution for the sparse regression coefficients consists of a mixture of a point mass at zero and a Laplace distribution, while an Inverse-Wishart prior is employed for the covariance parameter of the random effects. First, the effective dimension of this model is bounded with high posterior probabilities. Subsequently, we derive posterior contraction rates for both the covariance parameter and the prediction term of the response vector. Finally, under additional assumptions, the posterior distribution is shown to contract for recovery of the unknown sparse regression vector at the same rate as observed in the linear case.

Keywords: Posterior contraction rate, Sparse priors, Non-linear mixed-effects models, High-dimensional regression.

Acknowledgement: This work was funded by the Stat4Plant project ANR-20-CE45-0012. The authors thank Professor Ismaël Castillo (Sorbonne Université) for his helpful comments and recommendations about this work.

Contents

6.1	Introduction	99
6.2	Model description	101
6.2.1	Non-linear marginal mixed-effects model	101
6.2.2	Prior specification	102
6.2.3	Assumptions	102
6.3	Posterior contraction results	105
6.3.1	Support size control	105
6.3.2	Posterior contraction rates	105
6.4	Proofs of main theorems	108
6.4.1	Proof of Theorem 1	108
6.4.2	Proof of Theorem 2	110
6.4.3	Proof of Theorem 3	115
6.4.4	Proof of Theorem 4	117
6.5	Appendices	118
6.5.1	Proofs of technical lemmas	118
6.5.2	Useful lemmas	126

6.1 Introduction

Recent statistical literature has shown a keen interest in estimating high-dimensional models under sparsity assumptions, with different approaches proposed over the past few decades in both Bayesian and frequentist frameworks. The developed methodologies are numerous and use a large variety of techniques such as convex and non-convex penalization techniques, shrinkage methods and sparsity-inducing priors. In Bayesian analysis, a category of proposed priors includes those defined as mixtures of two distributions, commonly referred to as spike-and-slab priors. These priors have proven to be useful and relevant in many high-dimensional applications as demonstrated in (George and McCulloch, 1993, 1997; Tadesse and Vannucci, 2021).

The frequentist asymptotic properties of Bayesian sparse linear regression models with various types of mixture priors have been widely investigated, particularly in Narisetty and He (2014), Castillo et al. (2015), and Ročková and George (2018) with the spike-and-slab Gaussian prior, the discrete spike-and-slab prior, and the spike-and-slab lasso prior respectively. Subsequently, these investigations were extended to multivariate linear regression with an unknown residual covariance matrix, as discussed by Ning et al. (2020) with a discrete spike-and-slab prior and Shen and Deshpande (2022) with a multivariate spike-and-slab lasso prior. The classical techniques for determining posterior contraction rates (see *e.g.* (Castillo et al., 2015)) face limitations when the residual covariance matrix is unknown. In such scenarios, the general theory based on the average squared Hellinger distance proves inadequate for obtaining rates in terms of the Euclidean norm for parameters. To overcome this difficulty, an alternative approach has been introduced, leveraging the average Rényi divergence of order $1/2$. As underscored by Ning et al. (2020), this method enables the construction of exponentially powerful tests that are required by the general

theory (Ghosal and Van der Vaart, 2017), facilitating a more effective determination of posterior contraction rates in Bayesian analysis. Another theoretical aspect requires adaptation to the general theory when the residual covariance matrix is unknown. Indeed, classical proofs require lower bounds for prior mass around true parameter values, but when the residual covariance matrix is unknown, this condition can only be fulfilled if the true parameter set is bounded, as discussed in works of Ning et al. (2020) and Jeong and Ghosal (2021a,b).

Recent advancements have expanded the study of estimation and selection properties to more complex models than sparse linear regression models, such as sparse generalized linear models (Jiang, 2007; Jeong and Ghosal, 2021a) or sparse linear regression models with nuisance parameters (Jeong and Ghosal, 2021b). To our knowledge, there are no similar theoretical results for non-linear models in high-dimensional contexts. The absence of theoretical results in this domain may reflect the inherent challenges and complexities associated with extending such analyses to non-linear models. The present paper provides a contribution in this direction, focusing on a specific non-linear model which also contains random effects. Mixed-effects models have been introduced to analyze observations collected repeatedly on several individuals in a population of interest, commonly encountered in fields such as pharmacokinetics or biological growth modeling for example (Pinheiro and Bates, 2000; Lavielle, 2014). These models, which are generally non-linear, may use high-dimensional covariates to describe inter-individual variability. Our paper deals with a generalization of the linear mixed-effects model to a non-linear marginal version where the fixed effects are non-linear functions of the regression parameter, while the random effects are incorporated into the model in a linear manner (see *e.g.* Demidenko (2013)). Such non-linear marginal mixed-effects models are easier to handle than more general non-linear mixed-effects models because the mean and the covariance matrix of the response variable are explicit. However, despite their practical appeal, there has been a lack of theoretical exploration concerning non-linear marginal mixed models in high-dimensional context. In this paper, posterior contraction rates are obtained for both the covariance matrix and the prediction term in a high-dimensional setting by using a mixture of a point mass at zero and a Laplace distribution prior on the regression coefficients and an inverse Wishart prior on the covariance matrix. These results are extended to the regression coefficients themselves under additional assumptions.

This paper is organized as follows. Section 6.2 describes the non-linear marginal mixed model to introduce the notation, defines the prior distributions, along with the necessary assumptions. Section 6.3 provides the main results on the posterior contraction. Finally, the proofs of the theorems are given in Section 6.4. Proofs of technical lemmas are postponed in Appendix.

Notation This paragraph describes the notations used in this paper for a generic matrix A and a generic vector $\theta \in \mathbb{R}^k$. We note $S_\theta = \{j | \theta_j \neq 0\}$ the support of θ and $s_\theta = |S_\theta|$ its cardinal. The euclidean norm, the ℓ_1 -norm and the infinity norm are respectively noted $\|\theta\|_2 = \sqrt{\sum_{i=1}^k \theta_i^2}$, $\|\theta\|_1 = \sum_{i=1}^k |\theta_i|$, and $\|\theta\|_\infty = \max_i |\theta_i|$. The transpose of A is denoted by $A^\top$. For a square matrix A , we note $\rho_{\min}(A)$ and $\rho_{\max}(A)$ the minimum and maximum eigenvalues of A , respectively. The spectral norm of a matrix A is denoted $\|A\|_{sp} = \rho_{\max}^{1/2}(A^\top A)$, and the Frobenius norm is noted $\|A\|_F = \text{Tr}(A^\top A)^{1/2} = (\sum_{i,j} x_{ij}^2)^{1/2}$. The matrix norm $\|A\|_*$ is defined as $\|A\|_* = \max_j \|A_{\cdot j}\|_2$ for $A_{\cdot j}$ the j -th column of A . The identity matrix of size m is denoted I_m .

For sequences a_n and b_n , the notation $a_n \lesssim b_n$ means that for n large enough a_n is bounded

above by a constant multiple of b_n , i.e. $a_n \leq Cb_n$ for n large enough, where $C > 0$ is independent of n . We denote $a_n = o(b_n)$ if $\frac{a_n}{b_n} \xrightarrow{n \rightarrow \infty} 0$.

6.2 Model description

6.2.1 Non-linear marginal mixed-effects model

Mixed-effects models are sophisticated multivariate statistical models employed to analyze repeated observations, usually collected over time, on multiple statistical subjects, incorporating both fixed and random effects into the model for accurate description (Lavielle, 2014; Demidenko, 2013). We consider the following mixed-effects model:

$$Y_i = f_i(X_i\beta) + Z_i\xi_i + \varepsilon_i, \quad \varepsilon_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_{n_i}(0, \sigma^2 I_{n_i}), \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \quad i = 1, \dots, n. \quad (6.1)$$

In the above equation, n is the number of individuals, $Y_i \in \mathbb{R}^{n_i}$ and n_i are the vector and the number of observations for subject i respectively, $\xi_i \in \mathbb{R}^q$ is a vector composed of q random effects, $\varepsilon_i \in \mathbb{R}^{n_i}$ is the error term. $X_i \in \mathbb{R}^{q \times p}$ and $Z_i \in \mathbb{R}^{n_i \times q}$ are design matrices composed of explanatory variables that relate the observations to fixed-effects $\beta \in \mathbb{R}^p$ and to the random effects respectively. In the development of a mixed-effects model, a significant focus lies on the covariate selection process in X_i , which amounts to the identification of non-null components in β , as it allows establishing connections between inter-individual variability and measured individual characteristics. In the above model, the function f_i defines the relationship between observations of subject i and the explanatory variables in X_i which are pivotal in this stage of model construction and are commonly denoted as covariates in subsequent discussions. To draw parallels with the classical literature on mixed-effects models for longitudinally repeated data, we define

$$f_i(x) = (f(x, t_{i1}), f(x, t_{i2}), \dots, f(x, t_{in_i}))^\top,$$

where t_{ij} represents the j -th observation time for individual i , and $f : \mathbb{R}^q \times \mathbb{R} \rightarrow \mathbb{R}$ denotes a regression function chosen to effectively capture the longitudinally observed phenomenon. While opting for $f_i(X_i\beta) = X_i\beta$ yields the standard linear mixed-effects model, it's common in many applications to select a non-linear function f . Moreover, when f is non-linear, Model (6.1) is alternatively referred to as the non-linear marginal mixed-effects model (as discussed in (Demidenko, 2013)). The term "marginal mixed-effects model" is derived from the fact that, unlike numerous other non-linear mixed-effects models, both the expectation and variance of the observations possess an explicit expression. The distribution of Y_i defined through (6.1) is thus fully characterized:

$$Y_i \sim \mathcal{N}(f_i(X_i\beta), \Delta_{\Gamma,i}), \text{ where } \Delta_{\Gamma,i} = Z_i\Gamma Z_i^\top + \sigma^2 I_{n_i}. \quad (6.2)$$

In the following, the residual variance σ^2 and the number q of true random effects are assumed to be known. The aim is to estimate $(\beta, \Gamma) \in \mathcal{B} \times \mathcal{H}$ in an high-dimensional setting where $p \gg n$ and obtain posterior contraction results. We establish below appropriate priors to achieve these goals.

6.2.2 Prior specification

Drawing from classical literature in high-dimensional Bayesian analysis, this study adopts an approach employing priors that induce sparsity in β coefficients. For that purpose, we jointly consider a prior π_p on the number s of non-zero coefficients in β and a Laplace prior on the non-zero coefficients in β while setting the other components in β to zero:

$$(S, \beta) \mapsto \frac{\pi_p(s)}{\binom{p}{s}} g_S(\beta_S) \delta_0(\beta_{S^c}), \quad (6.3)$$

where S is a subset of s elements in $\{1, \dots, p\}$ that represents the support of β , *i.e.* the positions of its non-zero elements, S^c is the complementary subset of zero coefficients in β , $\beta_S = \{\beta_\ell | \ell \in S\}$ and $\beta_{S^c} = \{\beta_\ell | \ell \notin S\}$ are the non-zero and the zero coefficients in β respectively, δ_0 is the Dirac measure at zero on $\mathbb{R}^{p-s}$ and

$$g_S(\beta_S) = \prod_{\ell \in S} \frac{\lambda}{2} \exp(-\lambda |\beta_\ell|). \quad (6.4)$$

Concerning the random effects covariance matrix Γ , a conjugate inverse-Wishart prior is used:

$$\pi(\Gamma) \propto |\Gamma|^{-\frac{d+q+1}{2}} \exp\left(-\frac{1}{2} \text{Tr}(\Sigma \Gamma^{-1})\right),$$

where Σ is a positive definite matrix of size $q \times q$, and $d > q - 1$ the degree of freedom. This prior is chosen for a practical matter. Note that, as discussed in [Ning et al. \(2020\)](#), the inverse-Wishart prior may induce sub-optimal posterior contraction rate due to its weaker tail property when q increases to infinity. However, here q is assumed to be fixed so the rate should not be impacted by this property.

6.2.3 Assumptions

The frequentist assumption that the data, n independent observations $Y^{(n)} = (Y_i)_{1 \leq i \leq n} \in \mathbb{R}^N$, where $N = \sum_{i=1}^n n_i$, has been generated from model (6.1) for a given sparse regression parameter β_0 and a given random effects covariance matrix Γ_0 is made. The expectation under these true parameters is denoted $\mathbb{E}_0$. The support of the true parameter β_0 is denoted S_0 and its cardinal s_0 . The maximum number of observations per individual is defined as $J_n = \max_{1 \leq i \leq n} n_i$.

6.2.3.1 Assumptions on the non-linear model structure

Assumptions have to be made on the regression function f to obtain posterior contraction. A first natural condition is the Lipschitz assumption, allowing for easy control of the regression function from its inputs.

Assumption A1. $f : \mathbb{R}^q \times \mathbb{R} \rightarrow \mathbb{R}$ is K -Lipschitz with respect to its first component:

$$\forall x, y \in \mathbb{R}^q, \forall t \in \mathbb{R}, |f(x, t) - f(y, t)| \leq K \|x - y\|_2$$

Remark 9. Under assumption A1, notice we have that $f_i : \mathbb{R}^q \rightarrow \mathbb{R}^{n_i}$ is $\sqrt{K^2 n_i}$ -Lipschitz for $\|\cdot\|_2$.

As outlined in the introduction, to satisfy the condition of prior mass around the true parameters, they should be confined within a specific subset of the parameter space characterized by bounded norms.

Assumption A2. *The true value β_0 belongs to $\mathcal{B}_0 := \{\beta \in \mathbb{R}^p : \|\beta\|_\infty \lesssim \lambda^{-1} \log(p)\}$, where λ is the regularization parameter of the Laplace distribution defined in equation (6.4).*

Assumption A3. *The true covariance matrix of the random effects Γ_0 belongs to $\mathcal{H}_0 := \{\Gamma : 1 \lesssim \rho_{\min}(\Gamma) \leq \rho_{\max}(\Gamma) \lesssim 1\}$, and we denote $\underline{\rho}_\Gamma > 0$ and $\overline{\rho}_\Gamma > 0$ the bounds such that: $\underline{\rho}_\Gamma \leq \rho_{\min}(\Gamma_0) \leq \rho_{\max}(\Gamma_0) \leq \overline{\rho}_\Gamma$.*

Assumption A2 allows that the prior assigns sufficient mass on a Kullback-Leibler neighborhood of β_0 . In the same way, assumption A3 enables to put sufficient mass around the true parameter Γ_0 in terms of Frobenius norm. Similar conditions can be found in the work of Ning et al. (2020), Jeong and Ghosal (2021b), and Song and Liang (2023). This is in contrast to Castillo et al. (2015)'s work where they obtain a result uniformly over the entire parameter space because they have explicit expressions to satisfy this condition directly in their case of univariate regression with known variance. Also, it is assumed that β_0 is not the zero vector, and that p does not converge faster than exponential of n (see assumption A4).

Assumption A4. *The true support size satisfies $s_0 > 0$ and the following high-dimensional setting is consider $s_0 \log(p) = o(n)$.*

6.2.3.2 Assumptions on the prior distributions

The importance of the prior π_p lies in its essential role in representing the sparsity of the parameter. The crucial aspect of the prior π_p on model dimension is to appropriately reduce the influence of larger models while maintaining sufficient weight for the true one. It is revealed that an exponential decrease effectively fulfills this requirement (Castillo et al., 2015). The following assumption is made on π_p .

Assumption A5 (Prior dimension). *For some constants $A_1, A_2, A_3, A_4 > 0$,*

$$A_1 p^{-A_3} \pi_p(s-1) \leq \pi_p(s) \leq A_2 p^{-A_4} \pi_p(s-1) \quad , \text{ for } s = 1, \dots, p$$

Examples of priors satisfying this assumption A5 are given in Castillo and van der Vaart (2012) and Castillo et al. (2015). In fact, this type of prior is more generic than the discrete spike-and-slab prior. Indeed, if each coordinate β_ℓ is modeled as a mixture $(1-r)\delta_0 + rG$, where G follows the Laplace distribution, it can be realized as a prior of the form (6.3) by selecting π_p as a binomial distribution with parameters p and r . Since r controls the level of sparsity, which is unknown, a classical Bayesian strategy is to put a hyper-prior $Beta(1, p^u)$ with $u > 1$. Then, the overall prior satisfies the exponential decay rate A5. Furthermore, the regularization parameter of the Laplace prior λ must be bounded from below and above, as specified in the following assumption. Indeed, an excessively large value of λ would shrink non-zero coordinates of β towards 0, which is undesirable. Conversely, a too small value of λ may introduce false signals in the support, thereby slowing down the posterior contraction rate.

Assumption A6. The regularization parameter λ of the Laplace prior on the non-zero coordinates of β satisfies:

$$\frac{\|X\|_* K_n}{L_1 p^{L_2}} \leq \lambda \leq \frac{L_3 \|X\|_* K_n}{\sqrt{n}},$$

for some constants $L_1, L_2, L_3 > 0$, where $K_n = \sqrt{K^2 J_n}$ and $X = \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix} \in \mathbb{R}^{nq \times p}$.

Similar condition can be found in Jeong and Ghosal (2021b) for example. Note that, since the size of signal in β_0 is restricted (assumption A2), the Laplace density is not required to achieve the posterior contraction rates in Theorem 4. Other slab densities with similar tail properties, such as the Gaussian slab, can also be used with appropriate adjustments for the true signal size (see Jeong and Ghosal (2021a,b) for more details).

6.2.3.3 Assumptions about the experimental design

For $\Gamma_1, \Gamma_2 \in \mathcal{H}$, we define the pseudo distance $d_n^2(\Gamma_1, \Gamma_2) = \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2$, where $\Delta_{\Gamma, i}$ is defined in equation (6.2). The following three assumptions on the model A7, A8, A9 enable to control the maximum Frobenius norm of the difference between covariance matrices from the average Frobenius norm:

$$\max_i \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 \lesssim \|\Gamma_1 - \Gamma_2\|_F^2 \lesssim d_n^2(\Gamma_1, \Gamma_2) = \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2.$$

This point is demonstrated in Appendix 6.5.2, Lemma B3.

Assumption A7. The quantity $\frac{1}{n} \sum_{i=1}^n \mathbb{1}_{n_i \geq q}$ is bounded.

Assumption A7 means that the number of individuals i such as the number of observations n_i is greater than the number of random effects q is of the order of n , that is n_i is probably greater than q , which seems statistically reasonable to be able to estimate Γ and β . This is a necessary assumption for the identifiability of the model.

Assumption A8. For each $1 \leq i \leq n$ such that $n_i \geq q$, Z_i is of full rank, i.e. $\min_i \left\{ \rho_{\min}^{1/2}(Z_i^\top Z_i) : n_i \geq q \right\} \gtrsim 1$. We denote by $\underline{\rho_Z}$ the bound:

$$\min_i \left\{ \rho_{\min}^{1/2}(Z_i^\top Z_i) : n_i \geq q \right\} \geq \underline{\rho_Z}.$$

Assumption A9. For each $1 \leq i \leq n$, the maximum of $\|Z_i\|_{sp}$ is bounded, i.e. $\max_i \|Z_i\|_{sp} \lesssim 1$. We denote by $\overline{\rho_Z}$ the bound:

$$\max_i \|Z_i\|_{sp} \leq \overline{\rho_Z}.$$

Similar assumptions can be found in Theorem 10 of Jeong and Ghosal (2021b) for the linear mixed-effects model.

6.3 Posterior contraction results

In this section, we provide results on posterior contraction rates in sparse non-linear marginal mixed-effects model under suitable assumptions presented in Section 6.2.3. To achieve this, we first analyze a dimensionality property of the support of β . Then, we determine how quickly the posterior contracts based on the average Rényi divergence. Finally, we use this information about Rényi contraction to establish the rates for the parameters relative to more practical metrics.

6.3.1 Support size control

First, it is essential to examine the support size of β in order to then focus on models of relatively small sizes. The following theorem shows that the posterior distribution tends to concentrate on models of relatively small sizes, not much larger than the true one.

Theorem 1 (Effective dimension). *In model (6.1), with prior specifications outlined in Section 6.2.2, and assuming the validity of previous assumptions A1-A9, there exists a constant $C_1 > 0$ such that the following convergence holds:*

$$\sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : |S_\beta| > C_1 s_0 \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

Proof of this theorem is provided in Section 6.4.1. The derivation of the posterior contraction rate heavily relies on a technical lemma which provides a lower bound for the denominator of the posterior distribution with probability tending to 1, see Lemma 1 in Section 6.4.1. More precisely, this lemma is employed in deriving our main results on effective dimension and posterior contraction rates, as outlined in Theorems 1 and 2.

6.3.2 Posterior contraction rates

As discussed in the introduction, the classical approach for determining posterior contraction rates encounters limitations when dealing with the unknown nature of the random effects covariance matrix. Indeed, this approach based on the average squared Hellinger distance faces inadequacies in obtaining rates in terms of the Euclidean norm for the parameters in this context. Specifically, the issue arises from the fact that establishing proximity using the average squared Hellinger distance between multivariate normal densities with individual-specific mean and an unknown covariance does not guarantee average proximity in terms of the Euclidean distance for the mean parameters in these densities. To overcome this challenge, the proposed solution is a direct utilization of the average Rényi divergence of order 1/2 (see Definition 1). This approach is highlighted for its high manageability in the context of multivariate normal distributions and its ability to ensure closeness in terms of the desired Euclidean distance. Examples of the application of this theory can be found in the works of Ning et al. (2020) and Jeong and Ghosal (2021b), further supporting the efficacy of the average Rényi divergence in overcoming the limitations associated with the unknown covariance matrix for random effects.

Definition 1. *For two n -variables densities $f = \prod_{i=1}^n f_i$ and $g = \prod_{i=1}^n g_i$ of independent variables, the average Rényi divergence (of order 1/2) is defined by:*

$$R_n(f, g) = -\frac{1}{n} \sum_{i=1}^n \log \left(\int \sqrt{f_i g_i} \right)$$

Based on the result of Theorem 1, the following theorem establishes the rate of contraction of the posterior distribution towards the truth with respect to the average Rényi divergence.

Theorem 2 (Contraction rate, Rényi). *In model (6.1), with prior specifications outlined in Section 6.2.2, we denote by $p_{\beta, \Gamma} = \prod_{i=1}^n p_{\beta, \Gamma, i}$ the joint density, with $p_{\beta, \Gamma, i}$ representing the density of the i th observation vector y_i , and p_0 representing the true joint density. Assuming the previous assumptions A1-A9 hold, as well as $\log(nJ_n) \lesssim \log(p)$, then there exists a constant $C_2 > 0$ such that:*

$$\sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left((\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : R_n(p_{\beta, \Gamma}, p_0) > C_2 \frac{s_0 \log(p)}{n} \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 0.$$

The proof can be found in Section 6.4.2. This proof is based on the general theory of posterior contraction rate of Ghosal et al. (2000); Ghosal and Van Der Vaart (2007); Ghosal and Van der Vaart (2017), which relies on the construction and existence of exponentially powerful tests (see also Castillo (2024) for more details).

While Theorem 2 provides a fundamental result on posterior contraction, it does not offer precise interpretations for the parameters β and Γ . The following theorem relies on the form of the average Rényi divergence to obtain more concrete contraction rates. Specifically, it demonstrates that the posterior distribution of the prediction term and Γ contracts towards their true respective values at certain rates, relative to metrics more easily understandable than the average Rényi divergence.

Theorem 3 (Recovery). *In model (6.1), with prior specifications outlined in Section 6.2.2, and assuming Assumptions A1-A9, as well as $\log(nJ_n) \lesssim \log(p)$, then there exist constants $C_3, C_4, C_5 > 0$ such that:*

$$\begin{aligned} \sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\Gamma : d_n(\Gamma, \Gamma_0) > C_3 \sqrt{\frac{s_0 \log(p)}{n}} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\Gamma : \|\Gamma - \Gamma_0\|_F > C_4 \sqrt{\frac{s_0 \log(p)}{n}} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \sqrt{\frac{1}{n} \sum_{i=1}^n \|f_i(X_i \beta) - f_i(X_i \beta_0)\|_2^2} > C_5 \sqrt{\frac{s_0 \log(p)}{n}} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0. \end{aligned}$$

The proof can be found in Section 6.4.3. By comparing our theorem to Castillo et al. (2015)'s results in Bayesian, or Bühlmann and Van De Geer (2011)'s results in frequentist framework, in simple linear regression, it can be observed that the same rates are achieved for the prediction term. For the covariance term, the rate obtained in Theorem 3 coincides with that obtained for linear regression with nuisance parameters by Jeong and Ghosal (2021b).

The last theorem gives precise interpretations of the posterior contraction result for the parameter β . The posterior contraction rates with respect to more concrete metrics are derived based on additional conditions, summarized by Assumptions A10 and A11.

Assumption A10. *For all $i \in \{1, \dots, n\}$, $\delta > 0$ and $t \in \mathbb{R}$,*

$$\left\{ \beta \in \mathbb{R}^p : |f(X_i \beta, t) - f(X_i \beta_0, t)| \leq \delta \right\} \subset \left\{ \beta \in \mathbb{R}^p : |f(X_i \beta, t) - f(X_i \beta_0, t)| \gtrsim \|X_i(\beta - \beta_0)\|_2 \right\}.$$

This assumption, as well as Assumption A7, is a necessary condition for the identifiability of the model and allows to derive posterior contraction rate for β from the third assertion of Theorem 3. In particular, this assumption implies that for all $i \in \{1, \dots, n\}$ and $t \in \mathbb{R}$, the true $X_i\beta_0$ does not lie in a neighborhood of critical points of $f(\cdot, t)$, which seems a reasonable assumption. To ensure the identifiability of the parameter β , a kind of assumption of "local invertibility" for the Gram matrix $X^\top X$ is also required. For this purpose, we define the following compatibility numbers drawing from the literature (Castillo et al., 2015).

Definition 2. For all $s > 0$, the smallest scaled singular value of dimension s is defined as:

$$\phi_2(s) = \inf_{\beta: 1 \leq s_\beta \leq s} \frac{\|X\beta\|_2}{\|X\|_* \|\beta\|_2}.$$

Definition 3. For all $s > 0$, the uniform compatibility number in dimension s is defined as:

$$\phi_1(s) = \inf_{\beta: 1 \leq s_\beta \leq s} \frac{\|X\beta\|_2 \sqrt{s_\beta}}{\|X\|_* \|\beta\|_1}.$$

Assumption A11. For each $1 \leq i \leq n$, the maximum of $\|X_i\|_*$ is bounded, i.e. $\max_i \|X_i\|_* \lesssim 1$, and $\beta_0 \in \bar{\mathcal{B}}_0 := \left\{ \beta \in \mathcal{B}_0 : \frac{s_0^2 \log(p)}{\|X\|_*^2 \phi_1^2((C_1 + 1)s_0)} = o(1) \right\}$.

Typically, the first assertion in this assumption is commonly satisfied in practical scenarios. The second assertion concerns the true parameter β_0 and will be used for recovery of $X\beta$ in ℓ_2 -norm and β with respect to the ℓ_1 -norm and the ℓ_2 -norm. Since $\phi_2(s) \leq \phi_1(s)$ for all $s > 0$ by the Cauchy-Schwarz inequality, ϕ_1 can be removed if the smallest scaled singular value ϕ_2 is bounded away from zero. Note that under specific conditions on the design matrix, the compatibility numbers can be bounded away from zero (see Example 7 of Castillo et al. (2015) for further discussion). Thus, since by the first assertion of Assumption A11, $\|X\|_*^2 \lesssim n$, the second assertion implies that $s_0^2 \log(p) = o(n)$. This condition is similar to that obtained in Jeong and Ghosal (2021a).

Theorem 4 (Posterior contraction rate for β). In model (6.1), with prior specifications outlined in Section 6.2.2, and assuming Assumptions A1-A11, as well as $\log(nJ_n) \lesssim \log(p)$, then there exist constants $C_6, C_7, C_8 > 0$ such that:

$$\begin{aligned} \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|X(\beta - \beta_0)\|_2 > C_6 \sqrt{s_0 \log(p)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|\beta - \beta_0\|_2 > C_7 \frac{\sqrt{s_0 \log(p)}}{\|X\|_* \phi_2((C_1 + 1)s_0)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0, \\ \sup_{\beta_0 \in \bar{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0} \mathbb{E}_0 \left[\Pi \left(\beta : \|\beta - \beta_0\|_1 > C_8 \frac{s_0 \sqrt{\log(p)}}{\|X\|_* \phi_1((C_1 + 1)s_0)} \middle| Y^{(n)} \right) \right] &\xrightarrow{n \rightarrow \infty} 0. \end{aligned}$$

The proof can be found in Section 6.4.4. These rates coincide with those obtained by Jeong and Ghosal (2021a) or Jeong and Ghosal (2021b), respectively for generalized linear model and linear regression with nuisance parameters. Since the compatibility numbers can be bounded away from zero under some conditions (see above), they can be removed from the rates.

6.4 Proofs of main theorems

In this section, proofs of the main theorems are provided. First, additional notations used for the proofs are introduced. Let $\Lambda_n(\beta, \Gamma) = \prod_{i=1}^n p_{\beta, \Gamma, i} / p_{0, i}$ be the likelihood ratio of $p_{\beta, \Gamma} = \prod_{i=1}^n p_{\beta, \Gamma, i}$, where $p_{\beta, \Gamma, i}$ is the density of the i -th observation vector y_i , and $p_0 = \prod_{i=1}^n p_{0, i} = \prod_{i=1}^n p_{\beta_0, \Gamma_0, i}$ the density with the true parameters β_0 and Γ_0 . For two densities f and g , let $K(f, g) = \int f(x) \log(f(x)/g(x)) dx$ the Kullback-Leibler divergence, and $V(f, g) = \int f(x) |\log(f(x)/g(x)) - K(f, g)|^2 dx$ the Kullback-Leibler variation.

6.4.1 Proof of Theorem 1

The proof of Theorem 1 is based on the following technical lemma which provides a lower bound for the denominator of the posterior distribution, with the probability tending to 1.

Lemma 1. *Suppose that Assumptions A1-A9 are satisfied. Then, there exists a positive constant M such that*

$$\mathbb{P}_0 \left(\int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \geq \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))} \right) \longrightarrow 1, \quad (6.5)$$

when n tends to infinity.

This lemma is demonstrated in Appendix 6.5.1.1.

Proof of Theorem 1. Let $B = \{(\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : |S_\beta| > \tilde{s}\}$, with an integer $\tilde{s} \geq s_0$. First, by Bayes formula:

$$\Pi(B|Y^{(n)}) = \frac{\int_B \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)}{\int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)}. \quad (6.6)$$

Let us prove that $\mathbb{E}_0 [\Pi(B|Y^{(n)})]$ tends to 0 as n tends to infinity uniformly for $\beta_0 \in \mathcal{B}_0$ and $\Gamma_0 \in \mathcal{H}_0$, and choose a suitable $\tilde{s}$. Let $\mathcal{A}_n$ be the event that appears in Equation (6.5). We can write

$$\mathbb{E}_0 [\Pi(B|Y^{(n)})] = \mathbb{E}_0 [\Pi(B|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n}] + \mathbb{E}_0 [\Pi(B|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n^c}]. \quad (6.7)$$

where the second term tends to 0 by using Lemma 1.

Concerning the first term, by definition of $\mathcal{A}_n$, we have that

$$\begin{aligned} \mathbb{E}_0 [\Pi(B|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n}] &= \mathbb{E}_0 \left[\frac{\int_B \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)}{\int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)} \mathbf{1}_{\mathcal{A}_n} \right] \\ &\leq \mathbb{E}_0 \left[\int_B \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \pi_p(s_0)^{-1} e^{M(s_0 \log(p) + \log(n))} \mathbf{1}_{\mathcal{A}_n} \right] \\ &\leq \pi_p(s_0)^{-1} \exp \{M(s_0 \log(p) + \log(n))\} \mathbb{E}_0 \left[\int_B \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \mathbf{1}_{\mathcal{A}_n} \right] \end{aligned}$$

Now, we get that

$$\begin{aligned}\mathbb{E}_0 \left[\int_B \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \mathbf{1}_{\mathcal{A}_n} \right] &\leq \mathbb{E}_0 \left[\int_B \frac{p_{\beta, \Gamma}(y)}{p_0(y)} d\Pi(\beta, \Gamma) \right] \\ &= \int \int_B p_{\beta, \Gamma}(y) d\Pi(\beta, \Gamma) dy \\ &= \Pi(B)\end{aligned}$$

using Fubini-Tonelli theorem and since $p_{\beta, \Gamma}$ is a density. Thus,

$$\mathbb{E}_0 \left[\Pi \left(B|Y^{(n)} \right) \mathbf{1}_{\mathcal{A}_n} \right] \leq \pi_p(s_0)^{-1} \exp \{ M(s_0 \log(p) + \log(n)) \} \Pi(B),$$

and by Assumption A5,

$$\begin{aligned}\Pi(B) &= \Pi(|S_\beta| > \tilde{s}) = \sum_{s=\tilde{s}+1}^p \pi_p(s) \\ &\leq \pi_p(s_0) \sum_{s=\tilde{s}+1}^p (A_2 p^{-A_4})^{s-s_0} \\ &= \pi_p(s_0) (A_2 p^{-A_4})^{\tilde{s}+1-s_0} \sum_{k=0}^{p-\tilde{s}-1} (A_2 p^{-A_4})^k \\ &\leq \pi_p(s_0) (A_2 p^{-A_4})^{\tilde{s}+1-s_0} \frac{1}{1 - A_2 p^{-A_4}},\end{aligned}$$

for p large enough to ensure that $A_2 p^{-A_4} < 1$. Thus finally we have

$$\begin{aligned}\mathbb{E}_0 \left[\Pi \left(B|Y^{(n)} \right) \mathbf{1}_{\mathcal{A}_n} \right] &\leq \pi_p(s_0)^{-1} \exp \{ M(s_0 \log(p) + \log(n)) \} \Pi(B) \\ &\leq \exp \{ M(s_0 \log(p) + \log(n)) + (\tilde{s} + 1 - s_0) \log(A_2 p^{-A_4}) \} \frac{1}{1 - A_2 p^{-A_4}} \\ &= \exp \left\{ \log(p) \left(M s_0 + M \frac{\log(n)}{\log(p)} - A_4(\tilde{s} + 1 - s_0) \right) + (\tilde{s} + 1 - s_0) \log(A_2) \right\} \frac{1}{1 - A_2 p^{-A_4}}\end{aligned}$$

where $\log(n)/\log(p) \leq 1$ as $p > n$. Thus, as $(1 - A_2 p^{-A_4})^{-1}$ tends to 1 when $n \rightarrow \infty$, we choose $\tilde{s}$ as the largest integer that is smaller than $C_1 s_0$ (such as $\tilde{s} + 1 > C_1 s_0$), for some constant C_1 large enough to have $M s_0 + M - A_4(C_1 s_0 - s_0) < 0$, and then we have that

$$\begin{aligned}\mathbb{E}_0 \left[\Pi \left(B|Y^{(n)} \right) \mathbf{1}_{\mathcal{A}_n} \right] &\leq \exp \{ \log(p) (M s_0 + M - A_4(C_1 s_0 - s_0)) + (\tilde{s} + 1 - s_0) \log(A_2) \} \frac{1}{1 - A_2 p^{-A_4}} \xrightarrow{n \rightarrow \infty} 0.\end{aligned}$$

Finally, by Equation (6.7), we conclude that $\mathbb{E}_0 [\Pi(B|Y^{(n)})] \rightarrow 0$, for this well-chosen $\tilde{s}$. Thus, we have also that $\mathbb{E}_0 \left[\Pi \left(\beta : |S_\beta| > C_1 s_0 \middle| Y^{(n)} \right) \right] \rightarrow 0$, which concludes the proof of the theorem. $\square$

6.4.2 Proof of Theorem 2

Proof of Theorem 2. Let $\mathcal{B}_n = \{\beta \in \mathcal{B} : s_\beta \leq C_1 s_0\}$, $R_n^*(\beta, \Gamma) = R_n(p_{\beta, \Gamma}, p_0)$ and $\epsilon_n = \sqrt{\frac{s_0 \log(p)}{n}}$.

$$\begin{aligned} & \mathbb{E}_0 \left[\Pi \left((\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : R_n^*(\beta, \Gamma) > C_2 \epsilon_n^2 \middle| Y^{(n)} \right) \right] \\ & \leq \mathbb{E}_0 \left[\Pi \left((\beta, \Gamma) \in \mathcal{B}_n \times \mathcal{H} : R_n^*(\beta, \Gamma) > C_2 \epsilon_n^2 \middle| Y^{(n)} \right) \right] + \mathbb{E}_0 \left[\Pi \left(\mathcal{B}_n^c \middle| Y^{(n)} \right) \right] \end{aligned}$$

where the second term tends to 0 when n goes to infinity by Theorem 1.

Therefore, given $D = \{(\beta, \Gamma) \in \mathcal{B}_n \times \mathcal{H} : R_n^*(\beta, \Gamma) > C_2 \epsilon_n^2\}$, proving Theorem 2 consists in showing that $\mathbb{E}_0 [\Pi(D|Y^{(n)})]$ goes to 0 as n tends to infinity uniformly for $\beta_0 \in \mathcal{B}_0$ and $\Gamma_0 \in \mathcal{H}_0$.

This proof is based on the construction and existence of exponentially powerful tests to show contraction rates of posterior distributions (see Ghosal et al. (2000); Ghosal and Van der Vaart (2017) for more details). More precisely, we want to construct a test φ_n such that on an appropriate sieve $\mathcal{B}_n^* \times \mathcal{H}_n \subset \mathcal{B}_n \times \mathcal{H}$ we have, for some constants $M_1, M_2 > 0$:

$$\mathbb{E}_0[\varphi_n] \lesssim e^{-M_1 n \epsilon_n^2}, \quad \sup_{(\beta, \Gamma) \in \mathcal{B}_n^* \times \mathcal{H}_n : R_n^*(\beta, \Gamma) > C_2 \epsilon_n^2} \mathbb{E}_{(\beta, \Gamma)}[1 - \varphi_n] \leq e^{-M_2 n \epsilon_n^2} \quad (6.8)$$

where the sieve $\mathcal{B}_n^* \times \mathcal{H}_n$ shall satisfy that the prior mass of $\mathcal{B}_n \setminus \mathcal{B}_n^*$ and $\mathcal{H} \setminus \mathcal{H}_n$ decreases rapidly enough to balance the denominator of the posterior. Indeed, assuming that we have constructed such a test, then, for $\mathcal{A}_n$ the event that appears in Equation (6.5):

$$\begin{aligned} \mathbb{E}_0 [\Pi(D|Y^{(n)})] &= \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n}] + \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n^c}] \\ &= \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n} (1 - \varphi_n) + \Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n} \varphi_n] + \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n^c}] \\ &\leq \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n} (1 - \varphi_n)] + \mathbb{E}_0 [\varphi_n] + \mathbb{P}_0(\mathcal{A}_n^c) \end{aligned}$$

where by construction of φ_n , $\mathbb{E}_0[\varphi_n] \xrightarrow{n \rightarrow \infty} 0$, and $\mathbb{P}_0(\mathcal{A}_n^c) \xrightarrow{n \rightarrow \infty} 0$ by Lemma 1.

Now for the first term, by the Bayes formula (6.6), we have that

$$\begin{aligned} \mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbf{1}_{\mathcal{A}_n} (1 - \varphi_n)] &= \mathbb{E}_0 \left[\frac{\int_D \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)}{\int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)} \mathbf{1}_{\mathcal{A}_n} (1 - \varphi_n) \right] \\ &\leq \mathbb{E}_0 \left[\int_D \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \pi_p(s_0)^{-1} e^{M(s_0 \log(p) + \log(n))} (1 - \varphi_n) \right] \end{aligned}$$

But, grant Assumption A5, we have that: $\pi_p(s_0)^{-1} \leq A_1^{-1} p^{A_3} \pi_p(s_0 - 1)^{-1}$ and by iteration

$$-\log(\pi_p(s_0)) \lesssim s_0 \log(p) - \log(\pi_p(0)) \lesssim s_0 \log(p)$$

since $1 = \sum_{s=1}^p \pi_p(s) \leq \sum_{s=1}^p (A_2 p^{-A_4})^s \pi_p(0) \lesssim \pi_p(0)$ by assumption A5. Thus, for a constant C large enough, $\pi_p(s_0)^{-1} e^{M(s_0 \log(p) + \log(n))} \leq e^{C s_0 \log(p)} = e^{C n \epsilon_n^2}$, since $\log(n) \lesssim s_0 \log(p)$. So, by using the Fubini-Tonelli theorem,

$$\begin{aligned}
& \mathbb{E}_0 \left[\Pi \left(D|Y^{(n)} \right) \mathbb{1}_{\mathcal{A}_n} (1 - \varphi_n) \right] \\
& \leq \int_D \mathbb{E}_{(\beta, \Gamma)} [1 - \varphi_n] d\Pi(\beta, \Gamma) \times e^{Cn\epsilon_n^2} \\
& \leq \left(\int_{D \cap (\mathcal{B}_n^* \times \mathcal{H}_n)} \mathbb{E}_{(\beta, \Gamma)} [1 - \varphi_n] d\Pi(\beta, \Gamma) + \Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) + \Pi(\mathcal{H} \setminus \mathcal{H}_n) \right) \times e^{Cn\epsilon_n^2} \\
& \leq \left(\sup_{(\beta, \Gamma) \in D \cap (\mathcal{B}_n^* \times \mathcal{H}_n)} \{ \mathbb{E}_{(\beta, \Gamma)} [1 - \varphi_n] \} + \Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) + \Pi(\mathcal{H} \setminus \mathcal{H}_n) \right) \times e^{Cn\epsilon_n^2} \\
& \leq \left(e^{-M_2 n \epsilon_n^2} + \Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) + \Pi(\mathcal{H} \setminus \mathcal{H}_n) \right) \times e^{Cn\epsilon_n^2}
\end{aligned}$$

by construction of φ_n , equation (6.8). Then for M_2 large enough and by the condition on the prior mass of $\mathcal{B}_n \setminus \mathcal{B}_n^*$ and $\mathcal{H} \setminus \mathcal{H}_n$, we have that $\mathbb{E}_0 [\Pi(D|Y^{(n)}) \mathbb{1}_{\mathcal{A}_n} (1 - \varphi_n)] \xrightarrow{n \rightarrow \infty} 0$, and finally $\mathbb{E}_0 [\Pi(D|Y^{(n)})] \xrightarrow{n \rightarrow \infty} 0$, what was wanted to be demonstrated.

Thus, to complete the proof, we need to demonstrate the existence of such a test φ_n satisfying (6.8) on an appropriate sieve $\mathcal{B}_n^* \times \mathcal{H}_n$ such that the prior mass of $\mathcal{B}_n \setminus \mathcal{B}_n^*$ and $\mathcal{H} \setminus \mathcal{H}_n$ have an exponential decrease.

Construction of the test φ_n : To this end, we want to apply Lemma D.3 of Ghosal and Van der Vaart (2017), which directly allows to construct the test φ_n with appropriate control of error probabilities as described in (6.8) to test the true value against the whole of the alternative intersected with the sieve. To apply this lemma, we need to construct local tests with exponentially small errors to compare the true value with a subset of the alternative, centered at any $(\beta_1, \Gamma_1) \in \mathcal{B} \times \mathcal{H}$ which is adequately distant from the true value with respect to the average Rényi divergence. The other condition to apply this Lemma is that the minimum number N_n^* of these small subsets of the alternative needed to cover a sieve $\mathcal{B}_n^* \times \mathcal{H}_n$ is appropriately controlled in terms of ϵ_n .

First, the following lemma constructs an appropriate local test by employing the likelihood ratio to compare the true value with a subset of the alternative and by controlling the second order moment of the likelihood ratios in these small pieces of the alternative. For $(\beta_1, \Gamma_1) \in \mathcal{B} \times \mathcal{H}$, we denote by p_1 the associated density, and $\mathbb{E}_1$ and $\mathbb{P}_1$ the expectation and probability under p_1 .

Lemma 2. For a given positive sequence (γ_n) , $(\beta_1, \Gamma_1) \in \mathcal{B} \times \mathcal{H}$ such that $R_n(p_0, p_1) \geq \epsilon_n^2$, where $\epsilon_n = \sqrt{\frac{s_0 \log(p)}{n}}$, define

$$\begin{aligned}
\mathcal{F}_{1,n} &= \left\{ (\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : \frac{1}{n} \sum_{i=1}^n \|f_i(X_i \beta) - f_i(X_i \beta_1)\|_2^2 \leq \frac{\epsilon_n^2}{16\gamma_n}, \right. \\
&\quad \left. d_n(\Gamma, \Gamma_1) \leq \frac{\epsilon_n^2}{2J_n \gamma_n}, \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^{-1}\|_{sp} \leq \gamma_n \right\}.
\end{aligned}$$

Grant Assumptions A7-A9 and A4, then there exists a test $\bar{\varphi}_n$ such that

$$\mathbb{E}_0[\bar{\varphi}_n] \leq e^{-n\epsilon_n^2}, \quad \text{and} \quad \sup_{(\beta, \Gamma) \in \mathcal{F}_{1,n}} \mathbb{E}_{\beta, \Gamma}[1 - \bar{\varphi}_n] \leq e^{-n\epsilon_n^2/16}.$$

This lemma is demonstrated in Appendix 6.5.1.2.

Now, we still have to construct an appropriate sieve $\mathcal{B}_n^* \times \mathcal{H}_n$ such that the prior mass of $\mathcal{B}_n \setminus \mathcal{B}_n^*$ and $\mathcal{H} \setminus \mathcal{H}_n$ have an exponential decrease, and the minimum number N_n^* of the small subsets of the alternative needed to cover the sieve satisfies $\log(N_n^*) \lesssim n\epsilon_n^2$.

Define the sieve as follows:

$$\begin{aligned}\mathcal{B}_n^* &= \left\{ \beta \in \mathcal{B} \mid s_\beta \leq C_1 s_0, \|\beta\|_\infty \leq \frac{p^{L_2+2}}{K_n \|X\|_*} \right\}, \\ \mathcal{H}_n &= \left\{ \Gamma \in \mathcal{H} \mid n^{-M} \leq \rho_{\min}(\Gamma) \leq \rho_{\max}(\Gamma) \leq e^{Mn\epsilon_n^2} \right\},\end{aligned}$$

for a constant M , and define $\mathcal{F}_{1,n}$ as in Lemma 2 with $\gamma_n = n^M / \underline{\rho_Z}^2$. Remark that, with this choice of γ_n , the last condition $\max_{1 \leq i \leq n} \|\Delta_{\Gamma,i}^{-1}\|_{sp} \leq \gamma_n$ is always satisfy in the sieve. Indeed, $\|\Delta_{\Gamma,i}^{-1}\|_{sp} = \rho_{\max}(\Delta_{\Gamma,i}) = \rho_{\min}^{-1}(\Delta_{\Gamma,i})$. But by Assumption A8 and since $\Gamma \in \mathcal{H}_n$,

$$\rho_{\min}(\Delta_{\Gamma,i}) \geq \sigma^2 + \rho_{\min}(Z_i \Gamma Z_i^\top) \geq \rho_{\min}(\Gamma) \underline{\rho_Z}^2 \geq n^{-M} \underline{\rho_Z}^2.$$

So finally, $\max_{1 \leq i \leq n} \|\Delta_{\Gamma,i}^{-1}\|_{sp} \leq \gamma_n$ for $\Gamma \in \mathcal{H}_n$.

First, we show that $\mathcal{B}_n \setminus \mathcal{B}_n^*$ and $\mathcal{H} \setminus \mathcal{H}_n$ have an exponential decrease. Using Assumption A5, we obtain that

$$\begin{aligned}\Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) &= \Pi \left(\left\{ \beta \in \mathcal{B} \mid s_\beta \leq C_1 s_0, \|\beta\|_\infty > \frac{p^{L_2+2}}{K_n \|X\|_*} \right\} \right) \\ &= \sum_{S: s \leq C_1 s_0} \frac{\pi_p(s)}{\binom{p}{s}} \int_{\left\{ \beta_S: \|\beta_S\|_\infty > \frac{p^{L_2+2}}{K_n \|X\|_*} \right\}} g_S(\beta_S) d\beta_S \\ &\leq \sum_{S: s \leq C_1 s_0} \frac{(A_2 p^{-A_4})^s}{\binom{p}{s}} \int_{\left\{ \beta_S: \|\beta_S\|_\infty > \frac{p^{L_2+2}}{K_n \|X\|_*} \right\}} g_S(\beta_S) d\beta_S \\ &\leq \sum_{S: s \leq C_1 s_0} \frac{(A_2 p^{-A_4})^s}{\binom{p}{s}} \sum_{\ell \in S} \int_{\left\{ |\beta_\ell| > \frac{p^{L_2+2}}{K_n \|X\|_*} \right\}} \frac{\lambda}{2} e^{-\lambda|\beta_\ell|} d\beta_\ell\end{aligned}$$

Then, by using the tail probability of the Laplace distribution $\int_{|x|>t} \frac{\lambda}{2} e^{-\lambda|x|} dx = e^{-\lambda t}$ for every

$t > 0$, and since there is $\binom{p}{s}$ support S of size s , we obtain:

$$\begin{aligned}
\Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) &\leq \sum_{S: s \leq C_1 s_0} \frac{(A_2 p^{-A_4})^s}{\binom{p}{s}} s e^{-\lambda \frac{p^{L_2+2}}{K_n \|X\|_*}} \\
&\leq \sum_{s=1}^{C_1 s_0} s (A_2 p^{-A_4})^s e^{-\lambda \frac{p^{L_2+2}}{K_n \|X\|_*}} \\
&\leq C_1 s_0 e^{-\lambda \frac{p^{L_2+2}}{K_n \|X\|_*}} \sum_{s=1}^{C_1 s_0} (A_2 p^{-A_4})^s \\
&\lesssim s_0 e^{-\lambda \frac{p^{L_2+2}}{K_n \|X\|_*}} \lesssim s_0 e^{-\frac{\lambda}{L_1} p^2}
\end{aligned}$$

Thus, $\Pi(\mathcal{B}_n \setminus \mathcal{B}_n^*) e^{C n \epsilon_n^2} \xrightarrow{n \rightarrow \infty} 0$ for every $C > 0$ since $n \epsilon_n^2 = s_0 \log(p) = o(p^2)$.

Now,

$$\begin{aligned}
\Pi(\mathcal{H} \setminus \mathcal{H}_n) &= \Pi \left(\left\{ \Gamma \in \mathcal{H} \mid \rho_{\min}(\Gamma) < n^{-M} \text{ or } \rho_{\max}(\Gamma) > e^{M n \epsilon_n^2} \right\} \right) \\
&\leq \Pi \left(\left\{ \Gamma \in \mathcal{H} \mid \rho_{\min}(\Gamma) < n^{-M} \right\} \right) + \Pi \left(\left\{ \Gamma \in \mathcal{H} \mid \rho_{\max}(\Gamma) > e^{M n \epsilon_n^2} \right\} \right) \\
&= \Pi \left(\left\{ \Gamma \in \mathcal{H} \mid \rho_{\max}(\Gamma^{-1}) \geq n^M \right\} \right) + \Pi \left(\left\{ \Gamma \in \mathcal{H} \mid \rho_{\min}(\Gamma^{-1}) \leq e^{-M n \epsilon_n^2} \right\} \right) \\
&\leq b_1 e^{-b_2 n^{b_3 M}} \times b_4 e^{-b_5 M n \epsilon_n^2}
\end{aligned}$$

for some constants $b_1, b_2, b_3, b_4, b_5 > 0$ by Lemma 9.16 of Ghosal and Van der Vaart (2017) since $\Gamma^{-1} \sim \mathcal{W}_q(d, \Sigma^{-1})$. So, $\Pi(\mathcal{H} \setminus \mathcal{H}_n) e^{C n \epsilon_n^2} \xrightarrow{n \rightarrow \infty} 0$ for every $C > 0$, for M large enough.

Finally, we have to prove that the minimum number N_n^* of the small subsets of the alternative of the form $\mathcal{F}_{1,n}$ needed to cover the sieve satisfies $\log(N_n^*) \lesssim n \epsilon_n^2$. First, note that for every $\beta, \beta' \in \mathcal{B}$, by Assumption A1, the inequality $\|X\theta\|_2 \leq \|X\|_* \|\theta\|_1$, we have:

$$\begin{aligned}
\frac{1}{n} \sum_{i=1}^n \|f_i(X_i \beta) - f_i(X_i \beta')\|_2^2 &\leq \frac{1}{n} \sum_{i=1}^n K_n^2 \|X_i(\beta - \beta')\|_2^2 = \frac{K_n^2}{n} \|X(\beta - \beta')\|_2^2 \\
&\leq \frac{K_n^2}{n} \|X\|_*^2 \|\beta - \beta'\|_1^2 \\
&\leq \frac{p^2 K_n^2}{n} \|X\|_*^2 \|\beta - \beta'\|_\infty^2
\end{aligned}$$

Thus, we define

$$\mathcal{F}'_{1,n} = \left\{ (\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : \frac{p^2 K_n^2}{n} \|X\|_*^2 \|\beta - \beta'\|_\infty^2 + d_n^2(\Gamma, \Gamma_1) \leq \frac{1}{16 J_n^2 \gamma_n^2 n^3}, \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^{-1}\|_{sp} \leq \gamma_n \right\},$$

with the same (β_1, Γ_1) used in $\mathcal{F}_{1,n}$, and therefore we have $\mathcal{F}'_{1,n} \subset \mathcal{F}_{1,n}$. Thus, since N_n^* is the minimum number of the small subsets of the alternative of the form $\mathcal{F}_{1,n}$ needed to cover the

sieve $\mathcal{B}_n^* \times \mathcal{H}_n$, with $\mathcal{F}_{1,n} \supset \mathcal{F}'_{1,n}$, we have that N_n^* is bounded above by the minimum number of the small subsets of the alternative of the form $\mathcal{F}'_{1,n}$ needed to cover the sieve $\mathcal{B}_n^* \times \mathcal{H}_n$. This last minimum number is denoted by N'_n . In the following, for a pseudo-metric space $(\mathcal{F}, d)$, let $N(\epsilon, \mathcal{F}, d)$ denote the minimal number of ϵ -balls that cover $\mathcal{F}$.

Now, note that if $(\beta, \Gamma) \in \mathcal{B}_n^* \times \mathcal{H}_n$ with $\|\beta - \beta_1\|_\infty \leq \frac{1}{6npK_nJ_n\gamma_n\|X\|_*}$ and $d_n(\Gamma, \Gamma_1) \leq \frac{1}{6n^{3/2}J_n\gamma_n}$, then $(\beta, \Gamma) \in \mathcal{F}'_{1,n}$ (by using that the last condition is satisfy for all $\Gamma \in \mathcal{H}_n$). Thus,

$$N'_n \leq N\left(\frac{1}{6npK_nJ_n\gamma_n\|X\|_*}, \mathcal{B}_n^*, \|\cdot\|_\infty\right) \times N\left(\frac{1}{6n^{3/2}J_n\gamma_n}, \mathcal{H}_n, d_n\right).$$

Then,

$$\log(N_n^*) \leq \log N\left(\frac{1}{6npK_nJ_n\gamma_n\|X\|_*}, \mathcal{B}_n^*, \|\cdot\|_\infty\right) + \log N\left(\frac{1}{6n^{3/2}J_n\gamma_n}, \mathcal{H}_n, d_n\right). \quad (6.9)$$

Recall that $\mathcal{B}_n^* = \left\{ \beta \in \mathcal{B} \mid s_\beta \leq C_1 s_0, \|\beta\|_\infty \leq \frac{p^{L_2+2}}{K_n\|X\|_*} \right\}$, to cover $\mathcal{B}_n^*$, we have to choose at most $\lfloor C_1 s_0 \rfloor$ non-zero β coordinates and we need to recover a ball in $\mathbb{R}^{\lfloor C_1 s_0 \rfloor}$ with radius $\frac{p^{L_2+2}}{K_n\|X\|_*}$, with balls of radius $\frac{1}{6npK_nJ_n\gamma_n\|X\|_*}$. Therefore,

$$\begin{aligned} N\left(\frac{1}{6npK_nJ_n\gamma_n\|X\|_*}, \mathcal{B}_n^*, \|\cdot\|_\infty\right) &\leq \binom{p}{\lfloor C_1 s_0 \rfloor} (6p^{L_2+3}nJ_n\gamma_n)^{\lfloor C_1 s_0 \rfloor} \\ &\lesssim (6p^{L_2+4}nJ_n\gamma_n)^{\lfloor C_1 s_0 \rfloor} \text{ as } \binom{p}{\lfloor C_1 s_0 \rfloor} (\lfloor C_1 s_0 \rfloor)! \leq p^{\lfloor C_1 s_0 \rfloor} \end{aligned}$$

So, the first term in the right side of equation (6.9) is bounded by:

$$\log N\left(\frac{1}{6npK_nJ_n\gamma_n\|X\|_*}, \mathcal{B}_n^*, \|\cdot\|_\infty\right) \lesssim s_0 \log(p) = n\epsilon_n^2$$

since, by the assumption of Theorem 2, we have $\log(nJ_n) \lesssim \log(p)$.

Similarly, for the second term in the right side of equation (6.9), note that $\mathcal{H}_n \subset \left\{ \Gamma \in \mathcal{H} : \|\Gamma\|_F \leq \sqrt{q}e^{Mn\epsilon_n^2} \right\}$, and therefore by assumption A9:

$$\begin{aligned} \log N\left(\frac{1}{6n^{3/2}J_n\gamma_n}, \mathcal{H}_n, d_n\right) &\leq \log N\left(\frac{1}{6n^{3/2}J_n\gamma_n}, \left\{ \Gamma \in \mathcal{H} : \|\Gamma\|_F \leq \sqrt{q}e^{Mn\epsilon_n^2} \right\}, d_n\right) \\ &\leq \log N\left(\frac{1}{6n^{3/2}J_n\gamma_n\bar{\rho}_Z^2}, \left\{ \Gamma \in \mathcal{H} : \|\Gamma\|_F \leq \sqrt{q}e^{Mn\epsilon_n^2} \right\}, \|\cdot\|_F\right) \\ &\lesssim q^2 \log(\sqrt{q}e^{Mn\epsilon_n^2}n^{3/2}J_n\gamma_n) \lesssim n\epsilon_n^2 \end{aligned}$$

as $\log(nJ_n) \lesssim \log(p)$ and $\log(n) \lesssim \log(p)$.

Finally, $\log(N_n^*) \lesssim n\epsilon_n^2$. Thus, Lemma D.3 of [Ghosal and Van der Vaart \(2017\)](#) can be applied and gives that for every $\epsilon > \epsilon_n$, there exists a test φ_n satisfying

$$\mathbb{E}_0[\varphi_n] \leq 2e^{B_1 n \epsilon_n^2 - n \epsilon^2} \text{ and } \sup_{(\beta, \Gamma) \in \mathcal{B}_n^* \times \mathcal{H}_n : R_n^*(\beta, \Gamma) > \epsilon^2} \mathbb{E}_{(\beta, \Gamma)}[1 - \varphi_n] \leq e^{-n \epsilon^2 / 16}$$

for some constant $B_1 > 0$. Then, choosing $\epsilon = C_2 \epsilon_n$ for C_2 large enough, we obtain that the test φ_n satisfies (6.8), which concludes the proof, as demonstrated above. $\square$

6.4.3 Proof of Theorem 3

Proof of Theorem 3. The contraction rate of the posterior distribution with respect to the average Rényi divergence $R_n^*(\beta, \Gamma) = R_n(p_{\beta, \Gamma}, p_0)$ is provided by Theorem 2. Denote $\epsilon_n = \sqrt{\frac{s_0 \log(p)}{n}}$ this rate. We have that, for all $\beta_0 \in \mathcal{B}_0, \Gamma_0 \in \mathcal{H}_0$,

$$\mathbb{E}_0 \left[\Pi \left(\mathcal{R} | Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 1.$$

where $\mathcal{R} = \{(\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} : R_n^*(\beta, \Gamma) \leq C_2 \epsilon_n^2\}$. However, since for all $(\beta, \Gamma) \in \mathcal{B} \times \mathcal{H}$, $p_{\beta, \Gamma} = \prod_{i=1}^n p_{\beta, \Gamma, i}$, with $p_{\beta, \Gamma, i} = \mathcal{N}_{n_i}(f_i(X_i \beta), \Delta_{\Gamma, i})$ in Model (6.1), the average Rényi divergence is equal to:

$$\begin{aligned} R_n^*(\beta, \Gamma) &= R_n(p_{\beta, \Gamma}, p_0) = -\frac{1}{n} \sum_{i=1}^n \log \left(\int \sqrt{p_{\beta, \Gamma, i}(y_i) p_{0, i}(y_i)} dy_i \right) \\ &= -\frac{1}{n} \sum_{i=1}^n \log(1 - g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i})) + \frac{1}{4n} \sum_{i=1}^n \|(\Delta_{\Gamma, i} + \Delta_{\Gamma_0, i})^{-1/2} (f_i(X_i \beta) - f_i(X_i \beta_0))\|_2^2 \end{aligned}$$

where we used the Sherman-Morrison-Woodbury formula, with

$$g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i}) = 1 - \frac{\det(\Delta_{\Gamma, i})^{1/4} \det(\Delta_{\Gamma_0, i})^{1/4}}{\det((\Delta_{\Gamma, i} + \Delta_{\Gamma_0, i})/2)^{1/2}}.$$

Remark that for all $1 \leq i \leq n$, $g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i}) \geq 0$ since, with $\Delta_{\Gamma, i}^* = \Delta_{\Gamma_0, i}^{-1/2} \Delta_{\Gamma, i} \Delta_{\Gamma_0, i}^{-1/2}$:

$$\begin{aligned} \frac{\det(\Delta_{\Gamma, i})^{1/4} \det(\Delta_{\Gamma_0, i})^{1/4}}{\det((\Delta_{\Gamma, i} + \Delta_{\Gamma_0, i})/2)^{1/2}} &= \left(\frac{1}{2^{n_i}} \det(\Delta_{\Gamma, i}^{*1/2} + \Delta_{\Gamma, i}^{*-1/2}) \right)^{-1/2} \\ &= \left(\prod_{k=1}^{n_i} \frac{1}{2} (d_k^{1/2} + d_k^{-1/2}) \right)^{-1/2} \leq 1 \end{aligned}$$

where d_k are the eigenvalues of $\Delta_{\Gamma, i}^*$, and using that $\forall x \leq 0, x + x^{-1} \geq 2$.

Thus, by Theorem 2, this implies that, for $(\beta, \Gamma) \in \mathcal{R}$:

$$\begin{aligned} \epsilon_n^2 &\gtrsim -\frac{1}{n} \sum_{i=1}^n \log(1 - g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i})) + \frac{1}{4n} \sum_{i=1}^n \|(\Delta_{\Gamma, i} + \Delta_{\Gamma_0, i})^{-1/2} (f_i(X_i \beta) - f_i(X_i \beta_0))\|_2^2 \\ &\gtrsim -\frac{1}{n} \sum_{i=1}^n \log(1 - g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i})) \geq \frac{1}{n} \sum_{i=1}^n g^2(\Delta_{\Gamma, i}, \Delta_{\Gamma_0, i}) \end{aligned}$$

since $\log(1-x) \leq -x$ for all $x \geq 0$. Now, by Lemma 10 of Jeong and Ghosal (2021b), we obtain that $g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i}) \gtrsim \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$ if $g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i})$ is small enough for each $i \in \{1, \dots, n\}$. Thus, by defining $I_{n,\delta} = \{1 \leq i \leq n : g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i}) \geq \delta\}$, for some $\delta > 0$ small, we have:

$$\begin{aligned} \epsilon_n^2 &\gtrsim \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 + \frac{1}{n} \sum_{i \in I_{n,\delta}} g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i}) \\ &\gtrsim \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 = \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 - \frac{1}{n} \sum_{i \in I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq M_1 d_n^2(\Gamma, \Gamma_0) - M_1 \frac{|I_{n,\delta}|}{n} \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq M_1 d_n^2(\Gamma, \Gamma_0) - M_2 \epsilon_n^2 \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq (M_1 - M_3 \epsilon_n^2) d_n^2(\Gamma, \Gamma_0) \end{aligned}$$

where we used that $\frac{|I_{n,\delta}|}{n} \lesssim \epsilon_n^2$ since

$$\epsilon_n^2 \gtrsim \frac{1}{n} \sum_{i \notin I_{n,\delta}} g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i}) \gtrsim \frac{|I_{n,\delta}|}{n} \times \delta,$$

and thanks to Lemma B3 for the last inequality. Then, since $\epsilon_n^2 \xrightarrow{n \rightarrow \infty} 0$, $M_1 - M_3 \epsilon_n^2$ is bounded away from 0, the last inequation implies that $\epsilon_n \gtrsim d_n(\Gamma, \Gamma_0)$, which proves the first assertion of Theorem 3. Now, thanks to Lemma B3, we have also that $\epsilon_n \gtrsim \|\Gamma - \Gamma_0\|_F$, which proves the second assertion of Theorem 3.

Also, by Theorem 2, for $(\beta, \Gamma) \in \mathcal{R}$, since $g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i}) \geq 0$, we have that:

$$\begin{aligned} \epsilon_n^2 &\gtrsim -\frac{1}{n} \sum_{i=1}^n \log(1 - g^2(\Delta_{\Gamma,i}, \Delta_{\Gamma_0,i})) + \frac{1}{4n} \sum_{i=1}^n \|(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})^{-1/2} (f_i(X_i\beta) - f_i(X_i\beta_0))\|_2^2 \\ &\geq \frac{M_4}{4n} \sum_{i=1}^n \|(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})^{-1/2} (f_i(X_i\beta) - f_i(X_i\beta_0))\|_2^2 \\ &\geq \frac{M_4}{4n} \sum_{i=1}^n \rho_{\min}((\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})^{-1}) \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \end{aligned}$$

using that for A symmetric matrix, $\|Ax\|_2^2 \geq \rho_{\min}(A^2) \|x\|_2^2$ for all vector x .

Now, since $\rho_{\min}((\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})^{-1}) = \rho_{\max}^{-1}(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})$, we want to upper bound $\rho_{\max}(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i})$ uniformly across $i \in \{1, \dots, n\}$. Thus, by using Weyl's inequality:

$$\rho_{\max}(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i}) \leq \rho_{\max}(\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}) + 2\rho_{\max}(\Delta_{\Gamma_0,i}) \leq \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F + 2\overline{\rho_{\Delta}}$$

by Lemma B2 where $\overline{\rho_{\Delta}}$ denotes the uniform upper bound of the eigenvalues of $(\Delta_{\Gamma_0,i})_i$. Thus, by Lemma B3,

$$\max_{1 \leq i \leq n} \rho_{\max}(\Delta_{\Gamma,i} + \Delta_{\Gamma_0,i}) \leq \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F + 2\overline{\rho_{\Delta}} \leq d_n(\Gamma, \Gamma_0) + 2\overline{\rho_{\Delta}} \leq C_3 \epsilon_n + 2\overline{\rho_{\Delta}},$$

by the first assertion of Theorem 3.

Finally, we obtain that:

$$\epsilon_n^2 \geq \frac{M_4}{4n(C_3\epsilon_n + 2\overline{\rho\Delta})} \sum_{i=1}^n \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2,$$

and since $C_3\epsilon_n + 2\overline{\rho\Delta} \xrightarrow{n \rightarrow \infty} 2\overline{\rho\Delta}$, we finally obtain that for n large enough

$$\epsilon_n \gtrsim \sqrt{\frac{1}{n} \sum_{i=1}^n \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2},$$

which gives the last assertion of Theorem 3. □

6.4.4 Proof of Theorem 4

Proof of Theorem 4. Let us consider $\beta_0 \in \overline{\mathcal{B}}_0 \subset \mathcal{B}_0$ and $\Gamma_0 \in \mathcal{H}_0$. The contraction rate of the posterior distribution for the prediction term is provided by Theorem 3 and we have in particular that: for all $\beta_0 \in \overline{\mathcal{B}}_0, \Gamma_0 \in \mathcal{H}_0$,

$$\mathbb{E}_0 \left[\Pi \left(\mathcal{P}_n \middle| Y^{(n)} \right) \right] \xrightarrow{n \rightarrow \infty} 1,$$

where $\mathcal{P}_n = \{\beta : \frac{1}{n} \sum_{i=1}^n \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \lesssim \epsilon_n^2\}$ and $\epsilon_n = \sqrt{\frac{s_0 \log(p)}{n}}$.

Now, note that for $i \in \{1, \dots, n\}$, and $\delta > 0$, if we assume that $\|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2 \leq \delta$, then for all $j \in \{1, \dots, n_i\}$, $|f(X_i\beta, t_{ij}) - f(X_i\beta_0, t_{ij})| \leq \delta$ and so by Assumption A10

$$|f(X_i\beta, t_{ij}) - f(X_i\beta_0, t_{ij})| \gtrsim \|X_i(\beta - \beta_0)\|_2.$$

Therefore, $\|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \gtrsim n_i \|X_i(\beta - \beta_0)\|_2^2 \gtrsim \|X_i(\beta - \beta_0)\|_2^2$, as $n_i \geq 1$. Thus, using the same idea used in the proof of Theorem 3, we define

$$I_{n,\delta} = \{1 \leq i \leq n : \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 > \delta\},$$

for some $\delta > 0$ small. Thus, for $\beta \in \mathcal{P}_n$,

$$\begin{aligned} \epsilon_n^2 &\gtrsim \frac{1}{n} \sum_{i=1}^n \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \\ &\gtrsim \frac{1}{n} \sum_{i \in I_{n,\delta}} \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 + \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \\ &\gtrsim \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|X_i(\beta - \beta_0)\|_2^2 = \frac{1}{n} \sum_{i=1}^n \|X_i(\beta - \beta_0)\|_2^2 - \frac{1}{n} \sum_{i \in I_{n,\delta}} \|X_i(\beta - \beta_0)\|_2^2 \\ &\gtrsim \frac{1}{n} \|X(\beta - \beta_0)\|_2^2 - \frac{|I_{n,\delta}|}{n} \max_{1 \leq i \leq n} \|X_i(\beta - \beta_0)\|_2^2 \end{aligned}$$

Then, by the first assertion of Assumption A11, we have that

$$\max_{1 \leq i \leq n} \|X_i(\beta - \beta_0)\|_2^2 \leq \max_{1 \leq i \leq n} \|X_i\|_*^2 \|\beta - \beta_0\|_1^2 \lesssim \|\beta - \beta_0\|_1^2.$$

Remark that for β such as $s_\beta \leq C_1 s_0$, we have $s_{\beta - \beta_0} \leq s_\beta + s_0 \leq (C_1 + 1)s_0$, thus by Theorem 1 and by Definition 3, we obtain that: $\|\beta - \beta_0\|_1^2 \leq \frac{s_0 \|X(\beta - \beta_0)\|_2^2}{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}$, so

$$\max_{1 \leq i \leq n} \|X_i(\beta - \beta_0)\|_2^2 \lesssim \frac{s_0 \|X(\beta - \beta_0)\|_2^2}{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}.$$

Then, by using $|I_{n,\delta}| \lesssim n\epsilon_n^2$, since $\epsilon_n^2 \gtrsim \frac{1}{n} \sum_{i \in I_{n,\delta}} \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \gtrsim \frac{|I_{n,\delta}|}{n} \times \delta$, we have that:

$$\begin{aligned} \epsilon_n^2 &\gtrsim \left(1 - \frac{s_0 |I_{n,\delta}|}{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}\right) \frac{1}{n} \|X(\beta - \beta_0)\|_2^2 \\ &\gtrsim \left(1 - \frac{s_0 n \epsilon_n^2}{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}\right) \frac{1}{n} \|X(\beta - \beta_0)\|_2^2. \end{aligned}$$

Moreover, since $\beta_0 \in \bar{\mathcal{B}}_0$, $1 - \frac{s_0 n \epsilon_n^2}{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}$ is bounded away from 0. This implies that

$\|X(\beta - \beta_0)\|_2^2 \lesssim n\epsilon_n^2$ which gives the first assertion of Theorem 4.

Finally, by definition of the uniform compatibility number ϕ_1 and the smallest scaled singular value ϕ_2 , we obtain that:

$$\begin{aligned} \epsilon_n^2 &\gtrsim \frac{\|X\|_*^2 \phi_1^2 ((C_1 + 1)s_0)}{s_0 n} \|\beta - \beta_0\|_1^2, \\ \text{and } \epsilon_n^2 &\gtrsim \frac{\|X\|_*^2 \phi_2^2 ((C_1 + 1)s_0)}{n} \|\beta - \beta_0\|_2^2, \end{aligned}$$

which proves the last two assertions of the theorem. $\square$

6.5 Appendices

6.5.1 Proofs of technical lemmas

6.5.1.1 Proof of Lemma 1

First, we define a Kullback-Leibler neighbourhood around $p_{0,i}$

$$\mathcal{D}_n = \left\{ (\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} \mid \sum_{i=1}^n K(p_{0,i}, p_{\beta, \Gamma, i}) \leq c_1 \log(n), \sum_{i=1}^n V(p_{0,i}, p_{\beta, \Gamma, i}) \leq c_1 \log(n) \right\}$$

for a constant c_1 large enough. Then, by Lemma 10 of Ghosal and Van Der Vaart (2007), we have that, for every $C > 0$,

$$\mathbb{P}_0 \left(\int_{\mathcal{D}_n} \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \leq e^{-(1+C)c_1 \log(n)} \Pi(\mathcal{D}_n) \right) \leq \frac{1}{C^2 c_1 \log(n)}. \quad (6.10)$$

The proof of the lemma consists in showing that there exists $M > 0$ such that $e^{-(1+C)c_1 \log(n)} \Pi(\mathcal{D}_n) \gtrsim \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))}$. Indeed by combining this result with Inequality (6.10), since $\int_{\mathcal{D}_n} \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \leq \int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma)$, we have that

$$\begin{aligned} & \mathbb{P}_0 \left(\int \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \geq \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))} \right) \\ & \geq \mathbb{P}_0 \left(\int_{\mathcal{D}_n} \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \geq \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))} \right) \\ & \geq \mathbb{P}_0 \left(\int_{\mathcal{D}_n} \Lambda_n(\beta, \Gamma) d\Pi(\beta, \Gamma) \leq e^{-(1+C)c_1 \log(n)} \Pi(\mathcal{D}_n) \right) \\ & \geq 1 - \frac{1}{C^2 c_1 \log(n)} \xrightarrow{n \rightarrow \infty} 1, \end{aligned}$$

that concludes the proof. Thus, it remains to show that

$e^{-(1+C)c_1 \log(n)} \Pi(\mathcal{D}_n) \geq \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))}$, or, more precisely, we need to exhibit a lower bound of $\Pi(\mathcal{D}_n)$.

In Model (6.1), we have that $p_{\beta, \Gamma, i} = \mathcal{N}(f_i(X_i \beta), \Delta_{\Gamma, i})$, with $\Delta_{\Gamma, i} = Z_i \Gamma Z_i^\top + \sigma^2 I_{n_i}$. By Lemma 9 of Jeong and Ghosal (2021b), the Kullback-Leibler divergence and variation of the i -th individual are respectively expressed as:

$$\begin{aligned} K(p_{0,i}, p_{\beta, \Gamma, i}) &= \frac{1}{2} \left[\log \left(\frac{|\Delta_{\Gamma, i}|}{|\Delta_{\Gamma_0, i}|} \right) + \text{Tr}(\Delta_{\Gamma_0, i} \Delta_{\Gamma, i}^{-1}) - n_i + \left\| \Delta_{\Gamma, i}^{-1/2} (f_i(X_i \beta) - f_i(X_i \beta_0)) \right\|_2^2 \right], \\ V(p_{0,i}, p_{\beta, \Gamma, i}) &= \frac{1}{2} \left[\text{Tr} \left(\Delta_{\Gamma_0, i} \Delta_{\Gamma, i}^{-1} \Delta_{\Gamma_0, i} \Delta_{\Gamma, i}^{-1} \right) - 2 \text{Tr}(\Delta_{\Gamma_0, i} \Delta_{\Gamma, i}^{-1}) + n_i \right] + \\ & \quad \left\| \Delta_{\Gamma_0, i}^{1/2} \Delta_{\Gamma, i}^{-1} (f_i(X_i \beta) - f_i(X_i \beta_0)) \right\|_2^2. \end{aligned}$$

Then, by denoting $\rho_{i,k}$, for $k = 1, \dots, n_i$, the eigenvalues of $\Delta_{\Gamma_0, i}^{1/2} \Delta_{\Gamma, i}^{-1} \Delta_{\Gamma_0, i}^{1/2}$, we obtain that

$$\begin{aligned} K(p_{0,i}, p_{\beta, \Gamma, i}) &= \frac{1}{2} \left[- \sum_{k=1}^{n_i} \log(\rho_{i,k}) - \sum_{k=1}^{n_i} (1 - \rho_{i,k}) + \left\| \Delta_{\Gamma, i}^{-1/2} (f_i(X_i \beta) - f_i(X_i \beta_0)) \right\|_2^2 \right], \\ V(p_{0,i}, p_{\beta, \Gamma, i}) &= \frac{1}{2} \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 + \left\| \Delta_{\Gamma_0, i}^{1/2} \Delta_{\Gamma, i}^{-1} (f_i(X_i \beta) - f_i(X_i \beta_0)) \right\|_2^2. \end{aligned}$$

Our goal is to find a lower bound of $\Pi(\mathcal{D}_n)$, so we want to find an upper bound of

$\sum_{i=1}^n K(p_{0,i}, p_{\beta, \Gamma, i})$ and $\sum_{i=1}^n V(p_{0,i}, p_{\beta, \Gamma, i})$.

Let us first focus on the term $V(p_{0,i}, p_{\beta, \Gamma, i})$. By Lemma 10 of Jeong and Ghosal (2021b), we obtain that:

$$\sum_{k=1}^{n_i} (1 - \rho_{i,k}^{-1})^2 \leq \rho_{\min}^{-2}(\Delta_{\Gamma_0, i}) \|\Delta_{\Gamma, i} - \Delta_{\Gamma_0, i}\|_F^2.$$

By using Weyl's inequality and Assumptions A9 and A3, Lemma B2 shows that there exist $\underline{\rho}_0 > 0$ and $\overline{\rho}_0 > 0$ such that:

$$\underline{\rho}_0 \leq \min_i \rho_{\min}(\Delta_{\Gamma_0, i}) \leq \max_i \rho_{\max}(\Delta_{\Gamma_0, i}) \leq \overline{\rho}_0. \quad (6.11)$$

Thus

$$\max_i \sum_{k=1}^{n_i} (1 - \rho_{i,k}^{-1})^2 \leq \underline{\rho}_0^{-2} \max_i \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2,$$

and in particular, $\max_{i,k} (1 - \rho_{i,k}^{-1})^2 \leq \underline{\rho}_0^{-2} \max_i \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$. Thus, if $\max_i \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \rightarrow 0$ on $\mathcal{D}_n$, we will have that $\max_{i,k} |1 - \rho_{i,k}^{-1}| \rightarrow 0$, that is each $\rho_{i,k}$ tends to 1 and so:

$$\sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \lesssim \sum_{k=1}^{n_i} (1 - \rho_{i,k}^{-1})^2 \leq \underline{\rho}_0^{-2} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \lesssim \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2, \quad (6.12)$$

where the first inequality is due to $|1 - x^{-1}| \lesssim |1 - x| \lesssim |1 - x^{-1}|$ for $x \rightarrow 1$, which enables to bound the first term of $V(p_{0,i}, p_{\beta,\Gamma,i})$.

Now, prove that $\max_i \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$ tends to 0 on $\mathcal{D}_n$. We introduce the set $I_{n,\delta} = \{1 \leq i \leq n \mid \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \geq \delta\}$, for $\delta > 0$ small. We denote by $|I_{n,\delta}|$ its cardinal. Then for $(\beta, \Gamma) \in \mathcal{D}_n$, since $\sum_{i=1}^n V(p_{0,i}, p_{\beta,\Gamma,i}) \leq c_1 \log(n)$, we have that, on the one hand:

$$\sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \leq c_1 \log(n),$$

and on the other hand,

$$\begin{aligned} \sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 &= \sum_{i \in I_{n,\delta}} \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 + \sum_{i \notin I_{n,\delta}} \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \\ &\gtrsim \delta |I_{n,\delta}| + \sum_{i \notin I_{n,\delta}} \sum_{k=1}^{n_i} \left(1 - \frac{1}{\rho_{i,k}}\right)^2 \end{aligned}$$

since for $i \notin I_{n,\delta}$, $\sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 < \delta$, for $\delta > 0$ small, so each $|1 - \rho_{i,k}|$ is less than $\sqrt{\delta}$, and we have that $|1 - x^{-1}| \lesssim |1 - x| \lesssim |1 - x^{-1}|$ for $x \rightarrow 1$, so $|1 - \rho_{i,k}| \gtrsim |1 - \rho_{i,k}^{-1}|$. Then, by using Lemma 10 of Jeong and Ghosal (2021b), we obtain that

$$\sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \gtrsim \delta |I_{n,\delta}| + \sum_{i \notin I_{n,\delta}} \frac{1}{\rho_{\max}^2(\Delta_{\Gamma_0,i})} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$$

By (6.11), we obtain that

$$c_1 \log(n) \geq \sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \gtrsim \delta |I_{n,\delta}| + \frac{1}{\underline{\rho}_0^2} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$$

that is equivalent to

$$\frac{\log(n)}{n} \gtrsim \frac{1}{n} \sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 \gtrsim \delta \frac{|I_{n,\delta}|}{n} + \frac{1}{n \underline{\rho}_0^2} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$$

In particular, $\delta \frac{|I_{n,\delta}|}{n} \lesssim \frac{\log(n)}{n}$ for $\delta > 0$ small that implies $|I_{n,\delta}| \lesssim \log(n)$. We have also that $\frac{\log(n)}{n} \gtrsim \frac{1}{n\rho_0^2} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2$, that is

$$\frac{\log(n)}{n} \gtrsim \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2. \quad (6.13)$$

Then,

$$\begin{aligned} \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 &= \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 - \frac{1}{n} \sum_{i \in I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq d_n^2(\Gamma, \Gamma_0) - \frac{|I_{n,\delta}|}{n} \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \end{aligned}$$

By Lemma B3, with Assumptions A7, A8 and A9, for $\Gamma_1, \Gamma_2 \in \mathcal{H}$, we have that

$$\max_i \|\Delta_{\Gamma_1,i} - \Delta_{\Gamma_2,i}\|_F^2 \lesssim \|\Gamma_1 - \Gamma_2\|_F^2 \lesssim d_n^2(\Gamma_1, \Gamma_2).$$

Thus, there exist some constants $M_1 > 0$ and $M_2 > 0$:

$$\begin{aligned} \frac{1}{n} \sum_{i \notin I_{n,\delta}} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 &\geq M_1 \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 - \frac{|I_{n,\delta}|}{n} \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq \left(M_1 - \frac{|I_{n,\delta}|}{n} \right) \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \\ &\geq \left(M_1 - M_2 \frac{\log(n)}{n} \right) \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \end{aligned} \quad (6.14)$$

since $|I_{n,\delta}| \lesssim \log(n)$. Finally, by combining (6.13) and (6.14), since $\frac{\log(n)}{n} \rightarrow 0$, we obtain that, on $\mathcal{D}_n$,

$$\max_{1 \leq i \leq n} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \lesssim \frac{\log(n)}{n} \xrightarrow{n \rightarrow \infty} 0.$$

Hence, by using Equation (6.12), we obtain that:

$$\begin{aligned} \frac{1}{n} \sum_{i=1}^n V(p_{0,i}, p_{\beta,\Gamma,i}) &= \frac{1}{2n} \sum_{i=1}^n \sum_{k=1}^{n_i} (1 - \rho_{i,k})^2 + \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_0,i}^{1/2} \Delta_{\Gamma,i}^{-1} (f_i(X_i\beta) - f_i(X_i\beta_0))\|_2^2 \\ &\lesssim \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 + \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_0,i}^{1/2}\|_{sp}^2 \|\Delta_{\Gamma,i}^{-1}\|_{sp}^2 \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \end{aligned}$$

since $\|Ax\|_2 \leq \|A\|_{sp} \|x\|_2$. Then, by Lemma B2, we have that $\|\Delta_{\Gamma_0,i}^{1/2}\|_{sp}^2 = \rho_{\max}(\Delta_{\Gamma_0,i}) \lesssim 1$ for each $i \in \{1, \dots, n\}$. Moreover, on $\mathcal{D}_n$,

$$\begin{aligned} \|\Delta_{\Gamma,i}^{-1}\|_{sp} &= \rho_{\max}(\Delta_{\Gamma,i}^{-1}) = \rho_{\max}(\Delta_{\Gamma_0,i}^{-1/2} \Delta_{\Gamma_0,i}^{1/2} \Delta_{\Gamma,i}^{-1} \Delta_{\Gamma_0,i}^{1/2} \Delta_{\Gamma_0,i}^{-1/2}) \\ &\leq \rho_{\max}(\Delta_{\Gamma_0,i}^{1/2} \Delta_{\Gamma,i}^{-1} \Delta_{\Gamma_0,i}^{1/2}) \|\Delta_{\Gamma_0,i}^{-1/2}\|_{sp}^2 \text{ by Lemma B1} \\ &\lesssim \rho_{\max}(\Delta_{\Gamma_0,i}^{1/2} \Delta_{\Gamma,i}^{-1} \Delta_{\Gamma_0,i}^{1/2}) = \max_k \rho_{i,k} \lesssim 1 \end{aligned}$$

since $\|\Delta_{\Gamma_0,i}^{-1/2}\|_{sp}^2 = \rho_{\max}(\Delta_{\Gamma_0,i}^{-1}) = \rho_{\min}(\Delta_{\Gamma_0,i})^{-1} \leq \rho_0^{-1}$ by Lemma B2, and since each $\rho_{i,k}$ tends to 1 on $\mathcal{D}_n$. Then, using Assumption A1, that is each f_i is K_n -Lipschitz, we deduce that:

$$\frac{1}{n} \sum_{i=1}^n V(p_{0,i}, p_{\beta,\Gamma,i}) \lesssim d_n^2(\Gamma, \Gamma_0) + \frac{K_n^2}{n} \|X(\beta - \beta_0)\|_2^2.$$

Let us now focus on the term $K(p_{0,i}, p_{\beta,\Gamma,i})$. We have shown that on $\mathcal{D}_n$ each $|1 - \rho_{i,k}|$ is small for each i and k , so: $\log(\rho_{i,k}) = \log(1 - (1 - \rho_{i,k})) \sim -(1 - \rho_{i,k}) - \frac{(1 - \rho_{i,k})^2}{2}$, and so $-\log(\rho_{i,k}) - (1 - \rho_{i,k}) \sim \frac{(1 - \rho_{i,k})^2}{2}$. Thus, by using Inequality (6.12):

$$\begin{aligned} \frac{1}{n} K(p_{0,i}, p_{\beta,\Gamma,i}) &= \frac{1}{2n} \sum_{i=1}^n \left[-\sum_{k=1}^{n_i} \log(\rho_{i,k}) - \sum_{k=1}^{n_i} (1 - \rho_{i,k}) + \|\Delta_{\Gamma,i}^{-1/2} (f_i(X_i\beta) - f_i(X_i\beta_0))\|_2^2 \right] \\ &\lesssim \frac{1}{2n} \sum_{i=1}^n \left(\frac{1}{2} \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 + \|\Delta_{\Gamma,i}^{-1/2}\|_{sp}^2 \|f_i(X_i\beta) - f_i(X_i\beta_0)\|_2^2 \right) \\ &\lesssim d_n^2(\Gamma, \Gamma_0) + \frac{K_n^2}{n} \|X(\beta - \beta_0)\|_2^2, \end{aligned}$$

by Assumption A1 and since $\|\Delta_{\Gamma,i}^{-1}\|_{sp} \lesssim 1$ on $\mathcal{D}_n$. Thus, we obtain finally that $\frac{1}{n} K(p_{0,i}, p_{\beta,\Gamma,i})$ and $\frac{1}{n} V(p_{0,i}, p_{\beta,\Gamma,i})$ are bounded above by $d_n^2(\Gamma, \Gamma_0) + \frac{K_n^2}{n} \|X(\beta - \beta_0)\|_2^2$ up to a multiplicative constant. Then, for c_1 large enough,

$$\begin{aligned} \Pi(\mathcal{D}_n) &= \Pi \left((\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} \left| \frac{1}{n} \sum_{i=1}^n K(p_{0,i}, p_{\beta,\Gamma,i}) \leq c_1 \frac{\log(n)}{n}, \frac{1}{n} \sum_{i=1}^n V(p_{0,i}, p_{\beta,\Gamma,i}) \leq c_1 \frac{\log(n)}{n} \right. \right) \\ &\geq \Pi \left((\beta, \Gamma) \in \mathcal{B} \times \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) + \frac{K_n^2}{n} \|X(\beta - \beta_0)\|_2^2 \leq 2 \frac{\log(n)}{n} \right. \right) \\ &\geq \Pi \left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) \Pi \left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n} \|X(\beta - \beta_0)\|_2^2 \leq \frac{\log(n)}{n} \right. \right) \\ &\geq \Pi \left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) \Pi \left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n} \|X\|_*^2 \|\beta - \beta_0\|_1^2 \leq \frac{\log(n)}{n} \right. \right) \end{aligned}$$

since $\|X\theta\|_2 \leq \|X\|_* \|\theta\|_1$. For the first term, we have that:

$$\begin{aligned} \Pi \left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) &= \Pi \left(\Gamma \in \mathcal{H} \left| \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \leq \frac{\log(n)}{n} \right. \right) \\ &\geq \Pi \left(\Gamma \in \mathcal{H} \left| \max_i \|\Delta_{\Gamma,i} - \Delta_{\Gamma_0,i}\|_F^2 \leq \frac{\log(n)}{n} \right. \right) \\ &\geq \Pi \left(\Gamma \in \mathcal{H} \left| \|\Gamma - \Gamma_0\|_F^2 \leq \frac{1}{\rho_Z^4} \frac{\log(n)}{n} \right. \right) \\ &= \Pi \left(\Gamma \in \mathcal{H} \left| \|\Gamma - \Gamma_0\|_F \leq \frac{1}{\rho_Z^2} \sqrt{\frac{\log(n)}{n}} \right. \right) \end{aligned}$$

by Lemma B3 and Assumption A9. Thus, by Assumption A3:

$$\begin{aligned}\|\Gamma - \Gamma_0\|_F &= \|\Gamma_0^{1/2}(\Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2} - Id)\Gamma_0^{1/2}\|_F \\ &\leq \|\Gamma_0\|_{sp}\|\Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2} - Id\|_F \\ &\leq \bar{\rho}_\Gamma\|\Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2} - Id\|_F\end{aligned}$$

By using Lemma B4, we obtain that:

$$\begin{aligned}\Pi\left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) \\ \geq \Pi\left(\Gamma \in \mathcal{H} \left| \|\Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2} - Id\|_F \leq \bar{\rho}_\Gamma^{-1}\bar{\rho}_Z^{-2}\sqrt{\frac{\log(n)}{n}} \right. \right) \\ \geq \Pi\left(\Gamma \in \mathcal{H} \left| \bigcap_{k=1}^q \left\{ 1 \leq \rho_k(\Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2}) \leq 1 + \bar{\rho}_\Gamma^{-1}\bar{\rho}_Z^{-2}\sqrt{\frac{\log(n)}{qn}} \right\} \right. \right)\end{aligned}$$

Then, denoting by $A = \Gamma_0^{-1/2}\Gamma\Gamma_0^{-1/2} \in \mathbb{R}^{q \times q}$, since $\Gamma \sim \mathcal{IW}_q(d, \Sigma)$, we know that $A \sim \mathcal{IW}_q(d, \Gamma_0^{-1/2}\Sigma\Gamma_0^{-1/2})$, and so $A^{-1} \sim \mathcal{W}_q(d, \Gamma_0^{1/2}\Sigma^{-1}\Gamma_0^{1/2})$. Then, using Lemma 6.3 of Ning et al. (2020) on the eigenvalues of a Wishart distribution, we obtain that, for large n and since q is fixed:

$$\begin{aligned}\Pi\left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) &\geq \left(\frac{a_1 t e^2 d}{8\sqrt{\pi}} \right)^{-q} \left(\frac{2dq}{e a_1 t} \right)^{-dq/2} \left(\frac{d}{2e} \right)^{-q^2/2} \times \\ &\quad (\det(\Gamma_0^{1/2}\Sigma^{-1}\Gamma_0^{1/2}))^{-d/2} \exp\left(-\frac{a_1(1+t)\text{Tr}(\Gamma_0^{-1/2}\Sigma\Gamma_0^{-1/2})}{2} \right)\end{aligned}$$

with $a_1 = \left(1 + \bar{\rho}_\Gamma^{-1}\bar{\rho}_Z^{-2}\sqrt{\frac{\log(n)}{qn}} \right)^{-1}$ and $t = \bar{\rho}_\Gamma^{-1}\bar{\rho}_Z^{-2}\sqrt{\frac{\log(n)}{qn}}$.

Finally, for n large enough, we have that:

$$\log\left(\Pi\left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right)\right) \gtrsim -\log(n).$$

Concerning the second term $\Pi\left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n}\|X\|_*^2\|\beta - \beta_0\|_1^2 \leq \frac{\log(n)}{n} \right. \right)$ in the lower bound of $\Pi(\mathcal{D}_n)$, by defining $\mathcal{B}_{S_0, n} = \left\{ \beta_{S_0} \in \mathbb{R}^{s_0} \left| \frac{K_n}{\sqrt{n}}\|X\|_*\|\beta_{S_0} - \beta_{0, S_0}\|_1 \leq \sqrt{\frac{\log(n)}{n}} \right. \right\}$ we have that:

$$\begin{aligned}\Pi\left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n}\|X\|_*^2\|\beta - \beta_0\|_1^2 \leq \frac{\log(n)}{n} \right. \right) &\geq \Pi\left(S = S_0, \beta \in \mathcal{B} \left| \frac{K_n}{\sqrt{n}}\|X\|_*\|\beta - \beta_0\|_1 \leq \sqrt{\frac{\log(n)}{n}} \right. \right) \\ &\geq \frac{\pi_p(s_0)}{\binom{p}{s_0}} \int_{\mathcal{B}_{S_0, n}} g_{S_0}(\beta_{S_0}) d\beta_{S_0} \\ &\geq \frac{\pi_p(s_0)}{\binom{p}{s_0}} e^{-\lambda\|\beta_0\|_1} \int_{\mathcal{B}_{S_0, n}} g_{S_0}(\beta_{S_0} - \beta_{0, S_0}) d\beta_{S_0}\end{aligned}$$

because g_S is the Laplace distribution so satisfy the inequality $g_{S_0}(\beta_{S_0}) \geq e^{-\lambda\|\beta_0\|_1} g_{S_0}(\beta_{S_0} - \beta_{0,S_0})$. Then, since $s_0 > 0$ by Assumption A4 and using the equation (6.2) of Castillo et al. (2015), we obtain that:

$$\begin{aligned} \int_{\mathcal{B}_{S_0,n}} g_{S_0}(\beta_{S_0} - \beta_{0,S_0}) d\beta_{S_0} &\geq e^{-\lambda \frac{\sqrt{\log(n)}}{K_n \|X\|_*}} \frac{\left(\lambda \frac{\sqrt{\log(n)}}{K_n \|X\|_*} \right)^{s_0}}{s_0!} \\ &\geq e^{-L_3 \sqrt{\frac{\log(n)}{n}}} \frac{\left(\frac{\sqrt{\log(n)}}{L_1 p^{L_2}} \right)^{s_0}}{s_0!} \end{aligned}$$

by Assumption A6. We deduce that, by using $\binom{p}{s_0} s_0! \leq p^{s_0}$:

$$\begin{aligned} \Pi \left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n} \|X\|_*^2 \|\beta - \beta_0\|_1^2 \leq \frac{\log(n)}{n} \right. \right) \\ &\geq \frac{\pi_p(s_0)}{\binom{p}{s_0}} e^{-\lambda\|\beta_0\|_1} e^{-L_3 \sqrt{\frac{\log(n)}{n}}} \frac{\left(\frac{\sqrt{\log(n)}}{L_1 p^{L_2}} \right)^{s_0}}{s_0!} \\ &\geq \pi_p(s_0) \log(n)^{s_0/2} \exp \left(-(L_2 - 1)s_0 \log(p) - L_3 \sqrt{\frac{\log(n)}{n}} - s_0 \log(L_1) - \lambda\|\beta_0\|_1 \right) \\ &\geq \pi_p(s_0) \exp \left(-\tilde{C} s_0 \log(p) \right) \end{aligned}$$

for a constant $\tilde{C}$ since $s_0 + \sqrt{\frac{\log(n)}{n}} + s_0 \log(p) \lesssim s_0 \log(p)$ and $\lambda\|\beta_0\|_1 \lesssim s_0 \log(p)$ by Assumption A2. Finally, there exists a constant L such that:

$$\begin{aligned} \Pi(\mathcal{D}_n) &\geq \Pi \left(\Gamma \in \mathcal{H} \left| d_n^2(\Gamma, \Gamma_0) \leq \frac{\log(n)}{n} \right. \right) \Pi \left(\beta \in \mathcal{B} \left| \frac{K_n^2}{n} \|X\|_*^2 \|\beta - \beta_0\|_1^2 \leq \frac{\log(n)}{n} \right. \right) \\ &\geq \pi_p(s_0) e^{-L(s_0 \log(p) + \log(n))} \end{aligned}$$

Thus, we have shown that $e^{-(1+C)c_1 \log(n)} \Pi(\mathcal{D}_n) \geq \pi_p(s_0) e^{-M(s_0 \log(p) + \log(n))}$ for some constant M , as required to conclude the proof of Lemma 1.

6.5.1.2 Proof of Lemma 2

Let $(\beta_1, \Gamma_1) \in \mathcal{B} \times \mathcal{H}$ such that $R_n(p_0, p_1) \geq \epsilon_n^2$. First, for testing $H_0: p = p_0$ against $H_1: p = p_1$, consider the most powerful test $\bar{\varphi}_n = \mathbb{1}_{\Lambda_n(\beta_1, \Gamma_1) \geq 1}$ given by the Neyman-Pearson lemma, where $\Lambda_n(\beta_1, \Gamma_1) = \frac{p_1}{p_0}$ be the likelihood ratio of p_1 and p_0 . Thus,

$$\begin{aligned} \mathbb{E}_0[\bar{\varphi}_n] &= \mathbb{P}_0(\sqrt{\Lambda_n(\beta_1, \Gamma_1)} \geq 1) = \int \mathbb{1}_{\sqrt{p_1(y)} \geq \sqrt{p_0(y)}} p_0(y) dy \\ &\leq \int \sqrt{p_0(y) p_1(y)} dy = e^{-n R_n(p_0, p_1)} \leq e^{-n \epsilon_n^2} \end{aligned}$$

by assumption on (β_1, Γ_1) . This proves the first result of the lemma.

Then, for the second part of the lemma, note that:

$$\mathbb{E}_1[1 - \bar{\varphi}_n] = \mathbb{P}_1(\sqrt{\Lambda_n(\beta_1, \Gamma_1)} \leq 1) \leq \int \sqrt{p_0(y)p_1(y)} dy \leq e^{-n\epsilon_n^2} \quad (6.15)$$

However, by using Cauchy-Schwarz inequality:

$$\begin{aligned} \mathbb{E}_{\beta, \Gamma}[1 - \bar{\varphi}_n] &= \int (1 - \bar{\varphi}_n(y)) \frac{p_{\beta, \Gamma}(y)}{p_1(y)} dp_1(y) \\ &\leq \mathbb{E}_1[1 - \bar{\varphi}_n]^{1/2} \mathbb{E}_1 \left[\left(\frac{p_{\beta, \Gamma}}{p_1} \right)^2 \right]^{1/2} \\ &\leq e^{-n\epsilon_n^2/2} \mathbb{E}_1 \left[\left(\frac{p_{\beta, \Gamma}}{p_1} \right)^2 \right]^{1/2} \end{aligned}$$

by Equation (6.15). Therefore, the test $\bar{\varphi}_n$ can also have exponentially small error of type II at other alternatives if we can controlled the second term: we want to show that

$$\mathbb{E}_1 \left[\left(\frac{p_{\beta, \Gamma}}{p_1} \right)^2 \right]^{1/2} \leq e^{7n\epsilon_n^2/16} \text{ for every } (\beta, \Gamma) \in \mathcal{F}_{1,n}.$$

Recall that here $p_{\beta, \Gamma} = \prod_{i=1}^n \mathcal{N}_{n_i}(f_i(X_i\beta), \Delta_{\Gamma, i})$, where $\Delta_{\Gamma, i} = Z_i \Gamma Z_i^\top + \sigma^2 I_{n_i}$. By denoting $\Delta_{\Gamma, i}^* = \Delta_{\Gamma, i}^{-1/2} \Delta_{\Gamma_1, i} \Delta_{\Gamma, i}^{-1/2}$, then for $(\beta, \Gamma) \in \mathcal{F}_{1,n}$, if $2\Delta_{\Gamma, i}^* - Id$ and $2Id - \Delta_{\Gamma, i}^{*-1}$ are non-singular matrices for every $i \in \{1, \dots, n\}$, we can show that:

$$\begin{aligned} \mathbb{E}_1 \left[\left(\frac{p_{\beta, \Gamma}}{p_1} \right)^2 \right] &= \prod_{i=1}^n \left[\det(\Delta_{\Gamma, i}^*)^{1/2} \det(2Id - \Delta_{\Gamma, i}^{*-1})^{-1/2} \right] \times \\ &\quad \exp \left\{ \sum_{i=1}^n \|(2\Delta_{\Gamma, i}^* - Id)^{-1/2} \Delta_{\Gamma, i}^{-1/2} (f_i(X_i\beta) - f_i(X_i\beta_1))\|_2^2 \right\}. \end{aligned} \quad (6.16)$$

Let us now prove that these matrices are non-singular. We have, for all $k \leq n_i$,

$$\max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^* - Id\|_{sp} = \max_{1 \leq i \leq n} \rho_{\max}(|\Delta_{\Gamma, i}^* - Id|) \geq \max_{1 \leq i \leq n} |\rho_k(\Delta_{\Gamma, i}^*) - 1|. \quad (6.17)$$

Note that

$$\begin{aligned} \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^* - Id\|_{sp} &= \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^{-1/2} (\Delta_{\Gamma_1, i} - \Delta_{\Gamma, i}) \Delta_{\Gamma, i}^{-1/2}\|_{sp} \\ &\leq \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^{-1}\|_{sp} \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma, i}\|_F \\ &\leq \max_{1 \leq i \leq n} \|\Delta_{\Gamma, i}^{-1}\|_{sp} d_n(\Gamma, \Gamma_1) \\ &\leq \frac{\epsilon_n^2}{2J_n} \end{aligned}$$

by Lemma B3 and since $(\beta, \Gamma) \in \mathcal{F}_{1,n}$. Thus, for all $k \leq n_i$, $\max_{1 \leq i \leq n} |\rho_k(\Delta_{\Gamma, i}^*) - 1| \leq \frac{\epsilon_n^2}{2J_n}$. We deduce that

$$1 - \frac{\epsilon_n^2}{2J_n} \leq \min_{1 \leq i \leq n} \rho_{\min}(\Delta_{\Gamma,i}^*) \leq \max_{1 \leq i \leq n} \rho_{\max}(\Delta_{\Gamma,i}^*) \leq 1 + \frac{\epsilon_n^2}{2J_n}. \quad (6.18)$$

Therefore, since $\frac{\epsilon_n^2}{2J_n} \xrightarrow{n \rightarrow \infty} 0$ by Assumption A4, and for all $k \leq n_i$,

$\rho_k(2\Delta_{\Gamma,i}^* - Id) = 2\rho_k(\Delta_{\Gamma,i}^*) - 1$ and $\rho_k(2Id - \Delta_{\Gamma,i}^{*-1}) = 2 - \rho_k(\Delta_{\Gamma,i}^{*-1}) = 2 - \rho_k^{-1}(\Delta_{\Gamma,i}^*)$, we deduce that $2\Delta_{\Gamma,i}^* - Id$ and $2Id - \Delta_{\Gamma,i}^{*-1}$ are non-singular on $\mathcal{F}_{1,n}$ for every $i \in \{1, \dots, n\}$.

For concluding the proof, it remains to bound the right side term of (6.16). By using (6.18) and the inequalities $(1 - x^2)/(1 - 2x) \leq 1 + 3x$ for $x > 0$ small, and $1 + x \leq e^x$, we obtain for n large enough:

$$\begin{aligned} \det(\Delta_{\Gamma,i}^*)^{1/2} \det(2Id - \Delta_{\Gamma,i}^{*-1})^{-1/2} &= \left(\prod_{k=1}^{n_i} \rho_k(\Delta_{\Gamma,i}^*) \right)^{1/2} \left(\prod_{k=1}^{n_i} 2 - \rho_k^{-1}(\Delta_{\Gamma,i}^*) \right)^{-1/2} \\ &= \left(\prod_{k=1}^{n_i} \frac{\rho_k(\Delta_{\Gamma,i}^*)}{2 - \rho_k^{-1}(\Delta_{\Gamma,i}^*)} \right)^{1/2} \\ &\leq \left(\frac{1 - \frac{\epsilon_n^4}{4J_n^2}}{1 + \frac{\epsilon_n^2}{J_n}} \right)^{n_i/2} \leq \left(1 + 3\frac{\epsilon_n^2}{2J_n} \right)^{n_i/2} \leq \exp \left(3\frac{n_i \epsilon_n^2}{4J_n} \right) \leq e^{3\epsilon_n^2/4} \end{aligned}$$

Moreover, for n large enough,

$$\begin{aligned} &\sum_{i=1}^n \|(2\Delta_{\Gamma,i}^* - Id)^{-1/2} \Delta_{\Gamma,i}^{-1/2} (f_i(X_i \beta) - f_i(X_i \beta_1))\|_2^2 \\ &\leq \max_{1 \leq i \leq n} \|(2\Delta_{\Gamma,i}^* - Id)^{-1}\|_{sp} \max_{1 \leq i \leq n} \|\Delta_{\Gamma,i}^{-1}\|_{sp} \sum_{i=1}^n \|f_i(X_i \beta) - f_i(X_i \beta_1)\|_2^2 \\ &\leq 2\gamma_n \frac{n\epsilon_n^2}{16\gamma_n} = \frac{n\epsilon_n^2}{8} \end{aligned}$$

since $(\beta, \Gamma) \in \mathcal{F}_{1,n}$. Finally, by using (6.16), we conclude that

$$\mathbb{E}_1 \left[\left(\frac{p_{\beta, \Gamma}}{p_1} \right)^2 \right]^{1/2} \leq e^{3n\epsilon_n^2/8} e^{n\epsilon_n^2/16} = e^{7n\epsilon_n^2/16}, \text{ and so } \sup_{(\beta, \Gamma) \in \mathcal{F}_{1,n}} \mathbb{E}_{\beta, \Gamma}[1 - \bar{\varphi}_n] \leq e^{-n\epsilon_n^2/16},$$

which concludes the proof.

6.5.2 Useful lemmas

Lemma B1. For A and B two matrices, $\rho_{\min}(B)\|A\|_{sp}^2 \leq \rho_{\max}(ABA^\top) \leq \rho_{\max}(B)\|A\|_{sp}^2$.

Proof. By the Courant–Fischer–Weyl min-max principle,

$$\begin{aligned}
\rho_{\max}(ABA^\top) &= \max_{x \neq 0} \frac{\langle ABA^\top x, x \rangle}{\|x\|^2} \\
&= \max_{x \neq 0} \frac{\langle BA^\top x, A^\top x \rangle}{\|x\|^2} \\
&= \max_{x \neq 0} \frac{\langle BA^\top x, A^\top x \rangle}{\|A^\top x\|^2} \frac{\|A^\top x\|^2}{\|x\|^2} \\
&\leq \rho_{\max}(B) \max_{x \neq 0} \frac{\|A^\top x\|^2}{\|x\|^2} \\
&= \rho_{\max}(B) \rho_{\max}(AA^\top) \\
&= \rho_{\max}(B) \|A\|_{sp}^2
\end{aligned}$$

We obtain the other inequality with similar arguments. $\square$

Lemma B2. *Grant Assumptions A3 and A9. Thus, $\Delta_{\Gamma_0, i} := Z_i \Gamma_0 Z_i^\top + \sigma^2 I_{n_i}$ satisfies:*

$$1 \lesssim \min_i \rho_{\min}(\Delta_{\Gamma_0, i}) \leq \max_i \rho_{\max}(\Delta_{\Gamma_0, i}) \lesssim 1$$

Proof. By the Weyl's inequality, for $1 \leq i \leq n$,

$$\rho_{\min}(\Delta_{\Gamma_0, i}) \geq \rho_{\min}(Z_i \Gamma_0 Z_i^\top) + \sigma^2 \geq \sigma^2$$

since $Z_i \Gamma_0 Z_i^\top$ is a positive definite matrix. Thus $\min_i \rho_{\min}(\Delta_{\Gamma_0, i}) \geq \sigma^2$, otherwise, $\min_i \rho_{\min}(\Delta_{\Gamma_0, i}) \gtrsim 1$.

For the other inequality, by the Weyl's inequality, we have that

$$\rho_{\max}(\Delta_{\Gamma_0, i}) \leq \rho_{\max}(Z_i \Gamma_0 Z_i^\top) + \sigma^2.$$

Then, by Lemma B1, we have that:

$$\rho_{\max}(Z_i \Gamma_0 Z_i^\top) \leq \rho_{\max}(\Gamma_0) \|Z_i\|_{sp}^2$$

and by Assumptions A3 and A9,

$$\max_i \rho_{\max}(\Delta_{\Gamma_0, i}) \leq \rho_{\max}(\Gamma_0) \max_i \|Z_i\|_{sp}^2 + \sigma^2 \lesssim 1.$$

$\square$

Lemma B3. *For $\Gamma_1, \Gamma_2 \in \mathcal{H}$, under Assumptions A7, A8 and A9, we have that*

$$\max_i \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 \lesssim \|\Gamma_1 - \Gamma_2\|_F^2 \lesssim d_n^2(\Gamma_1, \Gamma_2) = \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2.$$

Proof. First,

$$\|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 = \|Z_i(\Gamma_1 - \Gamma_2)Z_i^\top\|_F^2 \leq \|Z_i\|_{sp}^4 \|\Gamma_1 - \Gamma_2\|_F^2,$$

since $\|AB\|_F \leq \|A\|_{sp}\|B\|_F$, and by Assumption A9, we have that

$$\max_{1 \leq i \leq n} \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 \lesssim \|\Gamma_1 - \Gamma_2\|_F^2.$$

Then, by Assumption A8, for each i such that $n_i \geq q$, $Z_i^\top Z_i$ is invertible and

$$\|\Gamma_1 - \Gamma_2\|_F^2 = \|(Z_i^\top Z_i)^{-1} Z_i^\top Z_i (\Gamma_1 - \Gamma_2) Z_i^\top Z_i (Z_i^\top Z_i)^{-1}\|_F^2 \leq \|Z_i (\Gamma_1 - \Gamma_2) Z_i^\top\|_F^2 \|(Z_i^\top Z_i)^{-1} Z_i^\top\|_{sp}^4.$$

Then, by Assumption A7, we have that

$$\begin{aligned} \max_{1 \leq i \leq n} \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 &\lesssim \|\Gamma_1 - \Gamma_2\|_F^2 \leq \frac{1}{\sum_{i=1}^n \mathbb{1}_{n_i \geq q}} \sum_{i: n_i \geq q} \|Z_i (\Gamma_1 - \Gamma_2) Z_i^\top\|_F^2 \|(Z_i^\top Z_i)^{-1} Z_i^\top\|_{sp}^4 \\ &\lesssim \frac{1}{n} \sum_{i: n_i \geq q} \|Z_i (\Gamma_1 - \Gamma_2) Z_i^\top\|_F^2 \|(Z_i^\top Z_i)^{-1} Z_i^\top\|_{sp}^4 \\ &\lesssim \frac{1}{n} \sum_{i=1}^n \|Z_i (\Gamma_1 - \Gamma_2) Z_i^\top\|_F^2 = \frac{1}{n} \sum_{i=1}^n \|\Delta_{\Gamma_1, i} - \Delta_{\Gamma_2, i}\|_F^2 = d_n^2(\Gamma_1, \Gamma_2). \end{aligned}$$

where the last inequality uses assumptions A8 and A9. $\square$

Lemma B4. For a positive definite symmetric matrix $A \in \mathbb{R}^{q \times q}$ such as its eigenvalues satisfy $1 \leq \rho_1(A) \leq \dots \leq \rho_q(A) \leq 1 + \frac{\epsilon}{\sqrt{q}}$, then $\|A - I_q\|_F \leq \epsilon$.

Proof. Observe that

$$\begin{aligned} \|A - I_q\|_F \leq \epsilon &\Leftrightarrow \text{Tr}((A - I_q)^2) \leq \epsilon^2 \quad \text{since } A \text{ is symmetric} \\ &\Leftrightarrow \sum_{k=1}^q \rho_k(A - I_q)^2 \leq \epsilon^2 \\ &\Leftrightarrow \sum_{k=1}^q (\rho_k(A) - 1)^2 \leq \epsilon^2. \end{aligned}$$

By assumption, for $1 \leq k \leq q$, we have that $0 \leq \rho_k(A) - 1 \leq \frac{\epsilon}{\sqrt{q}}$ and then

$\max_{1 \leq k \leq q} (\rho_k(A) - 1)^2 \leq \frac{\epsilon^2}{q}$. Hence, since $\sum_{k=1}^q (\rho_k(A) - 1)^2 \leq q \times \max_{1 \leq k \leq q} (\rho_k(A) - 1)^2 \leq \epsilon^2$ and so $\|A - I_q\|_F \leq \epsilon$. $\square$

Partie III

Travaux en cours, conclusion et perspectives

Travaux en cours : Métamodélisation pour la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes

En amélioration des plantes, les fonctions de régression employées dans les modèles non-linéaires à effets mixtes deviennent de plus en plus sophistiquées pour décrire avec précision les mécanismes biologiques, ce qui les rend coûteuses en termes de temps de calcul. Ce problème se pose notamment pour les modèles de culture qui s'appuient sur des modèles écophysiologiques précis pour reproduire les différentes phases de développement des plantes (voir chapitre 1). Au-delà de l'amélioration des plantes, les fonctions de régression utilisées en pharmacométrie ou en médecine par exemple deviennent également de plus en plus complexes, reposant sur des équations différentielles ordinaires multidimensionnelles ou des équations aux dérivées partielles (Chatterjee et al., 2012; Ribba et al., 2012).

Dans ces cas-là, la méthode SAEMVS proposée dans le chapitre 5 atteint ses limites. En effet, cette procédure repose sur l'algorithme MCMC-SAEM, qui est itératif et nécessite plusieurs évaluations de la fonction de régression non-linéaire à chaque itération, rendant le temps d'exécution prohibitif lorsque l'évaluation est coûteuse. Des approches de métamodélisation pour approcher la fonction de régression pourraient permettre de réduire drastiquement les temps de calcul.

Ces travaux sont en cours, et font notamment partie du sujet du stage de Kathya Viviana Gavilanes Guerrero, étudiante en Master 2 *Data Science* à l'Université d'Évry, co-encadrée par Maud Delattre et moi-même depuis mi-avril 2024. La section 7.1 commence par motiver ce travail à partir d'un jeu de données biologiques complexe. La section 7.2 présente des outils de métamodélisation déjà utilisés dans les modèles non-linéaires à effets mixtes. Enfin, la section 7.3 explore des premières pistes pour coupler la méthode SAEMVS avec ces approches de métamodélisation afin de permettre la sélection de variables dans des modèles non-linéaires à effets mixtes complexes.

Table des matières

7.1 Motivations pratiques	133
7.2 Outils de métamodélisation dans les modèles non-linéaires à effets mixtes complexes	134
7.2.1 Inférence dans les modèles non-linéaires à effets mixtes complexes	134
7.2.2 Méta-modèles par krigeage gaussien	135
7.3 Couplage SAEMVS et métamodélisation	137
7.3.1 Première approche	137

7.1 Motivations pratiques

Dans le cadre du projet ANR *Stat4Plant*, présenté dans le chapitre 1, nous collaborons avec Renaud Rincent (Université Paris-Saclay, INRAE, CNRS, AgroParisTech, GQE - Le Moulon) sur un ensemble de données biologiques comprenant les dates d'épiaison de 198 variétés de blé d'hiver dans 10 environnements différents. En plus de ces données, nous disposons d'informations génétiques pour chaque variété, couvrant environ 118000 marqueurs, ainsi que de données météorologiques pour chaque environnement. Par ailleurs, nous avons accès à un modèle de culture qui simule la date d'épiaison de chaque variété dans chaque environnement en se basant sur des paramètres génétiques et des données environnementales. Ce modèle de culture (Weir et al., 1984) est composé de cinq sous-modèles interagissant entre eux, qui modélisent le développement phénologique, la production de tiges et de feuilles, la production de racines, l'interception de la lumière et la photosynthèse, ainsi que la répartition de la matière sèche et la croissance des grains. La plupart des variables d'entrée du modèle de culture peuvent être fixées à partir des connaissances expertes. Les biologistes ont identifié deux paramètres d'intérêt comme étant spécifiques à chaque variété, et qui influencent la date d'épiaison. Ces deux paramètres sont la vernalisation, processus par lequel l'exposition prolongée des plantes à des températures froides stimule la floraison, et la photopériode, durée quotidienne d'exposition à la lumière nécessaire à la plante pour réguler son cycle de croissance et de floraison.

L'objectif de cette étude est de sélectionner les marqueurs génétiques impliqués dans le processus d'épiaison. La sélection génétique se concentre donc uniquement sur ces deux paramètres d'intérêt. Nous pouvons alors modéliser ce problème par un modèle non-linéaire à effets mixtes comme suit. Pour $1 \leq i \leq n$, $1 \leq j \leq n_i$, la date d'épiaison y_{ij} de la i -ème variété dans l'environnement j est modélisée par :

$$\begin{aligned} y_{ij} &= f(\varphi_i, E_j) + \varepsilon_{ij}, \quad \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i &= X_i \beta + \xi_i, \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_2(0, \Gamma), \end{aligned}$$

où $n = 198$, $n_1 = \dots = n_n = 10$, f est le modèle de culture dépendant de paramètres génétiques $\varphi_i \in \mathbb{R}^2$ non observés, respectivement relatif à la vernalisation et à la photopériode, et des données environnementales disponibles E_j . La matrice X_i est de dimension $2 \times p$, où $p = 118000$ est le nombre de marqueurs génétiques. Cependant, de par sa conception, le modèle de culture f est coûteux à évaluer.

Pour répondre à cette problématique, nous souhaitons appliquer la procédure SAEMVS, développée dans le chapitre 5. Cependant, l'algorithme MCMC-SAEM utilisé dans cette procédure nécessite de nombreuses évaluations de la fonction de régression lors des étapes de simulation et de mise à jour des statistiques exhaustives, de l'ordre de $N_{tot} \times K$, où N_{tot} représente le nombre total d'observations et K le nombre d'itérations de l'algorithme. Lorsque la fonction de régression est coûteuse à évaluer, ceci se traduit par un coût computationnel considérable et cette démarche devient prohibitive. En effet, évaluer le modèle de culture f en un point nécessite environ 260 fois plus de temps que l'évaluation d'une fonction logistique, comme utilisée dans l'étude de simulations du chapitre 5. Par exemple, pour seulement 500 covariables, une

itération de l'algorithme MCMC-SAEM avec le modèle de culture prendrait environ 5 minutes, contre seulement 1 seconde pour une fonction logistique. Si l'algorithme converge en $K = 300$ itérations, le calcul complet pour le modèle de culture nécessiterait plus de 24 heures pour un modèle de culture comme celui-ci et pour seulement 500 marqueurs génétiques parmi les 118000, comparé à seulement 5 minutes pour une fonction logistique.

La construction de modèles approximatifs devient alors une solution adéquate pour surmonter ce problème. Ces modèles approximatifs sont appelés méta-modèles (Fang et al., 2005). Ici, l'objectif d'un méta-modèle est donc d'être "proche" en un certain sens du modèle de culture réel mais plus rapide à exécuter, de sorte que l'algorithme MCMC-SAEM soit utilisable en pratique.

7.2 Outils de métamodélisation dans les modèles non-linéaires à effets mixtes complexes

Dans la suite, un modèle non-linéaire à effets mixtes est qualifié de "complexe" si l'évaluation de sa fonction de régression est coûteuse en temps de calculs. Cela inclut, par exemple, les fonctions de régression définies par un modèle de culture ainsi que les solutions d'équations aux dérivées partielles. Dans cette section, nous introduisons des outils de métamodélisation adaptés aux modèles non-linéaires à effets mixtes.

7.2.1 Inférence dans les modèles non-linéaires à effets mixtes complexes

Plusieurs travaux ont étudié l'inférence dans les modèles non-linéaires à effets mixtes complexes en approchant la fonction de régression pour rendre l'algorithme MCMC-SAEM utilisable en pratique. Donnet and Samson (2007) ont traité le cas où la fonction de régression est la solution d'une équation différentielle ordinaire, en l'approchant par un schéma numérique. Grenier et al. (2014) se sont concentrés sur les modèles basés sur un grand nombre d'équations différentielles ordinaires et sur les équations différentielles partielles, en utilisant également un schéma numérique, mais restreint à une grille prédéfinie. L'approximation de la solution globale est alors obtenue par interpolation linéaire entre deux points de la grille. Cependant, si la fonction de régression non-linéaire est complexe, cette interpolation linéaire peut entraîner un biais dans les estimations. Finalement, Barbillon et al. (2017) proposent de modéliser la fonction de régression comme la réalisation d'un processus gaussien et de l'approcher comme le processus gaussien conditionné à certaines évaluations pré-établies du modèle. Cette approche, connue sous le nom de méta-modèle par krigeage gaussien, est particulièrement intéressante pour les modèles de culture qui nous concernent car elle peut s'appliquer pour une fonction de régression générale. Par ailleurs, cette approche par processus gaussien semble particulièrement adaptée à notre algorithme d'inférence car nous pourrions rester dans le modèle exponentiel, ce qui facilite l'implémentation de l'algorithme SAEM. L'approche par krigeage gaussien est détaillée dans la section suivante.

7.2.2 Méta-modèles par krigeage gaussien

Considérons un modèle mathématique déterministe décrivant la relation entre des variables d'entrées $\mathbf{x} = (x_1, \dots, x_s)^\top$ et une variable de sortie $z \in \mathbb{R}$:

$$z = f(\mathbf{x}), \quad \mathbf{x} = (x_1, \dots, x_s)^\top \in T,$$

où T est l'espace des variables d'entrée, et f est une fonction de régression. Supposons que la fonction f n'a pas d'expression analytique et est coûteuse à évaluer en temps de calculs. L'objectif de la métamodélisation est de prédire la valeur du modèle $f(\mathbf{x})$, en un point $\mathbf{x}$ quelconque, par $g(\mathbf{x})$, où g est un modèle approximatif plus rapide à calculer. Les étapes d'une procédure de métamodélisation par krigeage gaussien (Matheron, 1963; Krige, 1951; Sacks et al., 1989) sont détaillées ci-après.

1. **Pré-calcul** : Nous supposons que nous avons les ressources nécessaires pour évaluer précisément la fonction f en n_D points distincts d'un plan d'expérience fixé $D = \{\mathbf{x}_1, \dots, \mathbf{x}_{n_D}\}$. Notons $z_k = f(\mathbf{x}_k)$ ces évaluations exactes, pour $k = 1, \dots, n_D$. Nous cherchons à construire une approximation g de f à partir de l'échantillon $\{(\mathbf{x}_k, z_k), k = 1, \dots, n_D\}$.
2. **Inférence du processus gaussien** : Supposons que f est la réalisation d'un processus gaussien F_λ : pour tout $\mathbf{x} \in T$,

$$F_\lambda(\mathbf{x}) = \sum_{j=1}^L \alpha_j h_j(\mathbf{x}) + \zeta(\mathbf{x}) = H(\mathbf{x})^\top \alpha + \zeta(\mathbf{x}),$$

où les fonctions de base h_j sont fixées a priori, $H = (h_1, \dots, h_L)^\top$ l'ensemble des fonctions de base, $\alpha = (\alpha_1, \dots, \alpha_L)^\top$ un vecteur de paramètres inconnu à estimer, et ζ un processus gaussien centré de fonction de covariance $\text{Cov}(\zeta(\mathbf{x}), \zeta(\mathbf{x}')) = \sigma_r^2 r_\gamma(\mathbf{x}, \mathbf{x}')$, où $\sigma_r^2 > 0$ est la variance du processus et r_γ est une fonction de corrélation prédéfinie dépendant de certains paramètres γ . Ce modèle est appelé modèle universel de krigeage. On note $\lambda = (\alpha, \sigma_r^2, \gamma)$ le vecteur des paramètres inconnus à estimer. Ainsi, la construction du méta-modèle passe par l'estimation de λ à partir de l'échantillon $\{(\mathbf{x}_k, z_k), k = 1, \dots, n_D\}$ par maximum de vraisemblance du processus gaussien. Notons $\hat{\lambda}$ cette estimation.

3. **Prédiction** : On suppose alors que f est la réalisation du processus gaussien $F = F_{\hat{\lambda}}$. Cependant, la fonction f n'est pas directement approchée par le processus F , mais par le processus conditionné à interpoler les évaluations non bruitées de f obtenues lors de l'étape de pré-calculs, noté F^D . Plus précisément, on définit le processus F^D comme le processus F conditionné à $F(\mathbf{x}_k) = z_k$, pour $k = 1, \dots, n_D$. Le processus F^D est encore un processus gaussien, défini par sa moyenne m_D et sa fonction de covariance C_D . Nous pouvons calculer explicitement ses fonctions de moyenne et de covariance en utilisant la propriété conditionnelle des distributions gaussiennes multivariées. Ainsi $m_D(\mathbf{x})$ fournit une approximation de $f(\mathbf{x})$ pour tout $\mathbf{x} \in T$ et la fonction de variance $\mathbf{x} \mapsto C_D(\mathbf{x}, \mathbf{x})$ mesure la confiance dans la précision de cette approximation.

Un exemple de prédiction par krigeage gaussien de la fonction logistique suivante est proposée sur la figure 1 :

$$f(x) = \frac{100}{1 + \exp\left(-\frac{20 - x}{2}\right)}.$$

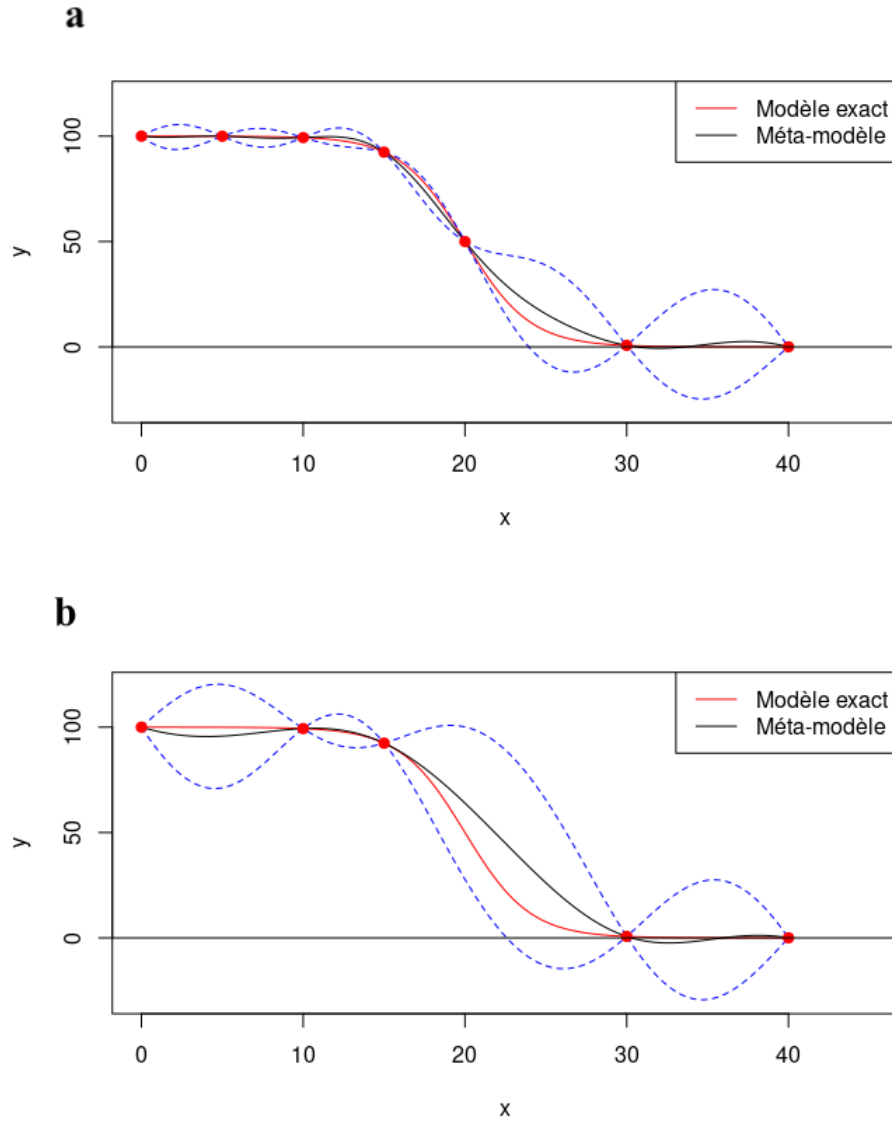


Fig. 1: Exemples de prédiction d'une fonction logistique avec le package *DiceKriging* pour deux plans d'expérience différents (fig. a et fig. b). La courbe rouge représente la fonction à approcher, la courbe noire est la valeur moyenne m_D de la prédiction par krigeage gaussien, et les courbes en pointillés bleues représentent les intervalles de confiance à 95%. Les points rouges correspondent au plan d'expérience D .

Plus précisément, la figure 1 présente deux prédictions de la fonction logistique par krigeage gaussien à partir d'un plan d'expérience différent. Pour la figure **a**, en utilisant 7 points bien choisis pour entraîner le méta-modèle, nous pouvons remarquer que le modèle exact et le modèle approché sont très proches. La qualité d'approximation pourrait encore être améliorée en ajoutant

plus de points dans le plan d'expérience vers la fin de la courbe notamment. Sur la figure **b**, nous pouvons remarquer que le choix du plan d'expérience est moins pertinent et dégrade la qualité de la prédiction.

Ainsi, cette approche par krigeage gaussien repose sur plusieurs éléments à bien calibrer pour garantir une prédiction satisfaisante. Premièrement, comme nous venons de le voir, le choix du plan d'expérience D est essentiel. Des critères de remplissage de l'espace des variables d'entrée sont proposés dans Fang et al. (2005). De plus, le choix de la fonction de corrélation r_γ et des fonctions de base H doit satisfaire la régularité et la tendance supposée de la fonction f si l'on dispose d'informations sur cette fonction. Voir Fang et al. (2005) pour plus de détails sur ces choix. Par exemple, une fonction de corrélation couramment utilisée est le noyau gaussien défini par $r_\gamma(\mathbf{x}, \mathbf{x}') = \exp(-\gamma \|\mathbf{x} - \mathbf{x}'\|^2)$, tandis que les fonctions de base sont souvent linéaires ou polynomiales de faible degré.

7.3 Couplage SAEMVS et métamodélisation

Les implémentations de l'algorithme MCMC-SAEM dans le cas des méta-modèles mixtes proposées par Barbillon et al. (2017) n'ont pas été évaluées dans le cadre de la sélection des effets fixes en grande dimension. Ainsi, pour permettre la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes complexes, nous proposons de coupler l'approche SAEMVS, développée dans le chapitre 5, avec une approche de métamodélisation par krigeage gaussien. Notons que l'ajout de priors, de la procédure de seuillage et de la grande dimension des covariables pourraient modifier les performances des implémentations proposées par Barbillon et al. (2017).

7.3.1 Première approche

Considérons le modèle non-linéaire à effets mixtes suivant avec les mêmes notations que le chapitre 5 : pour $1 \leq i \leq n$, $1 \leq j \leq n_i$,

$$\begin{cases} y_{ij} = f(\varphi_i, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \end{cases} \quad (7.1a)$$

$$\begin{cases} \varphi_i = \mu + X_i\beta + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma). \end{cases} \quad (7.1b)$$

Une première solution pour coupler ces deux approches consiste à (i) construire un méta-modèle pour f par krigeage gaussien, puis (ii) modifier l'algorithme SAEMVS en remplaçant les appels à f par des appels à son approximation. Nous avons commencé par considérer le méta-modèle mixte simple, qui approche f par la moyenne du processus gaussien conditionné obtenue à partir du package *DiceKriging* (Roustant et al., 2012). Ce méta-modèle est en effet dit "simple" car il permet de conserver le modèle dans la famille exponentielle courbe et qu'il ne prend pas en compte l'erreur d'approximation sur f .

Dans un premier temps, considérons le cas simple de la croissance logistique suivante, dans lequel l'algorithme d'inférence MCMC-SAEM peut directement être appliqué :

$$f(\varphi_i, t_{ij}) = \frac{200}{1 + \exp\left(-\frac{t_{ij} - \varphi_i}{400}\right)}.$$

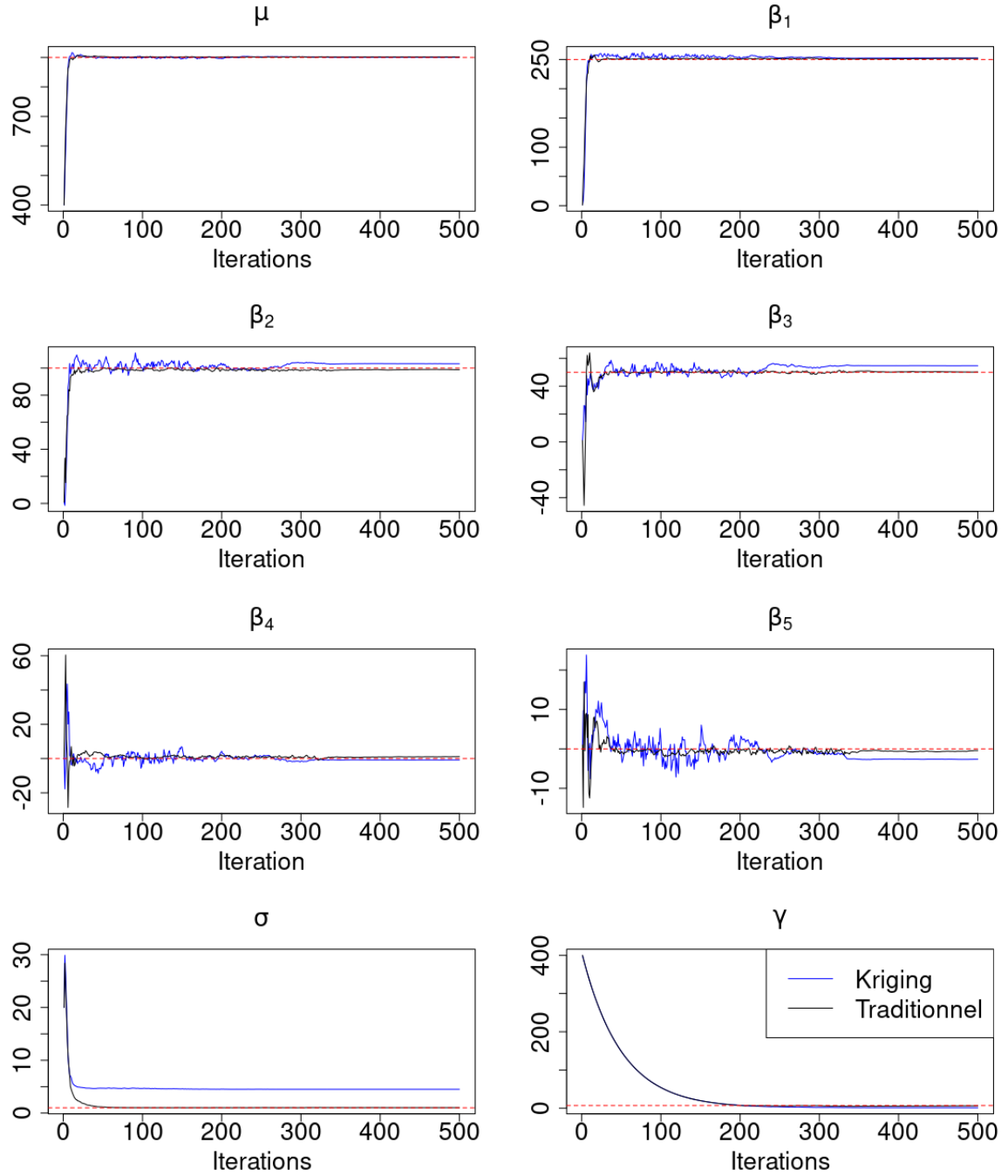


Fig. 2: Comparaison des graphes de convergence de l'algorithme MCMC-SAEM appliqué au modèle exact (courbes noires) et au méta-modèle (courbes bleues) pour différents paramètres : μ , les 5 premières coordonnées de β , σ et γ . Les courbes en pointillés rouges correspondent à la vraie valeur du paramètre.

La figure 2 compare les graphes de convergence de l'algorithme MCMC-SAEM à partir du modèle exact et à partir de l'approximation méta-modèle. Les valeurs des paramètres utilisées sont les suivantes : $n = 100$, $n_1 = \dots = n_{100} = 10$, $p = 10$, $q = 1$, $\mu = 900$, $\beta = (250, 100, 50, 0, \dots, 0)^\top$, $\sigma = 1$. Étant donné que $q = 1$, la matrice Γ est ici une constante, on la note γ et on choisit $\gamma = 7$. Pour le krigeage gaussien, nous avons utilisé un plan d'expérience composé de 30 points repartis selon une méthode d'échantillonnage par hypercube Latin (Fang et al., 2005). Nous avons choisi un noyau gaussien et des fonctions de base polynomiale de degré inférieur ou égal à 2. Il est important de noter que ces premiers tests de convergence de l'algorithme ont été effectués en petite dimension ($p < n$) et pour le calcul de l'estimateur du maximum de vraisemblance. Cette situation se retrouve dans la procédure SAEMVS lors de l'étape de calcul du critère eBIC (voir Algorithme 7 du chapitre 5).

Ainsi, sur les graphes de convergence de la figure 2, nous pouvons remarquer que, pour certains paramètres, la convergence de l'algorithme MCMC-SAEM est moins rapide lorsqu'il est appliqué sur le méta-modèle et peut mener à certains biais dans les estimations. Ceci était attendu étant donné que nous n'utilisons pas exactement la véritable valeur de la fonction de régression à chaque itération, surtout que nous utilisons pour le krigeage un plan d'expérience de seulement 30 points dans un espace de dimension 2. Toutefois, de manière générale, la figure 2 montre que les performances de convergence du couplage de l'algorithme MCMC-SAEM avec le méta-modèle sont très satisfaisantes. À noter qu'ici le nombre d'individus est seulement de 100. Le biais observé pourrait donc sans doute être corrigé en considérant plus d'individus, 200 par exemple.

7.3.2 Plan de travail

Nous avons identifié plusieurs aspects à explorer pour valider ce couplage.

- **Qualité d'approximation du méta-modèle :** Tout d'abord, indépendamment de la procédure SAEMVS, nous voulons tester la qualité d'approximation de f par le méta-modèle. En effet, cette première étape est essentielle pour espérer sélectionner les variables pertinentes. Pour cela, nous allons considérer un contexte plus complexe qu'une fonction logistique, comme par exemple un modèle de culture. Nous savons que le choix du plan d'expérience D est crucial pour obtenir une bonne approximation de f . Ainsi, nous envisageons de tester l'influence à la fois de la répartition des points et du nombre de points dans D . De même, le choix des fonctions de base H pourrait aussi être exploré en simulation. Nous espérons ainsi identifier une stratégie garantissant une certaine qualité d'approximation de f .
- **Performance de sélection :** Nous envisageons par la suite d'évaluer numériquement les performances en sélection de SAEMVS couplée au méta-modèle. En effet, dans un premier temps nous allons considérer une fonction de régression dont l'évaluation n'est pas trop coûteuse pour pouvoir appliquer la procédure SAEMVS sur le modèle exact et avoir une référence à laquelle se rapporter. Nous pourrions alors étudier la proximité entre les supports de β sélectionnés par SAEMVS appliqué au modèle exact (7.1) et au méta-modèle choisi. Cela permettrait d'étudier les performances de sélection, notamment en fonction du plan d'expérience D . Lors de cette étude de simulation, nous pourrions également tester d'autres méta-modèles mixtes que celui utilisé dans la section précédente, comme ceux proposés par Barbillon et al. (2017) par exemple, qui s'appuient sur les expressions explicites de la moyenne m_D et de la covariance C_D du

processus gaussien. Ces méta-modèles plus complexes permettent d'inclure l'incertitude sur l'approximation de f . La principale difficulté serait de réussir à conserver le modèle dans la famille exponentielle courbe pour ne pas alourdir l'implémentation de l'algorithme MCMC-SAEM.

- **Gain en temps de calcul :** Cette étude de simulation intensive permettrait également d'étudier le gain en temps de calcul obtenus grâce au couplage avec la métamodélisation. Il est sans doute considérable car, à chaque itération de l'algorithme MCMC-SAEM, nous utilisons uniquement m_D et C_D et non plus des appels à la fonction coûteuse f .
- **Étude théorique :** Par ailleurs, cette proximité entre les supports de β dans le modèle exact (7.1) et dans le méta-modèle choisi est un point essentiel à valider aussi théoriquement. Autrement dit, nous nous demandons si les variables sélectionnées dans le méta-modèle sont pertinentes pour le modèle exact. En effet, comme l'algorithme MCMC-SAEM est appliqué sur un modèle approché plutôt que sur le modèle exact (7.1), nous ne pouvons pas prouver la convergence de l'algorithme vers un maximum local de la distribution a posteriori exacte. Néanmoins, en appliquant les résultats de [Kuhn and Lavielle \(2005\)](#) au modèle approché, nous savons que l'algorithme converge vers un maximum local de la distribution a posteriori du modèle approché. Ainsi, il est nécessaire d'étudier la distance entre les distributions a posteriori du modèle exact et du modèle approché pour conclure à la proximité entre les maximums de ces distributions a posteriori. Sur le même principe, [Barbillon et al. \(2017\)](#) ont démontré que la distance entre les maximums de vraisemblance dans le modèle exact et dans les trois méta-modèles qu'ils proposent est très petite, à condition que le plan d'expérience D soit suffisamment riche. Cependant, dans notre cas, l'étude théorique de la proximité entre les supports serait plus délicate à obtenir étant donné que l'approche SAEMVS identifie le support de β par une étape de seuillage, et non directement par le maximum de la distribution a posteriori. Par contre, l'étude de la distance entre les distributions a posteriori de β pour le modèle exact et pour le méta-modèle pourrait déjà quantifier la qualité d'estimation de cette approche.

Conclusion et perspectives

Cette thèse explore la problématique de la sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Plus précisément, dans le chapitre 5, nous proposons une procédure répondant à cet objectif en combinant l'utilisation d'un prior bayésien de type *spike-and-slab* et l'algorithme d'inférence MCMC-SAEM. L'efficacité de cette procédure est démontrée par une étude de simulations intensive. À notre connaissance, elle constitue ainsi la première méthode disponible dans la littérature qui soit à la fois computationnellement efficace en grande dimension et tout en bénéficiant de la structure de population des modèles à effets mixtes. De plus, pour revenir aux motivations biologiques de ce travail de thèse évoquées dans la chapitre 1, l'utilité de la procédure est mise en évidence à travers l'identification de marqueurs moléculaires impliqués dans le processus de sénescence du blé tendre. Par ailleurs, d'un point de vue théorique, dans le chapitre 6, nous démontrons des taux de contraction de la distribution a posteriori autour des vraies valeurs des paramètres dans le cas d'un prior *spike-and-slab* discret pour un modèle non-linéaire à effets mixtes. Enfin, le chapitre 7 suggère d'associer la méthode développée dans le chapitre 5 à des approches de métamodélisation, pour la rendre utilisable en pratique quand l'évaluation de la fonction de régression est coûteuse en temps de calcul. Ainsi, cette thèse explore à la fois des aspects méthodologiques, algorithmiques, applicatifs mais aussi théoriques.

Dans ce dernier chapitre, nous explorons diverses perspectives de travail, tant sur le plan méthodologique que théorique. La section 8.1 vise à proposer des pistes pour étendre et améliorer la procédure SAEMVS développée dans le chapitre 5, en abordant la quantification de l'incertitude, la sélection conjointe des effets fixes et aléatoires, la sélection par groupes de variables, les modèles hétéroscédastiques et l'ajout d'une structure d'apparement entre les individus. La section 8.2 se concentre sur des perspectives théoriques, notamment la consistance en sélection et l'extension des résultats du chapitre 6 aux modèles non-linéaires à effets mixtes généraux et au prior *spike-and-slab* continu. Ces discussions visent à élargir et approfondir les travaux développés dans cette thèse, ouvrant ainsi de nouvelles voies pour des recherches futures.

Table des matières

8.1 Perspectives méthodologiques	143
8.1.1 Quantification de l'incertitude	143
8.1.2 Sélection des effets aléatoires	143
8.1.3 Sélection par groupe de variables	144
8.1.4 Modèle hétéroscédastique	145
8.1.5 Ajout d'une structure d'apparement	145
8.2 Perspectives théoriques	146
8.2.1 Consistance en sélection	146
8.2.2 Extension au modèle non-linéaire à effets mixtes général	146
8.2.3 Extension au prior <i>spike-and-slab</i> continu	147

8.1 Perspectives méthodologiques

Dans le chapitre 5, nous proposons une procédure de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes, appelée SAEMVS, qui combine l'utilisation d'un prior *spike-and-slab* gaussien et l'algorithme MCMC-SAEM. Parmi les perspectives identifiées dans ce chapitre, nous proposons d'adopter un prior plus flexible pour les variables indicatrices du mélange *spike-and-slab*, afin d'incorporer des informations structurelles a priori sur les covariables, ce qui aiderait à gérer la corrélation entre les covariables. De plus, nous suggérons également de relâcher l'hypothèse de distribution gaussienne des observations conditionnellement aux effets individuels (équation (2.1)) en considérant des distributions plus variées, telles que la loi de Poisson. Dans la suite de cette section, nous élargissons les perspectives de travail.

8.1.1 Quantification de l'incertitude

Une première piste de recherche future consisterait à développer la quantification de l'incertitude pour accompagner la sélection de variables dans SAEMVS. En effet, en raison de contraintes computationnelles liées au contexte de grande dimension, nous ne réalisons pas l'inférence de la distribution a posteriori complète. Nous nous limitons à résoudre le problème d'optimisation qui vise à maximiser cette distribution a posteriori. Par conséquent, nous perdons l'un des principaux avantages des méthodes d'inférence bayésiennes : leur capacité intrinsèque à quantifier l'incertitude. Une première solution pour quantifier l'incertitude de la sélection de variables dans SAEMVS serait d'effectuer des simulations a posteriori MCMC locales dans un voisinage du support sélectionné par SAEMVS pour évaluer l'accumulation relative de la posterior. En reprenant les notations du chapitre 5, cela reviendrait par exemple à simuler la distribution marginale a posteriori de δ sous le prior *spike-and-slab* discret ($\nu_0 = 0$), comme proposé par Ročková and George (2014). Cependant, dans le modèle non-linéaire à effets mixtes considéré dans cette thèse, cette quantité n'est pas explicite et elle est difficile à obtenir numériquement. Par contre, nous disposons déjà d'une information relative à l'incertitude de la sélection dans la procédure SAEMVS mais elle n'est pas pleinement exploitée pour le moment. En effet, le support sélectionné est identifié par une étape de seuillage des estimations des coefficients de β . Pour cette étape, nous calculons la probabilité d'inclusion a posteriori de chaque covariable ℓ étant données les estimations des paramètres $\hat{\Theta}$ pour chaque paramètre individuel m :

$$\mathbb{P}(\delta_{\ell m} = 1 | \mathbf{y}, \hat{\Theta}).$$

On en déduit alors une estimation $\hat{\delta}$ de δ comme suit :

$$\hat{\delta}_{\ell m} = 1 \iff \mathbb{P}(\delta_{\ell m} = 1 | \mathbf{y}, \hat{\Theta}) \geq 0.5.$$

Cependant, la valeur de la probabilité d'inclusion a posteriori de chaque covariable constitue une quantité précieuse qui n'est pas suffisamment exploitée ici, puisque seulement seuillée à 0.5. Elle pourrait fournir des informations sur le niveau de confiance que nous avons quant à l'inclusion de chaque variable.

8.1.2 Sélection des effets aléatoires

Une perspective plus ambitieuse consisterait à inclure la sélection des effets aléatoires dans la méthode SAEMVS, c'est-à-dire l'identification des zéros dans la matrice de covariance des effets

aléatoires Γ . En pratique, la sélection simultanée des effets fixes et aléatoires permettrait en plus d'identifier un effet variétale qui n'est pas lié aux marqueurs mesurés dans le modèle. Il est surtout important de noter que différentes structures de la matrice de covariance des effets aléatoires peuvent conduire à des sélections différentes, parfois même non pertinentes, des covariables pour les effets fixes. Ainsi, envisager une sélection conjointe des effets fixes et aléatoires, comme cela a été proposé dans le cadre linéaire par [Heuclin et al. \(2023\)](#), est particulièrement pertinent. Une solution pourrait consister à modifier le prior utilisé pour la matrice Γ dans SAEMVS afin de pénaliser les effets aléatoires. Plus précisément, en s'inspirant par exemple du prior proposé par [Chen and Dunson \(2003\)](#) dans le cas de la sélection d'effets aléatoires dans le modèle linéaire à effets mixtes en petite dimension, on pourrait décomposer la matrice de covariance Γ selon la décomposition de Cholesky modifiée suivante :

$$\Gamma = \Lambda \Omega \Omega^\top \Lambda,$$

où Λ est une matrice diagonale dont les éléments, notés $\lambda = (\lambda_\ell)_{1 \leq \ell \leq q}$, sont positifs ou nuls proportionnels à l'écart-type des effets aléatoires, et Ω est une matrice triangulaire inférieure avec des 1 sur la diagonale et dont les éléments, notés ω , se rapportent aux corrélations entre les effets aléatoires. Ainsi, si $\lambda_\ell = 0$, alors tous les éléments de la ℓ -ième ligne et de la ℓ -ième colonne de Γ sont nuls, ce qui signifie que l'effet aléatoire ℓ peut être exclu du modèle (sa variance vaut zéro). On choisit alors un prior gaussien sur ω conditionnellement à λ et un prior sur λ consistant en un mélange entre une masse de Dirac en zéro et une densité gaussienne tronquée en dessous de zéro. L'extension de la méthode SAEMVS à cette situation nécessite de revoir le calcul des statistiques exhaustives et sûrement la mise en place d'astuces pour rendre le modèle exponentiel. Sur ce dernier point, je pense notamment qu'il faudra remplacer la masse de Dirac en zéro, proposée par [Chen and Dunson \(2003\)](#), par une distribution continue mais très concentrée en zéro, nécessitant par la suite de mettre en place un critère de seuil pour sélectionner les effets aléatoires.

8.1.3 Sélection par groupe de variables

Une autre piste de recherche intéressante porte sur les performances de sélection de la méthode SAEMVS en présence d'une forte corrélation entre les covariables. En effet, dans le chapitre 5, nous avons montré par une étude de simulation que les performances de sélection de cette méthode se dégradent lorsque la corrélation entre les covariables est trop importante. La procédure SAEMVS est en effet sujette à un changement de sélection parmi les covariables fortement corrélées. Dans de telles situations, il est attendu que toute procédure de sélection de variables soit confrontée à des défis similaires. Cependant, surmonter cette difficulté est particulièrement pertinent pour les données génomiques discutées dans les motivations biologiques de cette thèse (voir chapitre 1), où les marqueurs génétiques adjacents au sein d'un chromosome partagent des informations similaires et sont donc fortement corrélés. L'approche habituelle en biologie consiste plutôt à identifier les régions du génome contenant des marqueurs associés au phénomène étudié. Une perspective intéressante serait donc d'étendre l'approche SAEMVS à une procédure de sélection par groupe de variables. Pour cela, une analyse statistique combinée à une expertise biologique serait nécessaire pour définir les régions du génome en identifiant plusieurs groupes de marqueurs similaires. Ensuite, nous pourrions appliquer un prior *spike-and-slab* sur les groupes de covariables, c'est-à-dire imposer de manière indépendante un prior de mélange de lois gaussi-

ennes multivariées pour chaque groupe de covariables, comme proposé par [Liquet et al. \(2017\)](#) dans le contexte de la régression gaussienne multivariée.

8.1.4 Modèle hétéroscédastique

Par ailleurs, lorsque les données répétées considérées ne sont pas au cours du temps mais dans différentes conditions environnementales, il est souvent supposé, dans les modèles hiérarchiques utilisés dans la littérature de la sélection assistée par marqueurs (voir [Messina et al. \(2018\)](#) par exemple, et les références du chapitre 1), que la variance résiduelle σ^2 dépend de l'environnement :

$$y_{ij} = g(\varphi_i, E_j) + \varepsilon_{ij}, \quad \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma_j^2), \quad 1 \leq i \leq n, \quad 1 \leq j \leq n_i,$$

où E_j représente la condition environnementale de la j -ième observation. La deuxième couche du modèle (2.2) est inchangée :

$$\varphi_i = X_i\beta + \xi_i, \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma).$$

Dans ce cas, les erreurs résiduelles ne sont plus identiquement distribuées mais on les suppose toujours indépendantes. La procédure SAEMVS pourrait aussi être facilement étendue à ce cas hétéroscédastique si $n_1 = n_2 = \dots = n_n = J$. En effet, il suffirait de modifier la statistique exhaustive associée à σ^2 , en une statistique vectorielle comme suit :

$$\left(\sum_{i=1}^n (y_{ij} - g(\varphi_i, E_j))^2 \right)_{1 \leq j \leq J}.$$

Cependant, si le nombre d'observations est différent d'un individu à l'autre, une technique permettant de rendre le modèle exponentiel, comme celle employée dans la section 5.5.2.1 du chapitre 5, sera nécessaire pour étendre la procédure SAEMVS à ce cas.

8.1.5 Ajout d'une structure d'apparentement

Un aspect que nous n'avons pas discuté dans cette thèse est la nécessité, dans certains cas d'application, de considérer une structure de dépendance génétique entre les individus, aussi appelée structure d'apparentement. Plus précisément, nous pourrions considérer le modèle suivant décrit dans [Jaffrézic et al. \(2006\)](#) :

$$\begin{cases} y_{ij} = g(\varphi_i, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \end{cases} \quad (8.1a)$$

$$\begin{cases} \varphi_i = X_i\beta + Z_iu + \xi_i, & u \sim \mathcal{N}_m(0, A \otimes G), \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \end{cases} \quad (8.1b)$$

où $X_i\beta$ modélise les effets fixes, Z_iu les effets génétiques et ξ_i les effets aléatoires environnementaux. La matrice A est connue, de taille $n \times n$, représentant les relations génétiques entre les individus, et G est une matrice inconnue de covariance génétique, de taille $d \times d$, où d est le nombre d'effets génétiques par individu ($m = d \times n$). Par exemple, [Delattre et al. \(2023\)](#) proposent une approche qui combine le modèle non-linéaire à effets mixtes (8.1) et l'algorithme SAEM pour analyser les effets génétiques et environnementaux sur la croissance des plantes. Ainsi, en choisissant un prior Inverse-Wishart pour la matrice G , la méthode SAEMVS doit pouvoir s'étendre relativement facilement à ce cas, permettant de distinguer les effets génétiques des effets environnementaux dans la sélection de variables.

8.2 Perspectives théoriques

Dans cette thèse, nous avons obtenu des taux de contraction de la distribution a posteriori autour des vraies valeurs des paramètres dans un modèle non-linéaire à effets mixtes comme suit :

$$Y_i = f_i(X_i\beta) + Z_i\xi_i + \varepsilon_i, \quad \varepsilon_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_{n_i}(0, \sigma^2 I_{n_i}), \quad \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma), \quad i = 1, \dots, n. \quad (8.2)$$

où σ^2 est supposée connue, $f_i(x) = (f(x, t_{i1}), f(x, t_{i2}), \dots, f(x, t_{in_i}))^\top$, avec t_{ij} le j -ième temps d'observation de l'individu i , et $f : \mathbb{R}^q \times \mathbb{R} \rightarrow \mathbb{R}$ une fonction pouvant être non-linéaire par rapport à sa première composante. Pour démontrer ce résultat, nous avons utilisé un prior *spike-and-slab* discret pour le paramètre de régression β et un prior *Inverse-Wishart* sur la matrice de covariance des effets aléatoires Γ (voir Chapitre 6 pour plus de détails sur ce résultat). Plusieurs perspectives de ce travail sont détaillées ci-après.

8.2.1 Consistance en sélection

Premièrement, bien que la propriété de contraction de la distribution a posteriori soit un critère minimal à vérifier pour valider une méthode d'estimation bayésienne, dans notre contexte de sélection de variables en grande dimension, notre objectif principal est d'identifier les coefficients non-nuls dans β . Plus précisément, et pour se ramener à un critère d'estimation, nous cherchons à estimer correctement le support de β :

$$S_\beta = \left\{ \ell \in \{1, \dots, p\} \mid \beta_\ell \neq 0 \right\}.$$

Ainsi, comme discuté dans la section 3.2.3 du chapitre 3, un résultat de consistance en sélection, établissant que le véritable support S_0 est retrouvé avec une probabilité tendant vers 1, serait plus précis. En ce sens, le travail réalisé dans le chapitre 6 établit les deux premières étapes du schéma classique de preuve de la consistance en sélection sous prior parcimonieux rappelé dans la section 3.2.3.2 du chapitre 3. Une perspective naturelle de ce travail serait donc de poursuivre dans la direction de ce schéma de preuve. Plus précisément, il faudrait établir un résultat d'approximation distributionnelle de la posterior du paramètre de régression pour quantifier l'incertitude via la dispersion de la posterior. Ce résultat permettrait d'établir que la distribution a posteriori ne concentre pas de masse sur des supports contenant strictement le vrai modèle. Enfin, en supposant une condition de type "beta-min" garantissant que les vrais coefficients de β_0 ont une valeur minimale suffisante, la propriété précédente assurerait la consistance en sélection désirée. À mon avis, dans le cas du modèle non-linéaire à effets mixtes (8.2), la principale difficulté résidera dans l'obtention d'une approximation distributionnelle. En effet, la non-linéarité du modèle empêche l'application directe de certaines étapes techniques des preuves classiques. Je pense cependant que les conditions imposées sur la fonction non-linéaire f permettant d'obtenir la contraction seront également suffisantes pour dériver la consistance en sélection.

8.2.2 Extension au modèle non-linéaire à effets mixtes général

Ensuite, comme mentionné dans la section 4.2.2 du chapitre 4, le modèle non-linéaire à effets mixtes, décrit par l'équation 8.2 et dans lequel nous avons démontré des taux de contraction de

la distribution a posteriori, n'est pas équivalent à la version générale du modèle non-linéaire à effets mixtes introduite dans le chapitre 2 par les équations :

$$\begin{cases} y_{ij} = f(\varphi_i, \psi, t_{ij}) + \varepsilon_{ij}, & \varepsilon_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \\ \varphi_i = X_i \beta + \xi_i, & \xi_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}_q(0, \Gamma). \end{cases} \quad (8.3a)$$

$$(8.3b)$$

En effet, la différence majeure entre ces deux modèles (8.2) et (8.3) est que dans le premier l'effet aléatoire est introduit de manière linéaire dans le modèle. Ainsi, contrairement au modèle (8.3), l'espérance et la covariance des observations dans le modèle (8.2) possèdent une expression explicite en fonction des paramètres, ce qui est utile pour conduire la preuve. De plus, il faut également noter que dans le modèle (8.2), nous supposons que les effets fixes ψ sont connus. Une suite naturelle au travail théorique proposé dans cette thèse serait d'adapter les preuves à une version complexifiée du modèle (8.2), pour se rapprocher du modèle non-linéaire à effets mixtes (8.3) dans lequel nous avons développé une méthode de sélection de variables sous un prior *spike-and-slab* (voir le chapitre 5). Par exemple, nous pourrions dans un premier temps rajouter des effets fixes ψ inconnus à estimer. En supposant f Lipschitzienne également par rapport aux effets fixes, je pense que la preuve doit se conduire de la même façon. À mon avis, la plus grande difficulté serait de gérer l'introduction de l'effet aléatoire de manière non-linéaire dans le modèle. En effet, la preuve proposée dans le chapitre 6 repose sur le fait que le modèle (8.2) peut se réécrire simplement sous la forme :

$$Y_i \sim \mathcal{N}(f_i(X_i \beta), \Delta_{\Gamma,i}), \text{ où } \Delta_{\Gamma,i} = Z_i \Gamma Z_i^\top + \sigma^2 I_{n_i},$$

ce qui ne sera pas possible ici. Pour contourner cette difficulté, une idée serait d'utiliser que le modèle (8.2) est en fait une approximation du modèle non-linéaire à effets mixtes général (8.3) (voir section 4.2.2 du chapitre 4). Ainsi, si nous pouvons montrer que les distributions a posteriori dans les modèles (8.2) et (8.3) sont "proches", en un certain sens à définir (potentiellement en distance en variation totale), alors un résultat de contraction de la posterior dans le modèle (8.3) pourrait être dérivé du résultat obtenu dans le chapitre 6 sur le modèle (8.2). Il en est de même pour un résultat de consistance en sélection en utilisant la distribution a posteriori des supports. Ce type de raisonnement est notamment utilisé dans les approches par approximation du modèle mixte par des méta-modèles mixtes en considérant les vraisemblances au lieu des distributions a posteriori (Barbillon et al., 2017).

8.2.3 Extension au prior *spike-and-slab* continu

Enfin, une dernière perspective intéressante serait de démontrer des propriétés asymptotiques fréquentistes dans un modèle non-linéaire à effets mixtes en utilisant un prior *spike-and-slab* continu pour le vecteur de régression β . En effet, dans le chapitre 6, nous avons utilisé un prior *spike-and-slab* discret, c'est-à-dire un mélange entre un Dirac en zéro et une loi continue diffuse. Ce prior *spike-and-slab* discret est considéré comme un idéal théorique puisqu'il assigne réellement la valeur zéro aux coefficients identifiés comme étant en dehors du support de β , permettant la parcimonie dans l'estimation du vecteur de régression. Cependant, comme détaillé dans la section 4.2.1 du chapitre 4, des résultats de contraction de la distribution a posteriori

et de consistance en sélection ont tout de même déjà été démontrés avec des priors *spike-and-slab* continu (Narisetty and He, 2014; Ročková and George, 2018; Jiang and Sun, 2019; Shen and Deshpande, 2022). Toutefois, à notre connaissance, ces résultats concernent seulement des modèles de régression linéaire simple ou multivariée. Une première étape serait d'étendre ces résultats au modèle linéaire à effets mixtes, en ajoutant des effets aléatoires dans le modèle de régression. Pour cela, on pourrait s'appuyer sur les techniques de démonstration proposées par Jeong and Ghosal (2021b) pour la régression linéaire avec paramètres de nuisance et les adapter au cas où la distribution *spike* suit une loi de Laplace, en utilisant les techniques de preuve de Ročková and George (2018).

Bibliographie

- Akaike, H. (1973). Information theory and an extension of the maximum likelihood principle. In *Proc. 2nd Int. Symp. on Information Theory*, pages 267–281.
- Allasonnière, S. and Debavelaere, V. (2021). On the curved exponential family in the Stochastic Approximation Expectation Maximization Algorithm. *ESAIM: Probability and Statistics*, 25:408–432.
- Ayral, G., Si Abdallah, J.-F., Magnard, C., and Chauvin, J. (2021). A novel method based on unbiased correlations tests for covariate selection in nonlinear mixed effects models: The cossac approach. *CPT: pharmacometrics & systems pharmacology*, 10(4):318–329.
- Baey, C., Delattre, M., Kuhn, E., Leger, J.-B., and Lemler, S. (2023). Efficient preconditioned stochastic gradient descent for estimation in latent variable models. In *International Conference on Machine Learning*, pages 1430–1453. PMLR.
- Bai, R., Ročková, V., and George, E. I. (2021). Spike-and-Slab Meets LASSO: A Review of the Spike-and-Slab LASSO. In *Handbook of Bayesian Variable Selection*. Chapman and Hall/CRC, Boca Raton.
- Barbieri, M. M. and Berger, J. O. (2004). Optimal predictive model selection. *The annals of statistics*, 32(3):870–897.
- Barbillon, P., Barthélémy, C., and Samson, A. (2017). Parameter estimation of complex mixed models based on meta-model approach. *Statistics and Computing*, 27(4):1111–1128. Publisher: Springer.
- Bard, Y. (1974). *Nonlinear Parameter Estimation*. Academic Press.
- Belitser, E. and Ghosal, S. (2020). Empirical bayes oracle uncertainty quantification for regression. *The Annals of Statistics*, 48(6):3113–3137.
- Berkey, C. S. and Laird, N. M. (1986). Nonlinear growth curve analysis: estimating the population parameters. *Annals of human biology*, 13(2):111–128.
- Bertrand, J. and Balding, D. J. (2013). Multiple single nucleotide polymorphism analysis using penalized regression in nonlinear mixed-effect pharmacokinetic models. *Pharmacogenetics and genomics*, 23(3):167–174.
- Bertrand, J., De Iorio, M., and Balding, D. J. (2015). Integrating dynamic mixed-effect modelling and penalized regression to explore genetic association with pharmacokinetics. *Pharmacogenetics and genomics*, 25(5):231–238.
- Bickel, P. J., Li, B., Tsybakov, A. B., van de Geer, S. A., Yu, B., Valdés, T., Rivero, C., Fan, J., and van der Vaart, A. (2006). Regularization in statistics. *Test*, 15:271–344.
- Birgé, L. (2006). Model selection via testing: an alternative to (penalized) maximum likelihood estimators. *Annales de l’Institut Henri Poincaré (B) Probability and Statistics*, 42(3):273–325.

- Bühlmann, P. and Van De Geer, S. (2011). *Statistics for high-dimensional data: methods, theory and applications*. Springer Science & Business Media.
- Carvalho, C. M., Polson, N. G., and Scott, J. G. (2009). Handling sparsity via the horseshoe. In *Artificial intelligence and statistics*, volume 5, pages 73–80. PMLR.
- Carvalho, C. M., Polson, N. G., and Scott, J. G. (2010). The horseshoe estimator for sparse signals. *Biometrika*, 97(2):465–480.
- Casella, G., Ghosh, M., Gill, J., and Kyung, M. (2010). Penalized regression, standard errors, and Bayesian lassos. *Bayesian Analysis*, 5(2):369 – 411.
- Castillo, I. (2024). Bayesian nonparametric statistics, St-Flour lecture notes. *arXiv preprint arXiv:2402.16422*.
- Castillo, I., Schmidt-Hieber, J., and Van der Vaart, A. (2015). Bayesian linear regression with sparse priors. *The Annals of Statistics*, 43(5):1986–2018. Publisher: Institute of Mathematical Statistics.
- Castillo, I. and van der Vaart, A. (2012). Needles and straw in a haystack: Posterior concentration for possibly sparse sequences. *The Annals of Statistics*, 40(4):2069–2101.
- Celeux, G. and Diebolt, J. (1986). L’algorithme sem: un algorithme d’apprentissage probabiliste pour la reconnaissance de mélange de densités. *Revue de statistique appliquée*, 34(2):35–52.
- Chae, M., Lin, L., and Dunson, D. B. (2019). Bayesian sparse linear regression with unknown symmetric error. *Information and Inference: A Journal of the IMA*, 8(3):621–653.
- Charmet, G., Tran, L.-G., Auzanneau, J., Rincet, R., and Bouchet, S. (2020). BWGS: a R package for genomic selection and its application to a wheat breeding programme. *PLoS One*, 15(4):e0222733.
- Chatterjee, A., Guedj, J., and Perelson, A. S. (2012). Mathematical modeling of HCV infection: what can it teach us in the era of direct antiviral agents? *Antiviral therapy*, 17(6 0 0):1171.
- Chen, J. and Chen, Z. (2008). Extended Bayesian information criteria for model selection with large model spaces. *Biometrika*, 95(3):759–771.
- Chen, Z. and Dunson, D. B. (2003). Random effects selection in linear mixed models. *Biometrics*, 59(4):762–769.
- Collard, B. C., Jahufer, M., Brouwer, J., and Pang, E. C. K. (2005). An introduction to markers, quantitative trait loci (QTL) mapping and marker-assisted selection for crop improvement: the basic concepts. *Euphytica*, 142:169–196.
- Cooper, M., Technow, F., Messina, C., Ghossein, C., and Totir, L. R. (2016). Use of crop growth models with whole-genome prediction: application to a maize multi-environment trial. *Crop Science*, 56(5):2141–2156.
- Davidian, M. and Gallant, A. R. (1993). The nonlinear mixed effects model with a smooth random effects density. *Biometrika*, 80(3):475–488.

- Davidian, M. and Giltinan, D. M. (1993). Analysis of repeated measurement data using the nonlinear mixed effects model. *Chemometrics and Intelligent laboratory systems*, 20(1):1–24.
- Davidian, M. and Giltinan, D. M. (1995). *Nonlinear Models for Repeated Measurement Data*. Chapman and Hall/CRC Monographs on Statistics and Applied Probability.
- de Valpine, P., Turek, D., Paciorek, C., Anderson-Bergman, C., Temple Lang, D., and Bodik, R. (2017). Programming with models: writing statistical algorithms for general model structures with NIMBLE. *Journal of Computational and Graphical Statistics*, 26:403–417.
- Delattre, M., Lavielle, M., Poursat, M.-A., et al. (2014). A note on BIC in mixed-effects models. *Electronic journal of statistics*, 8(1):456–475.
- Delattre, M. and Poursat, M.-A. (2020). An iterative algorithm for joint covariate and random effect selection in mixed effects models. *The International Journal of Biostatistics*, 16(2):2019–0082.
- Delattre, M., Toda, Y., Tressou, J., and Iwata, H. (2023). Modeling soybean growth: A mixed model approach. *bioRxiv*.
- Delyon, B., Lavielle, M., and Moulines, E. (1999). Convergence of a stochastic approximation version of the EM algorithm. *Annals of statistics*, 27(1):94–128.
- Demidenko, E. (2013). *Mixed models: theory and applications with R*. John Wiley & Sons.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1):1–22.
- Deshpande, S. K., Ročková, V., and George, E. I. (2019). Simultaneous variable and covariance selection with the multivariate spike-and-slab lasso. *Journal of Computational and Graphical Statistics*, 28(4):921–931.
- Diaconis, P. and Freedman, D. (1986). On the consistency of Bayes estimates. *The Annals of Statistics*, 14(1):1–26.
- Diebolt, J. and Ip, E. H. (1996). A stochastic em algorithm for approximating the maximum likelihood estimate. Technical report, Sandia National Lab.(SNL-CA), Livermore, CA (United States).
- Donnet, S. and Samson, A. (2007). Estimation of parameters in incomplete data models defined by dynamical systems. *Journal of Statistical Planning and Inference*, 137(9):2815–2831.
- Fan, J. and Li, R. (2001). Variable selection via nonconcave penalized likelihood and its oracle properties. *Journal of the American statistical Association*, 96(456):1348–1360.
- Fan, J. and Lv, J. (2010). A selective overview of variable selection in high dimensional feature space. *Statistica Sinica*, 20(1):101–148.
- Fan, Y. and Li, R. (2012). Variable selection in linear mixed effects models. *Annals of statistics*, 40(4):2043–2068.

- Fang, G. and Li, P. (2021). On estimation in latent variable models. In *International Conference on Machine Learning*, pages 3100–3110. PMLR.
- Fang, K.-T., Li, R., and Sudjianto, A. (2005). *Design and modeling for computer experiments*. CRC press.
- Fisher, R. A. (1922). On the mathematical foundations of theoretical statistics. *Philosophical transactions of the Royal Society of London. Series A, containing papers of a mathematical or physical character*, 222(594-604):309–368.
- Fort, G. and Moulines, E. (2003). Convergence of the Monte Carlo expectation maximization for curved exponential families. *The Annals of Statistics*, 31(4):1220–1259.
- Friedman, J., Hastie, T., and Tibshirani, R. (2010). Regularization paths for generalized linear models via coordinate descent. *Journal of statistical software*, 33(1):1–22.
- Frühwirth-Schnatter, S. and Tüchler, R. (2008). Bayesian parsimonious covariance estimation for hierarchical linear mixed models. *Statistics and Computing*, 18:1–13.
- George, E. I. and McCulloch, R. E. (1993). Variable selection via Gibbs sampling. *Journal of the American Statistical Association*, 88(423):881–889.
- George, E. I. and McCulloch, R. E. (1997). Approaches for Bayesian variable selection. *Statistica sinica*, 7(2):339–373.
- Geweke, J. (1996). Variable selection and model comparison in regression. In *Bayesian Statistics 5*.
- Ghosal, S. (1999). Asymptotic normality of posterior distributions in high-dimensional linear models. *Bernoulli*, 5(2):315–331.
- Ghosal, S., Ghosh, J. K., and Van Der Vaart, A. W. (2000). Convergence rates of posterior distributions. *Annals of Statistics*, 28(2):500–531. Publisher: JSTOR.
- Ghosal, S. and Van Der Vaart, A. (2007). Convergence rates of posterior distributions for noniid observations. *Annals of Statistics*, 35(1):192–223.
- Ghosal, S. and Van der Vaart, A. (2017). *Fundamentals of nonparametric Bayesian inference*, volume 44. Cambridge University Press.
- Giraud, C. (2021). *Introduction to high-dimensional statistics*. Chapman and Hall/CRC.
- Green, P. J. (1990). On use of the EM algorithm for penalized likelihood estimation. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 52(3):443–452.
- Grenier, E., Louvet, V., and Vigneaux, P. (2014). Parameter estimation in non-linear mixed effects models with SAEM algorithm: extension from ODE to PDE. *ESAIM: Mathematical Modelling and Numerical Analysis*, 48(5):1303–1329.
- Gu, M. G. and Li, S. (1998). A stochastic approximation algorithm for maximum-likelihood estimation with incomplete data. *Canadian Journal of Statistics*, 26(4):567–582.

- Hammer, G., Kropff, M., Sinclair, T., and Porter, J. (2002). Future contributions of crop modelling from heuristics and supporting decision making to understanding genetic regulation and aiding crop improvement. *European Journal of Agronomy*, 18(1-2):15–31.
- Hans, C., Dobra, A., and West, M. (2007). Shotgun stochastic search for “large p” regression. *Journal of the American Statistical Association*, 102(478):507–516.
- Hastie, T., Tibshirani, R., Friedman, J. H., and Friedman, J. H. (2009). *The elements of statistical learning: data mining, inference, and prediction*, volume 2. Springer, Stanford.
- Hastie, T., Tibshirani, R., and Wainwright, M. (2015). Statistical learning with sparsity. *Mono-graphs on statistics and applied probability*, 143(143):8.
- Hastings, W. K. (1970). Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, 57(1):97–109.
- Heslot, N., Akdemir, D., Sorrells, M. E., and Jannink, J.-L. (2014). Integrating environmental covariates and crop modeling into the genomic selection framework to predict genotype by environment interactions. *Theoretical and applied genetics*, 127:463–480.
- Heuclin, B., Denis, M., Trottier, C., Tisné, S., and Mortier, F. (2023). Continuous shrinkage priors for fixed and random effects selection in linear mixed models: application to genetic mapping. *preprint hal-04238536*.
- Heuclin, B., Mortier, F., Trottier, C., and Denis, M. (2020). Bayesian varying coefficient model with selection: An application to functional mapping. *Journal of the Royal Statistical Society: Series C Applied Statistics*, 70(1):24–50.
- Hoerl, A. E. and Kennard, R. W. (1970). Ridge regression: applications to nonorthogonal problems. *Technometrics*, 12(1):69–82.
- International Wheat Genome Sequencing Consortium, . (2018). Shifting the limits in wheat research and breeding using a fully annotated reference genome. *Science*, 361(6403):eaar7191.
- Jaffrézic, F., Meza, C., Lavielle, M., and Foulley, J.-L. (2006). Genetic analysis of growth curves using the SAEM algorithm. *Genetics Selection Evolution*, 38(6):583–600.
- Jamieson, P., Semenov, M., Brooking, I., and Francis, G. (1998). Sirius: a mechanistic model of wheat response to environmental variation. *European Journal of Agronomy*, 8(3-4):161–179.
- Jannink, J.-L., Bink, M. C., and Jansen, R. C. (2001). Using complex plant pedigrees to map valuable genes. *Trends in plant science*, 6(8):337–342.
- Jeong, S. and Ghosal, S. (2021a). Posterior contraction in sparse generalized linear models. *Biometrika*, 108(2):367–379. Publisher: Oxford University Press.
- Jeong, S. and Ghosal, S. (2021b). Unified Bayesian theory of sparse linear regression with nuisance parameters. *Electronic Journal of Statistics*, 15(1):3040–3111. Publisher: The Institute of Mathematical Statistics and the Bernoulli Society.

- Jiang, B. and Sun, Q. (2019). Bayesian high-dimensional linear regression with generic spike-and-slab priors. *arXiv preprint arXiv:1912.08993*.
- Jiang, W. (2007). Bayesian variable selection for high dimensional generalized linear models: convergence rates of the fitted densities. *Annals of Statistics*, 35(4):1487–1511.
- Jonsson, E. N. and Karlsson, M. O. (1998). Automated covariate model building within NON-MEM. *Pharmaceutical research*, 15(9):1463–1468.
- Kang, H. M., Zaitlen, N. A., Wade, C. M., Kirby, A., Heckerman, D., Daly, M. J., and Eskin, E. (2008). Efficient control of population structure in model organism association mapping. *Genetics*, 178(3):1709–1723.
- Krige, D. G. (1951). A statistical approach to some basic mine valuation problems on the witwatersrand. *Journal of the Southern African Institute of Mining and Metallurgy*, 52(6):119–139.
- Kuhn, E. and Lavielle, M. (2004). Coupling a stochastic approximation version of EM with an MCMC procedure. *ESAIM: Probability and Statistics*, 8:115–131.
- Kuhn, E. and Lavielle, M. (2005). Maximum likelihood estimation in nonlinear mixed effects models. *Computational statistics & data analysis*, 49(4):1020–1038.
- Laird, N. M. and Ware, J. H. (1982). Random-effects models for longitudinal data. *Biometrics*, 38(4):963–974.
- Lavielle, M. (2014). *Mixed effects models for the population approach: models, tasks, methods and tools*. CRC press, New York.
- Le Cam, L. (2012). *Asymptotic methods in statistical decision theory*. Springer Science & Business Media.
- Lee, S. Y. (2022). Bayesian Nonlinear Models for Repeated Measurement Data: An Overview, Implementation, and Applications. *Mathematics*, 10(6):898.
- Li, Y., Wang, S., Song, P. X.-K., Wang, N., Zhou, L., and Zhu, J. (2018). Doubly regularized estimation and selection in linear mixed-effects models for high-dimensional longitudinal data. *Statistics and its Interface*, 11(4):721.
- Lindstrom, M. J. and Bates, D. M. (1990). Nonlinear mixed effects models for repeated measures data. *Biometrics*, 46(3):673–687.
- Liquet, B., Mengersen, K., Pettitt, A. N., and Sutton, M. (2017). Bayesian variable selection regression of multivariate responses for group data. *Bayesian Analysis*, 12(4):1039–1067.
- Louis, T. A. (1982). Finding the observed information matrix when using the em algorithm. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 44(2):226–233.
- Malosetti, M., Van der Linden, C., Vosman, B., and Van Eeuwijk, F. (2007). A mixed-model approach to association mapping using pedigree information with an illustration of resistance to phytophthora infestans in potato. *Genetics*, 175(2):879–889.

- Malsiner-Walli, G. and Wagner, H. (2018). Comparing spike and slab priors for Bayesian variable selection. *Austrian Journal of Statistics*, 40(4):241–264.
- Martin, R., Mess, R., and Walker, S. G. (2017). Empirical Bayes posterior concentration in sparse high-dimensional linear models. *Bernoulli*, 23(3):1822 – 1847.
- Matheron, G. (1963). Principles of geostatistics. *Economic geology*, 58(8):1246–1266.
- McCulloch, C. E. (1997). Maximum likelihood algorithms for generalized linear mixed models. *Journal of the American statistical Association*, 92(437):162–170.
- Messina, C. D., Technow, F., Tang, T., Totir, R., Gho, C., and Cooper, M. (2018). Leveraging biological insight and environmental variation to improve phenotypic prediction: Integrating crop growth models (CGM) with whole genome prediction (WGP). *European Journal of Agronomy*, 100:151–162.
- Meuwissen, T. H., Hayes, B. J., and Goddard, M. (2001). Prediction of total genetic value using genome-wide dense marker maps. *Genetics*, 157(4):1819–1829.
- Mitchell, T. J. and Beauchamp, J. J. (1988). Bayesian variable selection in linear regression. *Journal of the american statistical association*, 83(404):1023–1032.
- Moreau, L., Charcosset, A., Hospital, F., and Gallais, A. (1998). Marker-assisted selection efficiency in populations of finite size. *Genetics*, 148(3):1353–1365.
- Narisetty, N. N. and He, X. (2014). Bayesian variable selection with shrinking and diffusing priors. *The Annals of Statistics*, 42(2):789–817. Publisher: Institute of Mathematical Statistics.
- Ning, B., Jeong, S., and Ghosal, S. (2020). Bayesian linear regression for multivariate responses under group sparsity. *Bernoulli*, 26(3):2353–2382. Publisher: International Statistical Institute.
- Ollier, E. (2022). Fast selection of nonlinear mixed effect models using penalized likelihood. *Computational Statistics & Data Analysis*, 167:107373.
- Ollier, E., Samson, A., Delavenne, X., and Viallon, V. (2016). A SAEM algorithm for fused lasso penalized nonlinear mixed effect models: application to group comparison in pharmacokinetics. *Computational statistics & data analysis*, 95:207–221.
- Onogi, A., Watanabe, M., Mochizuki, T., Hayashi, T., Nakagawa, H., Hasegawa, T., and Iwata, H. (2016). Toward integration of genomic selection with crop modelling: the development of an integrated approach to predicting rice heading dates. *Theoretical and Applied Genetics*, 129:805–817.
- Park, T. and Casella, G. (2008). The Bayesian lasso. *Journal of the American Statistical Association*, 103(482):681–686.
- Patterson, N., Price, A. L., and Reich, D. (2006). Population structure and eigenanalysis. *PLoS genetics*, 2(12):e190.

- Piepho, H.-P. (1997). Analyzing genotype-environment data by mixed models with multiplicative terms. *Biometrics*, 53(2):761–766.
- Pinheiro, J. C. and Bates, D. M. (1995). Approximations to the log-likelihood function in the nonlinear mixed-effects model. *Journal of computational and Graphical Statistics*, 4(1):12–35.
- Pinheiro, J. C. and Bates, D. M. (2000). *Mixed-effects Models in S and S-PLUS*. Springer, New York.
- Pocock, S. J., Cook, D. G., and Beresford, S. A. (1981). Regression of area mortality rates on explanatory variables: What weighting is appropriate? *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 30(3):286–295.
- Prague, M. and Lavielle, M. (2022). Samba: A novel method for fast automatic model building in nonlinear mixed-effects models. *CPT: Pharmacometrics & Systems Pharmacology*, 11(2):161–172.
- Purcell, S., Neale, B., Todd-Brown, K., Thomas, L., Ferreira, M. A., Bender, D., Maller, J., Sklar, P., De Bakker, P. I., Daly, M. J., et al. (2007). PLINK: a tool set for whole-genome association and population-based linkage analyses. *The American journal of human genetics*, 81(3):559–575.
- Ray, K. and Szabó, B. (2022). Variational bayes for high-dimensional linear regression with sparse priors. *Journal of the American Statistical Association*, 117(539):1270–1281.
- Reymond, M., Muller, B., Leonardi, A., Charcosset, A., and Tardieu, F. (2003). Combining quantitative trait loci analysis and an ecophysiological model to analyze the genetic variability of the responses of maize leaf growth to temperature and water deficit. *Plant physiology*, 131(2):664–675.
- Ribba, B., Kaloshi, G., Peyre, M., Ricard, D., Calvez, V., Tod, M., Čajavec-Bernard, B., Idbaih, A., Psimaras, D., Dainese, L., et al. (2012). A tumor growth inhibition model for low-grade glioma treated with chemotherapy or radiotherapy. *Clinical Cancer Research*, 18(18):5071–5080.
- Rimbert, H., Darrier, B., Navarro, J., Kitt, J., Choulet, F., Leveugle, M., Duarte, J., Rivière, N., Eversole, K., Consortium, I. W. G. S., et al. (2018). High throughput SNP discovery and genotyping in hexaploid wheat. *PloS one*, 13(1):e0186329.
- Rincent, R., Charpentier, J.-P., Faivre-Rampant, P., Paux, E., Le Gouis, J., Bastien, C., and Segura, V. (2018). Phenomic selection is a low-cost and high-throughput method based on indirect predictions: proof of concept on wheat and poplar. *G3: Genes, Genomes, Genetics*, 8(12):3961–3972.
- Rincent, R., Kuhn, E., Monod, H., Oury, F.-X., Rousset, M., Allard, V., and Le Gouis, J. (2017). Optimization of multi-environment trials for genomic selection based on crop models. *Theoretical and Applied Genetics*, 130:1735–1752.

- Rincent, R., Malosetti, M., Ababaei, B., Touzy, G., Mini, A., Bogard, M., Martre, P., Le Gouis, J., and Van Eeuwijk, F. (2019). Using crop growth model stress covariates and AMMI decomposition to better predict genotype-by-environment interactions. *Theoretical and Applied Genetics*, 132(12):3399–3411.
- Robert, C. (2006). *Le choix bayésien: Principes et pratique*. Springer Science & Business Media.
- Ročková, V. and George, E. I. (2014). EMVS: The EM approach to Bayesian variable selection. *Journal of the American Statistical Association*, 109(506):828–846.
- Rousseau, J. (2016). On the frequentist properties of bayesian nonparametric methods. *Annual Review of Statistics and Its Application*, 3:211–231.
- Roustant, O., Ginsbourger, D., and Deville, Y. (2012). DiceKriging, DiceOptim: Two R Packages for the Analysis of Computer Experiments by Kriging-Based Metamodeling and Optimization. *Journal of Statistical Software*, 51:1–55.
- Ročková, V. and George, E. I. (2018). The spike-and-slab lasso. *Journal of the American Statistical Association*, 113(521):431–444. Publisher: Taylor & Francis.
- Sacks, J., Schiller, S. B., and Welch, W. J. (1989). Designs for computer experiments. *Technometrics*, 31(1):41–47.
- Schelldorfer, J., Bühlmann, P., and de Geer, S. v. (2011). Estimation for high-dimensional linear mixed-effects models using l1-penalization. *Scandinavian Journal of Statistics*, 38(2):197–214.
- Schnabel, R. B., Koonatz, J. E., and Weiss, B. E. (1985). A modular system of algorithms for unconstrained minimization. *ACM Transactions on Mathematical Software (TOMS)*, 11(4):419–440.
- Schulz-Streeck, T., Ogutu, J. O., Gordillo, A., Karaman, Z., Knaak, C., and Piepho, H.-P. (2013). Genomic selection allowing for marker-by-environment interaction. *Plant Breeding*, 132(6):532–538.
- Schwarz, G. (1978). Estimating the dimension of a model. *The annals of statistics*, 6(2):461–464.
- Shaby, B. and Wells, M. T. (2010). Exploring an adaptive Metropolis algorithm. Technical report, Department of Statistical Science, Duke University.
- Sheiner, L. B. and Beal, S. L. (1980). Evaluation of methods for estimating population pharmacokinetic parameters. i. michaelis-menten model: routine clinical pharmacokinetic data. *Journal of pharmacokinetics and biopharmaceutics*, 8(6):553–571.
- Sheiner, L. B. and Beal, S. L. (1981). Evaluation of methods for estimating population pharmacokinetic parameters ii. biexponential model and experimental pharmacokinetic data. *Journal of pharmacokinetics and biopharmaceutics*, 9(5):635–651.
- Sheiner, L. B. and Beal, S. L. (1983). Evaluation of methods for estimating population pharmacokinetic parameters. iii. monoexponential model: routine clinical pharmacokinetic data. *Journal of pharmacokinetics and biopharmaceutics*, 11(3):303–319.

- Sheiner, L. B., Rosenberg, B., and Marathe, V. V. (1977). Estimation of population characteristics of pharmacokinetic parameters from routine clinical data. *Journal of pharmacokinetics and biopharmaceutics*, 5:445–479.
- Sheiner, L. B., Rosenberg, B., and Melmon, K. L. (1972). Modelling of individual pharmacokinetics for computer-aided drug dosage. *Computers and Biomedical Research*, 5(5):441–459.
- Shen, Y. and Deshpande, S. K. (2022). On the posterior contraction of the multivariate spike-and-slab LASSO. *arXiv preprint arXiv:2209.04389*.
- Smith, A., Cullis, B. R., and Thompson, R. (2005). The analysis of crop cultivar breeding and evaluation trials: an overview of current mixed model approaches. *The Journal of Agricultural Science*, 143(6):449–462.
- Song, Q. and Liang, F. (2023). Nearly optimal Bayesian shrinkage for high-dimensional regression. *Science China Mathematics*, 66(2):409–442.
- Spiegelhalter, D. J., Best, N. G., Carlin, B. P., and Van Der Linde, A. (2002). Bayesian measures of model complexity and fit. *Journal of the royal statistical society: Series B (statistical methodology)*, 64(4):583–639.
- Steimer, J.-L., Mallet, A., Golmard, J.-L., and Boisvieux, J.-F. (1984). Alternative approaches to estimation of population pharmacokinetic parameters: comparison with the nonlinear mixed-effect model. *Drug metabolism reviews*, 15(1-2):265–292.
- Stingo, F. C., Chen, Y. A., Vannucci, M., Barrier, M., and Mirkes, P. E. (2010). A Bayesian graphical modeling approach to microRNA regulatory network inference. *The annals of applied statistics*, 4(4):2024.
- Stingo, F. C. and Vannucci, M. (2011). Variable selection for discriminant analysis with Markov random field priors for the analysis of microarray data. *Bioinformatics*, 27(4):495–501.
- Stratelli, R., Laird, N., and Ware, J. H. (1984). Random-effects models for serial observations with binary response. *Biometrics*, 40(4):961–971.
- Tadesse, M. G. and Vannucci, M. (2021). *Handbook of Bayesian variable selection*. CRC Press, Boca Raton.
- Tang, Y. and Martin, R. (2023). Empirical Bayes inference in sparse high-dimensional generalized linear models. *arXiv preprint arXiv:2303.07854*.
- Technow, F., Messina, C. D., Totir, L. R., and Cooper, M. (2015). Integrating crop growth models with whole genome prediction through approximate bayesian computation. *PloS one*, 10(6):e0130855.
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society: Series B (Methodological)*, 58(1):267–288.
- Tibshirani, R., Saunders, M., Rosset, S., Zhu, J., and Knight, K. (2005). Sparsity and smoothness via the fused lasso. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 67(1):91–108.

- Touzy, G., Rincent, R., Bogard, M., Lafarge, S., Dubreuil, P., Mini, A., Deswarte, J.-C., Beauchêne, K., Le Gouis, J., and Praud, S. (2019). Using environmental clustering to identify specific drought tolerance QTLs in bread wheat (*t. aestivum* l.). *Theoretical and Applied Genetics*, 132(10):2859–2880.
- Van Berloo, R. and Stam, P. (1999). Comparison between marker-assisted selection and phenotypical selection in a set of *Arabidopsis thaliana* recombinant inbred lines. *Theoretical and Applied Genetics*, 98:113–118.
- Verbeke, G. and Molenberghs, G. (2000). *Linear mixed models for longitudinal data*. Springer, New York.
- Verbeke, J. and Cools, R. (1995). The newton-raphson method. *International Journal of Mathematical Education in Science and Technology*, 26(2):177–193.
- Vonesh, E. F. (1996). A note on the use of Laplace’s approximation for nonlinear mixed-effects models. *Biometrika*, 83(2):447–452.
- Vonesh, E. F. and Carter, R. L. (1992). Mixed-effects nonlinear regression for unbalanced repeated measures. *Biometrics*, 48(1):1–17.
- Vos, J., Marcelis, L., and Evers, J. (2007). Functional-structural plant modelling in crop production: adding a dimension. *Wageningen UR Frontis Series*, 22:1–12.
- Wakefield, J., Smith, A., Racine-Poon, A., and Gelfand, A. E. (1994). Bayesian analysis of linear and non-linear population models by using the gibbs sampler. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 43(1):201–221.
- Wei, G. C. and Tanner, M. A. (1990). A Monte Carlo implementation of the EM algorithm and the poor man’s data augmentation algorithms. *Journal of the American statistical Association*, 85(411):699–704.
- Weir, A. H., Bragg, P., Porter, J., and Rayner, J. (1984). A winter wheat crop simulation model without water or nutrient limitations. *The Journal of Agricultural Science*, 102(2):371–382.
- West, B. T., Welch, K. B., and Galecki, A. T. (2006). *Linear mixed models: a practical guide using statistical software*. Chapman and Hall/CRC.
- Wolfinger, R. (1993). Laplace’s approximation for nonlinear mixed models. *Biometrika*, 80(4):791–795.
- Wu, C. J. (1983). On the convergence properties of the EM algorithm. *The Annals of statistics*, 11(1):95–103.
- Yeap, B. Y. and Davidian, M. (2001). Robust two-stage estimation in hierarchical nonlinear models. *Biometrics*, 57(1):266–272.
- Yin, X., Struik, P. C., and Kropff, M. J. (2004). Role of crop physiology in predicting gene-to-phenotype relationships. *Trends in plant science*, 9(9):426–432.

-
- Yu, J., Pressoir, G., Briggs, W. H., Vroh Bi, I., Yamasaki, M., Doebley, J. F., McMullen, M. D., Gaut, B. S., Nielsen, D. M., Holland, J. B., et al. (2006). A unified mixed-model method for association mapping that accounts for multiple levels of relatedness. *Nature genetics*, 38(2):203–208.
- Yuan, M. and Lin, Y. (2006). Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 68(1):49–67.
- Zhao, K., Aranzana, M. J., Kim, S., Lister, C., Shindo, C., Tang, C., Toomajian, C., Zheng, H., Dean, C., Marjoram, P., et al. (2007). An Arabidopsis example of association mapping in structured samples. *PLoS genetics*, 3(1):e4.
- Zou, H. (2006). The adaptive lasso and its oracle properties. *Journal of the American statistical association*, 101(476):1418–1429.
- Zou, H. and Hastie, T. (2005). Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 67(2):301–320.

Titre: Procédures de sélection de variables en grande dimension dans les modèles non-linéaires à effets mixtes. Application en amélioration des plantes.

Mots clés: Modèles non-linéaires à effets mixtes, Données de grande dimension, Sélection de variables bayésienne, Algorithme SAEM, Priors *spike-and-slab*, Contraction a posteriori.

Résumé: Les modèles à effets mixtes analysent des observations collectées de façon répétée sur plusieurs individus, attribuant la variabilité à différentes sources (intra-individuelle, inter-individuelle, résiduelle). Prendre en compte cette variabilité est essentiel pour caractériser sans biais les mécanismes biologiques sous-jacents. Ces modèles utilisent des covariables et des effets aléatoires pour décrire la variabilité entre individus : les covariables décrivent les différences dues à des caractéristiques observées, tandis que les effets aléatoires représentent la variabilité non attribuable aux covariables mesurées. Dans un contexte de grande dimension, où le nombre de covariables dépasse celui des individus, identifier les covariables influentes est difficile, car la sélection porte sur des variables latentes du modèle. De nombreuses procédures ont été mises au point pour les modèles linéaires à effets mixtes, mais les contributions pour les modèles non-linéaires sont rares et manquent de fondements théoriques. Cette thèse vise à développer une procédure de sélection de covariables en grande dimen-

sion pour les modèles non-linéaires à effets mixtes, en étudiant leurs implémentations pratiques et leurs propriétés théoriques. Cette procédure est basée sur l'utilisation d'un prior *spike-and-slab* gaussien et de l'algorithme SAEM (*Stochastic Approximation of Expectation Maximisation Algorithm*). Des taux de contraction a posteriori autour des vraies valeurs des paramètres dans un modèle non-linéaire à effets mixtes sous prior *spike-and-slab* discret ont été obtenus, comparables à ceux observés dans des modèles linéaires. Les travaux conduits dans cette thèse sont motivés par des questions appliquées en amélioration des plantes, où ces modèles décrivent le développement des plantes en fonction de leurs génotypes et des conditions environnementales. Les covariables considérées sont généralement nombreuses puisque les variétés sont caractérisées par des milliers de marqueurs génétiques, dont la plupart n'ont aucun effet sur certains traits phénotypiques. La méthode statistique développée dans la thèse est appliquée à un jeu de données réel relatif à cette application.

Title: High-dimensional variable selection procedures in nonlinear mixed effects models. Application in plant breeding.

Keywords: Nonlinear mixed effects models, High dimensional data, Bayesian variable selection, SAEM Algorithm, Spike-and-slab priors, Posterior contraction.

Abstract: Mixed-effects models analyze observations collected repeatedly from several individuals, attributing variability to different sources (intra-individual, inter-individual, residual). Accounting for this variability is essential to characterize the underlying biological mechanisms without bias. These models use covariates and random effects to describe variability among individuals: covariates explain differences due to observed characteristics, while random effects represent the variability not attributable to measured covariates. In high-dimensional context, where the number of covariates exceeds the number of individuals, identifying influential covariates is challenging, as selection focuses on latent variables in the model. Many procedures have been developed for linear mixed-effects models, but contributions for non-linear models are rare and lack theoretical foundations. This thesis aims to develop a high-dimensional covariate selection proce-

duce for non-linear mixed-effects models by studying their practical implementations and theoretical properties. This procedure is based on the use of a gaussian spike-and-slab prior and the SAEM algorithm (*Stochastic Approximation of Expectation Maximisation Algorithm*). Posterior contraction rates around true parameter values in a non-linear mixed-effects model under a discrete spike-and-slab prior have been obtained, comparable to those observed in linear models. The work in this thesis is motivated by practical questions in plant breeding, where these models describe plant development as a function of their genotypes and environmental conditions. The considered covariates are generally numerous since varieties are characterized by thousands of genetic markers, most of which have no effect on certain phenotypic traits. The statistical method developed in the thesis is applied to a real dataset related to this application.