



Indices de Sobol généralisés pour variables dépendantes

Gaëlle Chastaing

► To cite this version:

Gaëlle Chastaing. Indices de Sobol généralisés pour variables dépendantes. Méthodologie [stat.ME]. Université de Grenoble, 2013. Français. NNT : . tel-00930229

HAL Id: tel-00930229

<https://theses.hal.science/tel-00930229>

Submitted on 17 Jan 2014

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.

THÈSE

Pour obtenir le grade de

DOCTEUR DE L'UNIVERSITÉ DE GRENOBLE

Spécialité : **Mathématiques appliquées**

Arrêté ministériel : du 7 août 2006

Présentée par

Gaelle Chastaing

Thèse dirigée par **Fabrice Gamboa** et **Clémentine Prieur**

préparée au sein du **Laboratoire Jean Kuntzmann**
et de l'**école doctorale Mathématiques, Sciences et Technologies de l'Information, Informatique**

Indices de Sobol généralisés pour variables dépendantes

Thèse soutenue publiquement le 23 Septembre 2013 à 14h
devant le jury composé de :

Hervé Monod

Directeur de recherche, INRA, Rapporteur

Sergei Kucherenko

Professeur, Imperial College London, Rapporteur

Éric Blayo

Professeur, Université Joseph Fourier/LJK, Examineur

François Coquet

Professeur, ENSAI, Examineur

Bertrand Iooss

Chercheur senior, EDF-R&D, Examineur

Fabrice Gamboa

Professeur, Université Paul Sabatier/IMT, Directeur de thèse

Clémentine Prieur

Professeur, Université Joseph Fourier/LJK, Directrice de thèse



À Armand Jean-Baptiste

Remerciements

Je remercie Clémentine pour sa rigueur et ses lectures minutieuses. Ses remarques pertinentes ont contribué à approfondir de nombreuses idées. Un grand merci à Fabrice, dont la chaleur humaine n'a d'égale que les compétences mathématiques. Ses nombreuses idées m'ont permises d'élargir mes connaissances en Statistique, et m'ont aidé à définir l'orientation de ma thèse, mais aussi mes préférences pour de futures orientations professionnelles. Je n'oublierai jamais ses sessions "Maths en course", ou encore ses recettes "café-pistoche-casse dalle-et après on bosse" dont lui seul a le secret. J'espère avoir l'occasion de le croiser entre Paris et Saïgon, avec escale au marathon du Médoc !

Je remercie aussi mes deux rapporteurs Sergei Kucherenko et Hervé Monod pour leur relecture attentive et leurs remarques pertinentes sur le manuscrit. Merci aussi à François Coquet, qui me suit depuis maintenant cinq ans, Bertrand Iooss, qui sait habilement concilier Sciences et animation pendant les écoles d'été. Enfin, merci aussi à Eric Blayo, chef d'équipe MOISE qui m'a accueilli et épaulé durant ces trois années.

Je ne saurai jamais assez remercier Anne de m'avoir libéré de très nombreuses et laborieuses procédures administratives, qu'elle traite avec une efficacité quasi inégalable. Je la remercie également pour son soutien et sa générosité. Merci à toute l'équipe MOISE, et particulièrement à mes cobureaux Federico, Bertrand et Vincent, qui ont su me divertir avec leurs batailles, vannes et taquineries en tout genre, sans oublier les soirées bières ! Je n'oublie pas non plus mes grands frères de thèse, Jean Yves, pour ses randonnées montagne et sa grande richesse humaine, Alexandre pour ses sessions flims et ses références inclassables et Thibault, pour son incroyable dévouement à la cause humaine.

Une standing ovation pour mes collègues de l'IMT ; Benoît, pour ses belles idées et nos débats parfois animés, Mélanie, pour son incroyable gentillesse et sa poésie, et Malika, pour sa couleur et nos sorties épicuriennes. J'ai eu l'extraordinaire chance de partager leur bureau, mais surtout des moments inoubliables. Le bureau 206 restera toujours le meilleur bureau à mes yeux !

Merci aussi à Claire C. (surtout pour sa façon bien à elle de repousser les englishmen dans le grand canyon !), Sébastien, Claire B., Nil, Laurent, Pierre, Julie, Chloé et Tatiana avec qui j'ai eu beaucoup de plaisir à discuter, et que je ne doute pas revoir bientôt.

Merci aussi à Sébastien Gadat et à Magali, avec qui j'ai eu beaucoup de plaisir à travailler, et à échanger.

Merci à mes parents et à mes soeurs et beaux frères de m'avoir soutenu en toutes circonstances, et d'avoir fait de moi ce que je suis aujourd'hui. Je leur dois cet accomplissement et tous les suivants. Plus que tout autre, ils méritent mes remerciements et ma reconnaissance. Je n'oublie pas non plus mes tantes et oncles, mes cousins/cousines et ma grand-mère, qui ont aussi contribué à mon succès pour le soutien et la confiance qu'ils m'ont attribué. Une autre partie de la "famille" sont mes chers amis de longue date, Laurence, sans qui cette dernière année n'aurait jamais aussi bien débuté (merci aussi à Anthony de m'avoir accueilli), Julie, pour nos inlassables discussions sur les terrasses de café, Fabien, mon ami d'école, de soirées et de rigolade, mais qui s'installe à Toulouse quand je m'en vais ! Merci aussi à Lionel, avec qui j'ai plaisir à échanger sur les voyages, et Babe, un ami de très longue date que j'ai toujours plaisir à voir. Merci aussi à Géraud et Sylvain, l'un pour ses soirées chansons et son amour du terroir, l'autre pour son soutien et les nombreuses escapades dans lesquelles il m'a amené.

Enfin, j'attribue l'essentiel de mes remerciements à Loïc, qui a su m'épauler pendant ces deux dernières années, et sans qui je n'aurais certainement pas écrit ces lignes. Il n'y a pas de mots assez fort pour le gratifier de la confiance et du soutien qu'il a su me donner, et qu'il continue à m'attribuer. J'espère partager encore longtemps avec lui les nombreuses aventures que la vie nous réserve.

Table des matières

Notations	13
Introduction	17
Cadre d'étude	25
I Etat de l'art	27
1 Analyse de sensibilité globale et indépendance	29
1.1 Méthodes basées sur la variance : un historique	30
1.1.1 Mesures d'importance	30
1.1.2 L'indice de Sobol	30
1.2 Décomposition de Hoeffding et indices de Sobol	33
1.3 Estimation des indices de Sobol	36
1.3.1 Méthodes de Monte Carlo	37
1.3.2 Méthodes FAST et ses variantes	39
1.3.3 Polynômes de chaos	42
2 Dépendance en Statistique	49
2.1 Théorie des copules	49
2.1.1 Définitions et propriétés	50
2.1.2 Exemples de copules	53
2.2 Détection et mesure de la dépendance en pratique	60
2.3 Techniques d'échantillonnage de variables dépendantes	66
2.3.1 Méthode basée sur la corrélation de rangs	66
2.3.2 Méthode d'échantillonnage des copules	68
2.4 Choix de modèles	70
2.4.1 Estimation de θ	71
2.4.2 évaluation de la dépendance	72

3	Analyse de sensibilité pour variables d'entrée dépendantes	75
3.1	Techniques d'échantillonnage	77
3.1.1	Technique de Monte Carlo	78
3.1.2	Technique d'échantillonnage FAST	79
3.1.3	Estimation par polynômes locaux	80
3.1.4	Procédé d'orthogonalisation des facteurs d'entrée	81
3.2	Mesures basées sur la distribution	82
3.2.1	Mesure d'importance basée sur les quantiles	83
3.2.2	Mesure d'importance shiftée	84
3.3	Indices basés sur la décomposition fonctionnelle	85
3.3.1	Procédé d'orthogonalisation	86
3.4	Reconstruction par méta-modélisation	87
II	Méthodes d'optimisation et de sélection de variables	93
4	Méthodes d'optimisation	95
4.1	Résolution de système linéaire	95
4.2	Formule de Sherman-Morrison et polynômes locaux	99
5	Méthodes de sélection de variables	103
5.1	Introduction	103
5.2	L'approximation <i>greedy</i>	108
5.2.1	L'approche <i>forward</i>	108
5.2.2	L'approche <i>backward</i>	109
5.2.3	L'approche <i>forward-backward</i>	109
5.3	Algorithmes pour régression Lasso	111
5.3.1	Algorithme <i>forward stagewise regression</i>	112
5.3.2	Méthode LARS adaptée au Lasso	113
III	Indices généralisés pour variables dépendantes	117
	Introduction to Chapter 6	119
6	Generalized decomposition for dependent variables	121
6.1	Introduction	121
6.2	Generalized Hoeffding decomposition-application to SA	124
6.2.1	Notation and first assumptions	124
6.2.2	Sobol sensitivity indices	125
6.2.3	Generalized decomposition for dependent inputs	126
6.2.4	Generalized sensitivity indices	128
6.3	Examples of distribution function	129
6.3.1	Boundedness of the inputs density function	129
6.3.2	Examples of distribution of two inputs	131

6.4	Estimation	132
6.4.1	Models of $p = 2$ input variables	133
6.4.2	Generalized IPDV models	135
6.5	Numerical examples	136
6.5.1	Two-dimensional IPDV model	137
6.5.2	Linear four-dimensional model	138
6.5.3	River flood inundation	139
6.6	Conclusions and perspectives	140
Introduction to Chapter 7		143
7	Generalized sensitivity indices and numerical methods	149
7.1	Introduction	149
7.2	Generalized Sobol sensitivity indices	152
7.2.1	First settings	153
7.2.2	Generalized Sobol sensitivity indices	154
7.3	The hierarchical orthogonal Gram-Schmidt procedure	156
7.4	Estimation of the generalized sensitivity indices	161
7.4.1	Least-Squares estimation	161
7.4.2	The variable selection methods	162
7.4.3	Summary of the estimation procedure	163
7.5	Asymptotic results	163
7.6	Application	165
7.6.1	A test case	165
7.6.2	The tank pressure model	166
Appendices		172
A	Generalized decomposition for dependent variables	173
A.1	Generalized Hoeffding decomposition	173
A.1.1	Generalized decomposition for dependent inputs	173
A.1.2	Generalized sensitivity indices	177
A.2	Examples of distribution function	177
A.2.1	Boundedness of the inputs density function	177
A.2.2	Examples of distribution of two inputs	178
A.3	Estimation	179
A.3.1	Model of $p = 2$ input variables	179
B	Generalized sensitivity indices and numerical methods	181
B.1	Notation	181
B.1.1	Reminder	181
B.1.2	New settings	182
B.2	Preliminary results	183
B.3	Main convergence results	185

IV	Conclusions and Perspectives	195
	Bibliography	206

Liste de symboles

$\subset$	Inclusion stricte, c'est-à-dire $A \subset B \implies A \cap B \neq B$
$\subseteq$	Inclusion avec égalité possible
$\not\subset, (\not\subseteq)$	Non inclus (strictement)
$[1 : k]$	Ensemble $\{1, \dots, k\}$ pour $k \in \mathbb{N}^*$
S	Collection de tous les sous-ensembles de $[1 : p]$
$S \setminus \{\emptyset\}$	Collection S privé de l'ensemble vide
u, v, w	Éléments de S
$-u$	Complémentaire de l'ensemble u
$ u $	Cardinal de l'ensemble u
$u \setminus v$	L'ensemble u privé de l'ensemble v
$\mathbf{X}, \mathbf{x}$	Vecteur aléatoire $\mathbf{X} = (X_i)_{i \in [1:p]}$ et vecteur $\mathbf{x} = (x_i)_{i \in [1:p]} \in \mathbb{R}^p$
$\mathbf{X}_{\mathbf{u}}, \mathbf{x}_{\mathbf{u}}$	Sous-vecteur de $\mathbf{X}$ tel que $\mathbf{X}_{\mathbf{u}} = (X_i)_{i \in u}$ (resp. $\mathbf{x}_{\mathbf{u}} = (x_i)_{i \in u}$)
$\mathbf{X}_{-\mathbf{u}}, \mathbf{x}_{-\mathbf{u}}$	Sous-vecteur de $\mathbf{X}$ tel que $\mathbf{X}_{-\mathbf{u}} = (X_i)_{i \notin u}$ (resp. $\mathbf{x}_{-\mathbf{u}} = (x_i)_{i \notin u}$)
$P_{\mathbf{X}}, p_{\mathbf{X}}$	Distribution jointe (resp. densité jointe) de $\mathbf{X}$
$P_{\mathbf{X}_{\mathbf{u}}}, p_{\mathbf{X}_{\mathbf{u}}}$	Distribution marginale (resp. densité marginale) de $\mathbf{X}_{\mathbf{u}}$
$P_{\mathbf{X}_{-\mathbf{u}}}, p_{\mathbf{X}_{-\mathbf{u}}}$	Distribution marginale (resp. densité marginale) de $\mathbf{X}_{-\mathbf{u}}$
$ \beta $	Valeur absolue du nombre β
$P_H(T)$	Projection de T sur l'espace H
$F_{\mathbf{X}}, F_Y$	Fonction de répartition de $\mathbf{X}$ (resp. Y)
$\langle \cdot, \cdot \rangle, \ \cdot\ $	Produit scalaire et norme de l'espace $L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$
$\mathcal{M}_{a,b}(K)$	Ensemble de matrices à a lignes et b colonnes à valeurs dans K , pour $a, b \in \mathbb{N}^*$
${}^t A$	Transposée de la matrice A
I_p	Matrice identité de taille p
$\xrightarrow{a.s.}$	Convergence presque sûre
$\text{Span}\{H\}$	Espace engendré par H

Introduction

Introduction

Pour décrire des systèmes complexes, les scientifiques ont recours à un modèle mathématique qu'ils utilisent comme une version simplifiée, mais néanmoins représentative, du système à étudier. L'expansion des performances numériques des machines de calcul nous a poussé à considérer des modèles de plus en plus complexes, dans le but de se rapprocher au mieux de la réalité physique. Cette complexité se traduit par un nombre très important d'entrées/sorties-parfois plusieurs milliers- mais aussi sur les interactions et les dépendances qu'il peut exister entre ces paramètres. Néanmoins, un modèle trop complexe peut engendrer un certain nombre de problèmes qu'il convient de résoudre pour sa bonne utilisation :

- Les paramètres d'entrée peuvent être sujets à de nombreuses sources d'incertitude, attribuées le plus souvent aux erreurs de mesure, ou à une absence d'information. Cette incertitude génère de la variabilité dans le modèle qui contribue à réduire sa robustesse. Pour augmenter la fiabilité des prédictions du modèle, il est essentiel d'établir l'origine de cette variabilité dans le but de la réduire.
- Un nombre de paramètres très élevé peut rendre le modèle quasi inutilisable en pratique, en particulier si son étude nécessite un grand nombre d'évaluations. Cette complexité peut également rendre le modèle difficilement interprétable et donc sans intérêt pour l'utilisateur.
- L'interaction entre les facteurs constitue un enjeu de taille dans certaines disciplines (par exemple l'industrie pharmaceutique). Il peut être primordial de bien les appréhender.

En outre, l'objectif est de proposer un modèle assurant un bon compromis entre une proximité avec la réalité, mais qui engendre souvent la complexité, et une simplification qui assure de rendre le modèle utilisable et interprétable.

Répondre à ce compromis en offrant des solutions aux problèmes posés fait l'objet de l'analyse de sensibilité [SCS00, STCR04, SRA⁺08]. Plus précisément, dans le modèle entrées-sortie

$$y = \eta(\mathbf{x}), \quad \mathbf{x} \in \mathbb{R}^p, \quad y \in \mathbb{R},$$

le but premier de l'analyse de sensibilité est de comprendre comment les variations des paramètres d'entrée impactent sur la sortie du modèle. De

cette étude, nous pouvons ainsi détecter les variables qui contribuent le plus à l'incertitude dans le modèle, et par conséquent les variables les moins influentes, qui peuvent être fixées à une valeur nominale.

Les méthodes d'analyse de sensibilité sont grossièrement catégorisées en trois grandes classes, que nous décrivons brièvement ici.

1. Les méthodes de criblage regroupent des méthodes qui ont pour objectif d'étudier les aspects qualitatifs de la sensibilité. En outre, l'essentiel est de savoir quelles variables ont un effet sur la sortie par exemple, mais en ignorant si cet effet est important. Ces méthodes sont essentiellement utilisées pour des modèles dépendant d'un très grand nombre de facteurs et pour lesquels il s'agit de réduire de manière assez drastique ce nombre. La méthode la plus standard consiste à considérer, pour les valeurs extrêmes de chaque paramètre d'entrée et d'étudier un à un l'impact de ce différentiel sur la sortie. Ces techniques ont l'avantage d'être simples à implémenter et très peu coûteuses. Néanmoins, les méthodes de criblage sont d'usage très limité de par l'apport d'une information qualitative, mais aussi parce que ces méthodes ne prennent pas en compte les effets d'interaction, dont le rôle est parfois prédominant dans un modèle.
2. Les méthodes locales sont des techniques quantitatives, et consistent à étudier l'impact d'une perturbation locale autour de la valeur nominale d'une entrée. Cela revient à s'intéresser aux dérivées partielles,

$$\frac{\partial \eta}{\partial x_i}(\mathbf{x}^0), \quad i = 1, \dots, p,$$

puis à les comparer entre elles pour comprendre comment ces variations influent sur la sortie. Une telle analyse est cependant limitée, car elle ne prend en compte aucune connaissance relative aux paramètres, à l'exception d'une valeur nominale. Ainsi, une très grande variation des paramètres risque d'être ignorée.

3. Les méthodes globales reposent sur le fait que les entrées et sortie ne sont plus déterministes, mais aléatoires. Cette classe de méthodes est la plus générale et permet de considérer les plages de variation des paramètres. Nous considérons que les entrées-sortie sont des variables aléatoires $(Y, \mathbf{X})$ à valeurs dans $\mathbb{R} \times \mathbb{R}^p$. Nous nous intéressons ici à la sous-classe de techniques basées sur la variance, qui se révèle être très exploitée. L'origine de cette approche remonte aux travaux de Fisher [Fis25] sur l'analyse de la variance (ANOVA). Aujourd'hui, l'ANOVA consiste à mettre en place un modèle mathématique composé d'effets principaux (c'est-à-dire dépendants d'une seule variable) et d'effets d'interaction (dépendant de plusieurs variables), puis de décomposer la variance globale en variances partielles dans le but d'ana-

lyser les différences entre les différents groupes constitués. L'objectif est donc de connaître la significativité des effets sur la sortie. Bien que l'ANOVA soit utilisée pour des variables discrètes (on parle de facteurs et de leurs modalités), la décomposition de la sortie du modèle peut être étendue à des fonctions continues, qui s'écrit

$$\eta(\mathbf{X}) = \eta_{\emptyset} + \sum_{i=1}^p \eta_i(X_i) + \sum_{i < j=1}^p \eta_{ij}(X_i, X_j) + \cdots + \eta_{1 \dots p}(\mathbf{X}),$$

où $\eta_{\emptyset}$ est une constante, η_i un effet principal, et η_u , $|u| \geq 2$ est un effet d'interaction. Cependant, il existe une infinité de décompositions de η si aucune propriété n'est imposée à ses composantes. En imposant des conditions d'orthogonalité notamment, on peut espérer obtenir l'unicité des composantes, et, étant donnée une norme, en déduire une décomposition ANOVA qui rendra possible une évaluation de l'importance des effets sur la sortie.

Dans le cas où la distribution des entrées est indépendante, l'unicité de la décomposition ANOVA est assurée par le fait que les termes sont deux à deux orthogonaux. Dans ce cas, nous obtenons que $\|\eta\|^2$ est exactement la décomposition $\sum_u \|\eta_u\|^2$. Ces quantités constituent

l'indice de Sobol [\[Sob93\]](#), une des mesures de sensibilité les plus utilisées. Les termes η_u s'expriment de manière analytique en fonction de η , réduisant ainsi le problème d'analyse de sensibilité à un problème numérique d'estimation des composantes η_u . Le calcul de telles quantités, qui se résume à un calcul d'intégrales, peut être abordé de différentes manières. Une représentation de l'analyse de sensibilité globale est donnée en Figure [1](#).

Néanmoins, lorsque l'hypothèse d'indépendance des entrées n'est plus vérifiée, il n'y a plus de contraintes évidentes à imposer pour assurer l'unicité de la décomposition, et l'exploitation des indices de Sobol peut alors entraîner une analyse erronée des contributions dans le modèle. C'est dans ce contexte que nous plaçons nos contributions.

En se concentrant sur la décomposition ANOVA, objet mathématique rigoureux et pertinent pour l'analyse de sensibilité, l'objectif de cette thèse est de proposer des solutions pour traiter la sensibilité d'entrées dépendantes dans un modèle stochastique. Notre travail vise à introduire un indice de sensibilité qui repose sur une décomposition exacte de la fonction modèle, en prenant compte du caractère non indépendant des entrées. L'objectif de ces travaux est de proposer des méthodes numériques qui assurent la mise en œuvre pratique de cet indice.

Dans la Partie I, nous détaillons les outils de l'analyse de sensibilité globale sous l'hypothèse d'indépendance des variables d'entrée. En particulier, nous nous concentrons sur la décomposition ANOVA précédemment évoquée, et de l'indice de Sobol qui en découle. Nous revenons sur les travaux de Hoeffding à l'origine de cette décomposition. Nous mettons l'accent sur la preuve rigoureuse de l'unicité, ce qui constitue le point de départ de nos travaux. Par la suite, nous présentons les principales méthodes du calcul de l'indice de Sobol : nous traitons ici des méthodes directes d'intégration numérique et des méthodes indirectes, reposant sur la construction préalable d'un méta-modèle.

Dans un second temps, nous explicitons les différentes notions de dépendance, c'est-à-dire les liens structurels qu'il peut exister entre les variables lorsqu'elles sont dépendantes entre elles. Il s'agira notamment d'aborder la théorie des copules, qui constitue un puissant outil théorique. Nous nous attachons aussi à montrer que les copules donnent lieu à une très vaste littérature sur leur utilisation pratique. Nous donnons quelques unes de ces techniques pratiques, étayées par des exemples.

Enfin, la dernière partie introductive fait état des travaux réalisés en analyse de sensibilité pour des entrées dépendantes. Les différentes solutions apportées se classent en plusieurs catégories, pour lesquelles nous détaillerons les tenants et les aboutissants.

En Partie II, nous explicitons des méthodes d'optimisation classiques et les méthodes de sélection de variables, qui seront exploitées dans nos travaux contributifs. Les techniques d'optimisation, relatées au Chapitre 4, sont des outils usuels de résolution de systèmes linéaires, mais constituent les fondements de la méthode d'estimation développée au Chapitre 6. Les méthodes de sélection de variables font l'objet du Chapitre 5. Dans ce chapitre, nous introduisons d'abord les motivations et le principe de sélection de variables. Il s'agira ensuite de présenter les différentes solutions pratiques auxquelles nous nous intéressons dans le Chapitre 7.

Dans la Partie III, nous regroupons nos contributions traitant du problème de dépendance en analyse de sensibilité. Cette partie se compose d'une publication, qui constitue le Chapitre 6, et d'une prépublication, qui constitue le Chapitre 7.

Dans le Chapitre 6, nous introduisons une décomposition fonctionnelle de la fonction modèle. Sous des conditions d'orthogonalité dite hiérarchique des composantes, nous assurons l'existence et l'unicité d'une telle décomposition. Une décomposition ANOVA permettra d'en déduire un indice de sensibilité, généralisant celui de Sobol. Enfin, une méthode numérique basée sur les propriétés de projection met en application cet indice. Cependant, cette méthode n'est numériquement utilisable que pour des modèles de faible dimension et/ou des modèles dont les entrées sont dépendantes par paires.

Dans le Chapitre 7, nous apportons une extension au Chapitre 6 en proposant une méthode beaucoup moins restrictive que la précédente, c'est-à-dire utilisable pour tout type de modèles. La procédure numérique que nous détaillons repose sur un procédé d'orthogonalisation du type Gram-Schmidt qui permet de construire pas à pas des fonctions approchées des composantes de notre décomposition.

Par le biais de ces travaux, nous tentons de pallier les difficultés liées à la prise en compte de la dépendance dans un modèle, par une approche théorique qui se veut rigoureuse et cohérente. Nous tentons également de rendre cette approche pratique, et utilisable par tous.

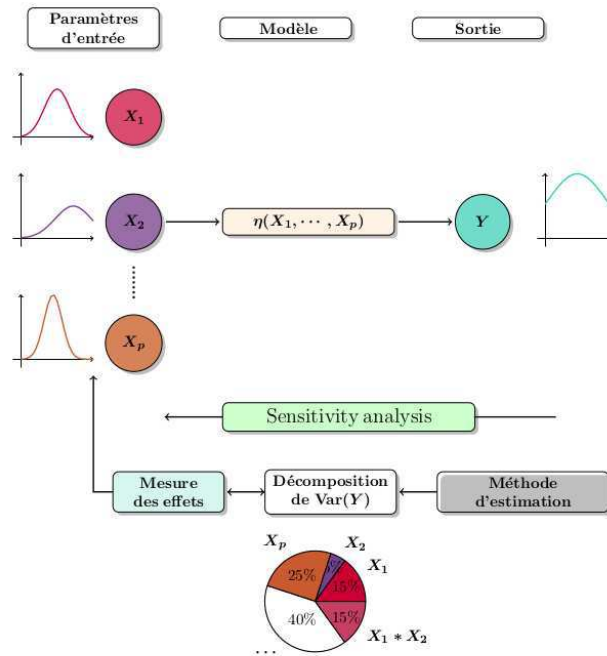


FIGURE 1 – Représentation schématique de l'analyse de sensibilité basée sur la décomposition ANOVA

Cadre d'étude

Cadre d'étude

Cette section a pour objectif d'introduire les notations qui serviront à l'ensemble du manuscrit. Nous donnons également les hypothèses principales sous lesquelles les résultats qui vont suivre seront vérifiés.

Soit $(\Omega, \mathcal{A}, P)$ un espace de probabilité, et $\mathbf{X} = (X_1, \dots, X_p)$ un vecteur de variables aléatoires définies sur $(\Omega, \mathcal{A}, P)$ et à valeurs dans $\mathbb{R}^p$. Nous définissons Y , une variable aléatoire à valeurs réelles comme

$$\begin{array}{ccccc} Y : & (\Omega, \mathcal{A}, P) & \rightarrow & (\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}) & \rightarrow & (\mathbb{R}, \mathcal{B}(\mathbb{R})) \\ & \omega & \mapsto & \mathbf{X}(\omega) & \mapsto & \eta(\mathbf{X}(\omega)) \end{array}$$

où $\eta : \mathbb{R}^p \rightarrow \mathbb{R}$ est une fonction mesurable, et la loi du vecteur $\mathbf{X}$, notée $P_{\mathbf{X}}$, est la mesure image de P par $\mathbf{X}$.

Soit ν une mesure σ -finie produit définie sur $(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p))$. Supposons que $P_{\mathbf{X}}$ est absolument continue par rapport à ν . Dans ce cas, nous définissons la densité $p_{\mathbf{X}}$ de la loi de $\mathbf{X}$, comme

$$p_{\mathbf{X}} := \frac{dP_{\mathbf{X}}}{d\nu}$$

Nous supposons que $\eta \in L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$ (noté $L^2_{\mathbb{R}}$), dont le produit scalaire et la norme associée s'expriment comme

$$\langle h_1, h_2 \rangle = \int h_1(\mathbf{x})h_2(\mathbf{x})p_{\mathbf{X}}d\nu(\mathbf{x}) = \mathbb{E}(h_1(\mathbf{X})h_2(\mathbf{X})), \quad \forall h_1, h_2 \in L^2_{\mathbb{R}}$$

$$\|h\|^2 = \langle h, h \rangle = \mathbb{E}(h(\mathbf{X})^2), \quad \forall h \in L^2_{\mathbb{R}}$$

où $\mathbb{E}(\cdot)$ est l'espérance. La variance et la covariance de variables aléatoires se définissent respectivement comme $V(\cdot) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))^2]$ et $\text{Cov}(\cdot, *) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))(* - \mathbb{E}(*))]$. L'espérance et la variance conditionnelle se notent comme $\mathbb{E}(\cdot|*)$ et $V(\cdot|*)$ respectivement.

L'ensemble discret $\{1, \dots, p\}$ est noté $[1 : p]$, et u, v, w dénotent des sous-ensembles d'indices distincts de $[1 : p]$. S représente la collection de tous

les sous-ensembles de $[1 : p]$. Nous définissons aussi le complémentaire de u comme $-u := [1 : p] \setminus u$, pour tout $u \in S$.

Pour $u = (u_1, \dots, u_t) \in S$ de taille $|u| = t \geq 0$, $\mathbf{X}_{\mathbf{u}}$ est un sous-vecteur aléatoire de $\mathbf{X}$ à valeurs dans $\mathbb{R}^t$, c'est-à-dire $\mathbf{X}_{\mathbf{u}} = (X_{u_1}, \dots, X_{u_t})$. Par convention, si $t = 0$, alors $u = \emptyset$ et $X_{\emptyset} = 1$. Pour $t \geq 1$, la mesure ν_u définie sur $(\mathbb{R}^t, \mathcal{B}(\mathbb{R}^t))$ est donnée par

$$\nu_u(d\mathbf{x}_{\mathbf{u}}) = \nu_{u_1}(dx_{u_1}) \otimes \dots \otimes \nu_{u_t}(dx_{u_t})$$

La densité $p_{\mathbf{X}_{\mathbf{u}}}$ et la distribution $P_{\mathbf{X}_{\mathbf{u}}}$ de $\mathbf{X}_{\mathbf{u}}$ sont données par

$$p_{\mathbf{X}_{\mathbf{u}}} = \int_{\mathbb{R}^{p-t}} p_{\mathbf{X}}(\mathbf{x}) d\nu_{-u}(\mathbf{x}_{-u}), \quad P_{\mathbf{X}_{\mathbf{u}}}(d\mathbf{x}_{\mathbf{u}}) = p_{\mathbf{X}_{\mathbf{u}}}(\mathbf{x}_{\mathbf{u}}) \nu_u(d\mathbf{x}_{\mathbf{u}})$$

Enfin, nous supposons que Y est de variance non nulle.

Première partie

Etat de l'art

Chapitre 1

Analyse de sensibilité globale et indépendance

La compréhension de l'analyse de sensibilité pour des facteurs d'entrée dépendants a pour préalable la maîtrise des techniques de sensibilité pour des systèmes à entrées indépendantes. Ce chapitre fait l'objet d'une discussion sur des mesures de sensibilité globale. L'accent est mis sur des indices basés sur la variance globale du modèle.

La première partie de la Section 1.1 introduit brièvement les mesures basées sur la variance, et plus particulièrement l'indice de Sobol. La seconde partie discute des indices de Sobol [Sob93] tels qu'ils sont présentés dans les ouvrages classiques de sensibilité. Nous verrons notamment qu'ils sont issus de la décomposition de Hoeffding-Sobol, dont les propriétés assurent une définition exacte et univoque de ces indices de sensibilité basés sur la variance. Ces indices feront l'objet d'une étude approfondie à la Section 1.2, où un éclairage différent est adopté. Nous développons ici la décomposition de Hoeffding [Hoe48], antérieure aux travaux de Sobol. Mes travaux de thèse sont une extension de cette décomposition dans le cas d'entrées dépendantes. Enfin, la Section 1.3 donne un état de l'art des méthodes d'estimation des indices de Sobol. De la méthode la plus standard, la méthode de Monte Carlo [Sob93], aux techniques plus élaborées, avec les décompositions spectrales [CLS78, Tis12, Sud08], nous essayerons ici de donner un large aperçu des techniques numériques les plus utilisés dans le domaine. Nous verrons d'ailleurs par la suite que ces approches ont largement inspiré les nouvelles contributions en analyse de sensibilité pour des modèles à entrées dépendantes, dont le travail contributif de cette thèse.

1.1 Méthodes basées sur la variance : un historique

1.1.1 Mesures d'importance

Pour étudier la variation de la réponse d'un modèle, la technique la plus naturelle consiste à "fixer" un/des paramètre(s) d'entrée, et d'étudier l'effet de cette action sur la variabilité de la sortie. Suivant cette idée, Hora & Iman [HI86] introduisent la *mesure d'importance* d'une variable X_i comme

$$I_i = \sqrt{V(Y) - \mathbb{E}[V(Y|X_i)]}$$

Pour des questions de robustesse, cette mesure sera plus tard modifiée par les mêmes auteurs au profit du nouvel indice [IH90],

$$\frac{V[\mathbb{E}(\log Y|X_i)]}{V(\log Y)}$$

Plus tard, McKay [McK97] exploite le modèle de régression classique et la décomposition de la variance totale [WHH06] pour introduire le *rapport de corrélation*, défini comme suit

$$\eta^2 = \frac{V[\mathbb{E}(Y|X_i)]}{V(Y)}$$

Notons que cet indice permet une interprétation simple de la contribution de X_i dans le modèle. En effet, étant donné que $\mathbb{E}(Y|X_i)$, fonction de X_i , est la meilleure approximation de Y sur l'espace engendré par X_i , la variance de cette quantité mesure la dispersion de Y due à la variation de X_i . Le ratio de $V[\mathbb{E}(Y|X_i)]$ avec $V(Y)$ permet ainsi de comparer cette dispersion à celle de Y seul.

Cette mesure coïncide exactement avec l'indice de Sobol du premier ordre [SCS00], que nous introduisons dans la suite en toute généralité.

1.1.2 L'indice de Sobol

Dans des travaux antérieurs, Sobol montre l'utilité des bases de Fourier pour la décomposition de la fonction du modèle η [Sob67], décomposition qu'il généralisera plus tard, et qui constitue aujourd'hui la base de la construction de l'indice de Sobol [Sob93].

Cette décomposition se présente dans les nombreux ouvrages d'analyse de sensibilité [SCS00, STCR04, SRA⁺08, CIBN05] comme suit.

Soit $\mathbf{x} = (x_1, \dots, x_p)$ un vecteur de l'hypercube unité

$$\Omega = \{\mathbf{x} \in \mathbb{R}^p, 0 \leq x_i \leq 1, \forall i \in [1 : p]\}.$$

La décomposition fonctionnelle ANOVA de $\eta(\mathbf{x})$ consiste à exprimer η comme une somme de fonctions de dimension croissante :

$$\begin{aligned}\eta(\mathbf{x}) &= \eta_\emptyset + \sum_{i=1}^p \eta_i(x_i) + \sum_{1 \leq i < j \leq p} \eta_{i,j}(x_i, x_j) + \cdots + \eta_{1,\dots,p}(\mathbf{x}) \\ &= \sum_{u \in S} \eta_u(\mathbf{x}_u).\end{aligned}\tag{1.1}$$

La fonction η se décompose ainsi en une somme de 2^p fonctions, où $\eta_\emptyset$ est une constante, η_i , $i \in [1 : p]$ sont les effets principaux, les fonctions $\eta_{i,j}$, $i < j \in [1 : p]$ représentent les effets d'interaction d'ordre 2, etc. La décomposition (1.1), aussi connue sous le nom de *High Dimensional Model Representations* (HDMR) [LRR01] dans d'autres circonstances, n'est pas unique, car il peut exister une infinité de choix de ses composantes. Pour assurer l'unicité de (1.1), il est nécessaire d'imposer des contraintes d'orthogonalité sur les composantes. Elles s'expriment comme,

$$\int_{[0,1]} \eta_u(\mathbf{x}_u) dx_i = 0 \quad \forall i \in u, \forall u \in S.\tag{1.2}$$

Une conséquence de la condition (1.2) est que les termes de la décomposition (1.1) sont deux à deux orthogonaux, c'est-à-dire

$$\int_{[0,1]^p} \eta_u(\mathbf{x}_u) \eta_v(\mathbf{x}_v) d\mathbf{x} = 0 \quad \forall u \neq v \in S.$$

Ainsi, toute fonction de carré intégrable admet une unique représentation en fonctionnelle ANOVA si ses variables sont définies sur un hypercube unité. La propriété sympathique est que les composantes sont deux à deux orthogonales.

Revenons maintenant à un cadre stochastique. Considérons que ν est la mesure de Lebesgue sur $(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p))$ et $\mathbf{X}$ est un vecteur de variables aléatoires indépendantes uniformes, c'est-à-dire $\mathbf{X} \sim \mathcal{U}([0, 1]^p)$ et $p_{\mathbf{X}} = \mathbb{I}_{[0,1]^p}$, où $\mathbb{I}(\cdot)$ dénote la fonction indicatrice. La décomposition (1.1), sous les contraintes données en (1.2) reste vraie sur le modèle $\eta(\mathbf{X})$. Celle-ci assure la décomposition de la variance globale en une somme de variances partielles comme donnée ci-dessous

$$\begin{aligned}V(Y) &= \sum_{i=1}^p V[\eta_i(X_i)] + \sum_{1 \leq i < j \leq p} V[\eta_{i,j}(X_i, X_j)] + \cdots + V[\eta_{1,\dots,p}(\mathbf{X})] \\ &= \sum_{u \in S} V[\eta_u(\mathbf{X}_u)].\end{aligned}$$

Il en résulte la définition des indices de Sobol explicitée dans ses travaux [Sob93, Sob01], et redonnée dans la Proposition 1.1.

Proposition 1.1. *L'indice de Sobol S_u d'ordre $|u|$ mesurant la contribution du groupe de variables $\mathbf{X}_u$ dans le modèle $Y = \eta(\mathbf{X})$ est donné par*

$$S_u = \frac{V(\eta(\mathbf{X}_u))}{V(Y)} = \frac{V[\mathbb{E}(Y|\mathbf{X}_u)] + \sum_{v \subset u} (-1)^{|u|-|v|} V[\mathbb{E}(Y|\mathbf{X}_v)]}{V(Y)}. \quad (1.3)$$

En particulier, l'indice d'ordre 1 mesurant la contribution de X_i s'exprime comme

$$S_i = \frac{V[\mathbb{E}(Y|X_i)]}{V(Y)}. \quad (1.4)$$

La Proposition 1.1 témoigne de l'attrait de la décomposition fonctionnelle ANOVA (1.1) pour obtenir une définition claire et précise de ce qui constitue les indices de Sobol. Les projections de Y sur les différents espaces engendrés par des groupes de variables $\mathbf{X}_u$ rendent, comme précédemment évoqué, l'interprétation de tels indices aisée. De plus, par la Proposition 1.1, il découle les propriétés suivantes,

$$0 \leq S_u \leq 1, \quad \forall u \in S \setminus \{\emptyset\},$$

et

$$\sum_{u \in S \setminus \{\emptyset\}} S_u = 1.$$

Par conséquent, si S_u est proche de 1, cela signifie que la part de variabilité induite par $\mathbf{X}_u$ est très importante dans le modèle. Au contraire, $\mathbf{X}_u$ est très peu influente si S_u est proche de 0. Dans ce cas, les variables constituant l'ensemble $\mathbf{X}_u$ peuvent être fixées à une valeur nominale, et peuvent ne plus être considérées comme des variables agissant sur la variance de la sortie.

Pour construire ces indices, Sobol suppose que les facteurs d'entrée sont uniformes sur $[0, 1]^p$. Signalons cependant que ce raisonnement est généralisable à toute distribution produit par simple transformation (en passant par les fonctions de répartition inverse par exemple). De plus, la décomposition (1.1) relève des travaux antérieurs de Hoeffding sur les U-statistiques [Hoe48], et qui se basent essentiellement sur la projection de Y sur des espaces orthogonaux. Nous présentons maintenant ces travaux en détail, car ils se rapprochent de la démarche adoptée dans cette thèse pour construire des indices de sensibilité généralisés dans le cas dépendant.

1.2 Décomposition de Hoeffding et indices de Sobol

Dans cette section, nous supposons toujours que les variables d'entrée sont indépendantes, avec $P_{\mathbf{X}}$ une distribution produit quelconque c'est-à-dire $P_{\mathbf{X}} = P_{X_1} \otimes \cdots \otimes P_{X_p}$. L'objectif de cette partie est de présenter la décomposition ANOVA par le biais du lemme de Hoeffding, tel qu'il est présenté dans [VDV98] pages 158-159. En effet, la décomposition ANOVA fonctionnelle fut initiée par les travaux de Hoeffding sur les U-statistiques [Hoe48], et reprise par divers auteurs au cours du temps [Nel65, ES81]. Ces travaux sont donc antérieurs à la Proposition 1.1.

Définissons d'abord les espaces suivants,

Définition 1.1. Soit $u \in S$. Pour $u = \emptyset$, les espaces $H_{\emptyset}$ et $H_{\emptyset}^0$ représentent l'espace des constantes. Pour $u \neq \emptyset$, les espaces H_u et H_u^0 se définissent comme suit

$$H_u = \{h_u(\mathbf{X}_u), \mathbb{E}(h_u(\mathbf{X}_u)^2) < +\infty\} \subset L_{\mathbb{R}}^2,$$

et

$$H_u^0 = \{h_u(\mathbf{X}_u) \in H_u, \mathbb{E}(h_u(\mathbf{X}_u)|\mathbf{X}_v) = 0, \forall v, |v| < |u|\}.$$

Avec, pour convention, $\mathbb{E}(h_u(\mathbf{X}_u)|\emptyset) = \mathbb{E}(h_u(\mathbf{X}_u))$.

Les espaces H_u^0 , $u \in S$ bénéficient de quelques propriétés, que nous énonçons maintenant. Ces propriétés sont à l'origine du lemme de Hoeffding, qui sera énoncé ensuite.

Proposition 1.2. Soit $u \in S$. Alors, par indépendance des $X_1, \dots, X_p$,

$$H_u^0 \perp H_v^0, \quad \forall u \neq v \in S.$$

Démonstration de la Proposition 1.2. Soit $u \in S$, et $h_u \in H_u^0$. Montrons tout d'abord l'égalité suivante :

$$\int_{\mathbb{R}} h_u(\mathbf{x}_u) dP_{X_i}(x_i) = 0, \quad i \in u. \quad (1.5)$$

Si $u = \{i\}$, alors l'égalité (1.5) est immédiate par définition de l'espace H_u^0 . Maintenant, si $|u| \geq 2$, $-i$ s'écrit nécessairement comme $-i = (u \setminus i, -u)$. Par indépendance des composantes de $\mathbf{X}$, il vient

$$\begin{aligned} \int_{\mathbb{R}} h_u(\mathbf{x}_u) dP_{X_i}(x_i) &= \int \left[\int_{\mathbb{R}} h_u(\mathbf{x}_u) dP_{X_i} \right] dP_{\mathbf{X}_{-u}} \\ &= \int h_u(\mathbf{x}_u) dP_{\mathbf{X}_{-w}} \\ &= \mathbb{E}(h_u(\mathbf{X}_u)|\mathbf{X}_w) = 0 \quad \text{car } w = u \setminus i, |w| < |u|. \end{aligned}$$

Montrons maintenant que $H_u^0 \perp H_v^0$. Soient $v \neq u \in S$, $h_u \in H_u^0$, $h_v \in H_v^0$ et $i \in v \setminus u$, alors

$$\begin{aligned} \int h_u(\mathbf{x}_u) h_v(\mathbf{x}_v) dP_{\mathbf{X}} &= \int h_u(\mathbf{x}_u) \left(\int h_v(\mathbf{x}_v) dP_{X_i} \right) dP_{\mathbf{X}_{-i}} \\ &= 0 \quad \text{par (1.5).} \end{aligned} \quad (1.6)$$

Pour $i \in u \setminus v$, on peut appliquer le même raisonnement en inversant les rôles de h_u et h_v dans (1.6). ■

Enonçons le lemme à l'origine de la décomposition de Hoeffding :

Lemme 1.1 (Décomposition de Hoeffding). *Soient $X_1, \dots, X_p$ des variables indépendantes, et soit T une variable aléatoire telle que $\mathbb{E}(T^2) < \infty$. Alors, $\forall u \in S$, la projection de T dans H_u^0 s'exprime comme*

$$P_{H_u^0}(T) = \sum_{v \subseteq u} (-1)^{|u|-|v|} \mathbb{E}(T | \mathbf{X}_v). \quad (1.7)$$

De plus, si, pour un $u \in S$ donné, $T \perp H_v^0$, $\forall v \subseteq u$, alors $\mathbb{E}(T | \mathbf{X}_u) = 0$. Par conséquent,

$$H_u \subset \bigoplus_{v \subseteq u} H_v^0. \quad (1.8)$$

Démonstration du Lemme 1.1.

Soit $u \in S$. D'une part, montrons que $P_{H_u^0}(T)$ tel qu'il est défini en (1.7) est un opérateur de projection orthogonale.

– Montrons que $P_{H_u^0}(T) \in H_u^0$. Pour cela, considérons, en toute généralité, $w \subset u$,

$$\begin{aligned} \mathbb{E}(P_{H_u^0}(T) | \mathbf{X}_w) &= \sum_{v \subseteq u} (-1)^{|u|-|v|} \mathbb{E}[\mathbb{E}(T | \mathbf{X}_v) | \mathbf{X}_w] \\ &= \sum_{v \subseteq u} (-1)^{|u|-|v|} \mathbb{E}(T | \mathbf{X}_v \cap \mathbf{X}_w) \quad \text{par indépendance} \\ &= \sum_{t \subseteq w} \sum_{j=0}^{|u|-|w|} (-1)^{|u|-|t|-j} \binom{|u|-|w|}{j} \mathbb{E}(T | \mathbf{X}_t) \\ &= \sum_{t \subseteq w} \left(\sum_{j=0}^{|u|-|w|} (-1)^{|u|-|w|-j} \binom{|u|-|w|}{j} \right) (-1)^{|w|-|t|} \cdot \mathbb{E}(T | \mathbf{X}_t) \\ &= 0 \quad \text{par la formule du binôme.} \end{aligned}$$

– Montrons que $P_{H_u^0}(T)$ est un opérateur orthogonal. Soit $h_u \in H_u^0$, alors

$$\begin{aligned} \mathbb{E}[(T - P_{H_u^0}(T))h_u(\mathbf{X}_u)] &= \mathbb{E}(Th_u(\mathbf{X}_u)) - \mathbb{E}[\mathbb{E}(T|\mathbf{X}_u)h_u(\mathbf{X}_u)] \\ &\quad - \sum_{v \subset u} (-1)^{|u|-|v|} \mathbb{E}[\mathbb{E}(T|\mathbf{X}_v)h_u(\mathbf{X}_u)] \\ &= \mathbb{E}(Th_u(\mathbf{X}_u)) - \mathbb{E}[\mathbb{E}(Th_u(\mathbf{X}_u)|\mathbf{X}_u)] \\ &\quad - \sum_{v \subset u} (-1)^{|u|-|v|} \mathbb{E}[\mathbb{E}(T|\mathbf{X}_v)\mathbb{E}(h_u(\mathbf{X}_u)|\mathbf{X}_v)] \\ &= 0. \end{aligned}$$

Maintenant, supposons que $T \perp H_v^0$, $\forall v \subseteq u$. Montrons que $\mathbb{E}(T|\mathbf{X}_u) = 0$ par récurrence sur $|u|$. Posons,

$$(\mathcal{H}_t), \forall u, |u| = t, \quad (T \perp H_v^0, \forall v \subseteq u) \implies (\mathbb{E}(T|\mathbf{X}_u) = 0).$$

Si $u = \emptyset$ et $T \perp H_\emptyset$, alors $\mathbb{E}(T|\emptyset) = \mathbb{E}(T) = 0$, donc $(\mathcal{H}_0)$ est vraie. Supposons maintenant que $(\mathcal{H}_{k-1})$ est vraie. Prenons $|u| = k$ et supposons que $T \perp H_v^0$, $\forall v \subseteq u$. Par conséquent, $T \perp H_w^0$, $\forall w \subseteq v \subseteq u$, et par hypothèse de récurrence, $\mathbb{E}(T|\mathbf{X}_v) = 0$, $\forall v \subset u$. Comme $T \perp H_u^0$, $P_{H_u^0}(T) = 0$, et par la formule (1.7), $P_{H_u^0}(T) = \mathbb{E}(T|\mathbf{X}_u)$, d'où $\mathbb{E}(T|\mathbf{X}_u) = 0$.

La dernière assertion du lemme est vérifiée si $h(\mathbf{X}_u) := h_u(\mathbf{X}_u) - P_{\oplus_{v \subseteq u} H_v^0}(h_u(\mathbf{X}_u))$ est égal à 0 pour $h_u \in H_u$. Notons que par la Proposition 1.2, $P_{\oplus_{v \subseteq u} H_v^0} = \sum_{v \subseteq u} P_{H_v^0}$ et $P_{H_v^0}(h_u(\mathbf{X}_u)) = 0$ pour tout $v \subset u$. D'où $h(\mathbf{X}_u) = h_u(\mathbf{X}_u) - h_u(\mathbf{X}_u) = 0$. $\blacksquare$

Ainsi, nous pouvons en déduire que la projection d'une variable de carré intégrable dans un espace H_u , pour $u \in S$ donné, peut s'écrire en une somme de projections dans les espaces contraints H_v^0 , $v \subseteq u$.

Revenons sur le modèle $Y = \eta(\mathbf{X}) \in L_{\mathbb{R}}^2$. Nous pouvons en déduire que

$$Y = \mathbb{E}(Y|\mathbf{X}) = \sum_{u \in S} \eta_u(\mathbf{X}_u), \quad \text{avec } \eta_u(\mathbf{X}_u) = P_{H_u^0}(Y).$$

De plus, les composantes $(\eta_u)_{u \in S}$ s'expriment comme

$$\begin{aligned} \eta_\emptyset &= \mathbb{E}(Y) \\ \eta_u(\mathbf{X}_u) &= \mathbb{E}(Y|\mathbf{X}_u) + \sum_{v \subset u} (-1)^{|u|-|v|} \mathbb{E}(Y|\mathbf{X}_v). \end{aligned}$$

Par orthogonalité des $(H_u^0)_{u \in S}$ (Proposition 1.2), il en résulte que

$$V(\eta_u(\mathbf{X}_u)) = V[\mathbb{E}(Y|\mathbf{X}_u)] + \sum_{v \subset u} (-1)^{|u|-|v|} V[\mathbb{E}(Y|\mathbf{X}_v)], \quad u \in S,$$

et

$$V(Y) = \sum_{u \in S} \{V[\mathbb{E}(Y|\mathbf{X}_u)] + \sum_{v \subset u} (-1)^{|u|-|v|} V[\mathbb{E}(Y|\mathbf{X}_v)]\}.$$

En normalisant chacune des composantes $V(\eta_u(\mathbf{X}_u))$ par $V(Y)$, nous retrouvons l'indice de Sobol tel qu'il a été défini à la Proposition 1.1.

Notons que la décomposition de Hoeffding, et les indices de Sobol qui en découlent, est très générale en ce sens qu'il n'est pas nécessaire d'assurer la linéarité ou la monotonie de la fonction η . Cependant, notons aussi que cette décomposition, telle qu'elle a été réalisée au travers de la preuve, n'est pas possible sans l'hypothèse d'indépendance des variables d'entrée. Néanmoins, nous verrons dans la seconde partie qu'il est possible de proposer une décomposition sous des hypothèses moins fortes sur la distribution de $\mathbf{X}$.

Signalons aussi que les indices définis à la Proposition 1.1 permettent la construction d'indices totaux, qui sont aussi souvent utilisés. L'indice total associé à X_i est donné par

$$S_i^T = \sum_{i \in u} S_u = 1 - \frac{V[\mathbb{E}(Y|\mathbf{X}_{[1:p] \setminus \{i\}})]}{V(Y)}.$$

Cet indice permet de mesurer l'influence de X_i , en incluant l'effet de ses interactions avec les autres variables du modèle. Cependant, il ne sera pas étudié dans le présent manuscrit, bien que les développements proposés dans la suite peuvent être étendus aux indices totaux.

La suite de ce chapitre est consacrée au développement de plusieurs méthodes d'estimation des indices de Sobol. Nous avons opté pour trois méthodes d'estimation, qui sont pour l'essentiel en relation avec le Chapitre 3 de cette première partie.

1.3 Estimation des indices de Sobol

En pratique, il est rare de pouvoir calculer analytiquement les indices de Sobol. En effet, le calcul des indices pour un système complexe implique la connaissance des lois des paramètres d'entrée, ainsi que la connaissance de la réponse ou de la forme du modèle. Il faut donc avoir recours à des méthodes numériques pour estimer les indices de Sobol. Dans ce qui suit, nous nous intéresserons à plusieurs techniques classiques pour l'estimation de l'indice de Sobol du premier ordre, soit, $\forall i \in [1 : p]$,

$$S_i = \frac{V[\mathbb{E}(Y|X_i)]}{V(Y)}. \quad (1.9)$$

Notons que la forme explicite donnée en (1.9) peut s'exprimer en termes d'intégrales multiples

$$S_i = \frac{\int [\int \eta(\mathbf{x}) dP_{\mathbf{X}_{-i}}]^2 dP_{X_i} - (\int \eta(\mathbf{x}) dP_{\mathbf{X}})^2}{\int \eta(\mathbf{x})^2 dP_{\mathbf{X}} - (\int \eta(\mathbf{x}) dP_{\mathbf{X}})^2}. \quad (1.10)$$

Ainsi, l'estimation des indices de Sobol consiste à proposer une bonne approximation numérique de ces intégrales.

Nous proposons ici de détailler trois méthodes d'approximation, regroupées en fonction du type de plans d'expérience qu'elles utilisent. Cependant, nous supposons pour ces trois méthodes que la mesure de référence est la mesure de Lebesgue, et que la densité des variables $\mathbf{X}$ s'écrit comme le produit des densités marginales, soit $p_{\mathbf{X}}(\mathbf{x}) = p_{X_1}(x_1) \times \cdots \times p_{X_p}(x_p)$, pour tout $\mathbf{x} \in \mathbb{R}^p$.

Dans un premier temps, nous étudierons la méthode de Monte Carlo, qui fut la méthode proposée par Sobol lui-même [Sob01]. Nous présenterons brièvement quelques unes de ses variantes, largement explorées dans la littérature [ST02]. La deuxième partie est consacrée à la méthode FAST et ses variantes, techniques reposant sur la décomposition de Fourier [CLS78]. Enfin, dans la troisième partie, nous discutons de l'estimation des indices par polynômes de chaos [Sud08], et notamment du calcul de coefficients par la technique de régression.

1.3.1 Méthodes de Monte Carlo

La technique de Monte Carlo pour estimer des indices de Sobol repose sur une réécriture de l'indice. En effet, l'indice S_i peut se réécrire de la manière suivante

$$S_i = \frac{\text{Cov}(Y, Y^*)}{V(Y)}, \quad (1.11)$$

où $Y^* = \eta(X_1^*, \dots, X_{i-1}^*, X_i, X_{i+1}^*, \dots, X_p^*)$, et X_j^* , pour $j \neq i$, est une copie indépendante et de même loi que X_j . Cette caractéristique, notamment démontrée par Janon *et al.* [Jan12, JKLR⁺12], permet d'en déduire une technique d'échantillonnage qui génèrera les entrées utilisées pour calculer Y^* à partir de celles utilisées pour calculer Y .

Considérons maintenant deux n -échantillons indépendants et identiquement distribués (i.i.d) issus de $P_{\mathbf{X}}$,

$$(x_1^l, \dots, x_p^l)_{l=1, \dots, n} \quad \text{et} \quad (x_1^{l*}, \dots, x_p^{l*})_{l=1, \dots, n}.$$

Posons ensuite, $\forall l = 1, \dots, n$,

$$y^l = \eta(x_1^l, \dots, x_p^l) \quad \text{et} \quad y^{l*} = \eta(x_1^{l*}, \dots, x_{i-1}^{l*}, x_i^l, x_{i+1}^{l*}, \dots, x_p^{l*}).$$

En remplaçant la covariance et la variance dans (1.11) par leur version empirique, on obtient alors un estimateur de S_i ,

$$\hat{S}_i = \frac{\frac{1}{n} \sum_{l=1}^n y^l y^{l*} - (\frac{1}{n} \sum_{l=1}^n y^l)(\frac{1}{n} \sum_{l=1}^n y^{l*})}{\frac{1}{n} \sum_{l=1}^n (y^l)^2 - (\frac{1}{n} \sum_{l=1}^n y^l)^2}. \quad (1.12)$$

En remarquant que $\mathbb{E}(Y) = \mathbb{E}(Y^*)$, nous pouvons remplacer l'estimateur empirique de $\mathbb{E}(Y^*)$ par l'estimateur empirique de $\mathbb{E}(Y)$, et ainsi obtenir l'estimateur proposé par Sobol [Sob93], puis repris par [ST02, HS96].

$$\hat{S}_i = \frac{\frac{1}{n} \sum_{l=1}^n y^l y^{l*} - (\frac{1}{n} \sum_{l=1}^n y^l)^2}{\frac{1}{n} \sum_{l=1}^n (y^l)^2 - (\frac{1}{n} \sum_{l=1}^n y^l)^2}. \quad (1.13)$$

Cependant, la méthode de Monte Carlo classique souffre d'un coût de calcul très élevé, car elle doit faire appel à $(p+1)n$ évaluations pour estimer l'ensemble des indices de premier ordre. De plus, elle nécessite l'évaluation d'un grand nombre de sorties pour être assez précise [DDFG01]. Dans des systèmes de grande dimension, la méthode de Monte Carlo peut donc s'avérer difficile voire impossible à utiliser. Pour remédier à ce problème, plusieurs solutions ont été proposées.

D'abord, l'utilisation d'un métamodèle (aussi appelé modèle réduit ou surface de réponse suivant l'utilisateur) permet d'approcher convenablement le modèle analytique, tout en présentant des gains majeurs en coût d'évaluations. Dans [JNP11], les auteurs s'intéressent à ce cas particulier. Plus précisément, ils s'intéressent à l'erreur d'estimation des indices, en quantifiant l'erreur provenant du métamodèle, et celle due à la méthode de Monte Carlo.

D'autres techniques ont été proposées comme des variantes à la méthode de Monte Carlo. Dans (1.13), on note que Y et Y^* ont même loi. On peut donc estimer $\mathbb{E}(Y)$ et $\mathbb{E}(Y^2)$ par $\frac{1}{n} \sum_{l=1}^n \frac{y^l + y^{l*}}{2}$ et $\frac{1}{n} \sum_{l=1}^n \frac{(y^l)^2 + (y^{l*})^2}{2}$ respectivement. Ceci donne lieu à un estimateur asymptotiquement efficace étudié dans [JKLR⁺12, MNM06].

Ces estimateurs sont souvent calculés en utilisant un échantillonnage aléatoire (i.i.d. ou non i.i.d), le plus fréquemment utilisé étant l'échantillonnage par hypercubes latins. Cependant, un échantillonnage déterministe de la loi de $\mathbf{X}$ peut aussi être considéré. On parle alors de méthodes de quasi-Monte Carlo qui utilisent des séquences de points à discrétion faible. En effet, l'estimation des indices par la méthode de Monte Carlo nécessite la génération d'échantillons aléatoires. Cependant, les quantités générées par divers logiciels sont en fait pseudo-aléatoires, et il arrive parfois que les points générés ne couvrent pas entièrement le domaine de définition des variables.

Les suites à discrédance faible permettent de remédier à ce genre de problèmes. Cette technique permet de construire des suites de points dont l'équidistribution est assurée par une mesure de discrédance, que l'on désire être aussi faible que possible [DP10]. Parmi les suites les plus connues, on peut citer la suite de Halton [KN06] ou la suite LP_τ de Sobol [Sob67, Sob76]. Néanmoins, nous ne détaillerons pas ce type de techniques ici. Nous renvoyons le lecteur à [SRA⁺08, Lem09, Owe05] pour l'utilisation de la méthode de quasi-Monte Carlo pour l'estimation des indices de Sobol.

Note 1.1. La méthode de Monte Carlo et ses variantes ont ici été présentées pour l'indice de premier ordre S_i , mais ces techniques sont généralisables pour un indice d'ordre supérieur S_u , $|u| \geq 2$. En effet, nous pouvons réécrire S_u donné en (1.3) comme

$$S_u = \frac{V[\mathbb{E}(Y|\mathbf{X}_u)]}{V(Y)} + \sum_{v \subset u} (-1)^{|u|-|v|} \frac{V[\mathbb{E}(Y|\mathbf{X}_v)]}{V(Y)}.$$

Pour estimer chaque ratio de variances (parfois appelé *closed index* dans [Jan12]), il suffit donc de considérer deux n -échantillons

$$\mathbf{x} = (\mathbf{x}_v^l, \mathbf{x}_{-v}^l)_l \quad \text{et} \quad \mathbf{x}^* = (\mathbf{x}_v^{l*}, \mathbf{x}_{-v}^{l*})_l, \quad \forall v \subseteq u,$$

et

$$y^l = \eta(\mathbf{x}^l) \quad \text{et} \quad y^{l*} = \eta(\mathbf{x}_v^{l*}, \mathbf{x}_{-v}^{l*}), \quad \forall v \subseteq u.$$

Dans ce cadre plus général, les propriétés asymptotiques de tels estimateurs sont cependant plus difficiles à établir car il faut tenir compte des covariances entre les différents estimateurs de *closed indices*. Cette étude fait notamment l'objet de [GJK⁺13].

1.3.2 Méthodes FAST et ses variantes

Introduite par les travaux initiaux de Cukier *et al.* [CS75, CLS78] afin d'étudier la cinétique de concentration dans des réactions chimiques, la méthode FAST a été enrichie au cours du temps, et continue à susciter un grand intérêt en analyse de sensibilité [STC99, XG07, TP12b, TP12a, Tis12]. En effet, les estimateurs FAST correspondent rigoureusement aux estimations des indices de Sobol définis en Section 1.1, lorsque la fonction multivariée η est étudiée du point de vue harmonique. Comme nous allons le voir, le principal intérêt de la méthode FAST est de réduire le coût de calcul des intégrales multiples (1.10) en les remplaçant par une intégrale unidimensionnelle.

Dans la suite, nous restreignons le modèle $Y = \eta(\mathbf{X})$ à des variables indépendantes uniformément distribuées sur $[0, 1]$. Rappelons néanmoins que les résultats qui suivent peuvent être appliqués à des variables $\mathbf{X}$ indépendantes quelconques, comme mentionné à la Section 1.1.2.

Considérons d'abord la décomposition spectrale de $\eta(\mathbf{x})$,

$$\eta(\mathbf{x}) = \sum_{\mathbf{k} \in \mathbb{Z}^p} c_{\mathbf{k}}(\eta) \Phi_{\mathbf{k}}(\mathbf{x}), \quad (1.14)$$

où, $\Phi_{\mathbf{k}}$ est une base de Fourier, et $c_{\mathbf{k}}(\eta)$ est le $\mathbf{k}$ -ème coefficient de Fourier de η . Il est donné par

$$c_{\mathbf{k}}(\eta) = \int_{[0,1]^p} \eta(\mathbf{x}) \exp(-2i\pi \mathbf{k} \cdot \mathbf{x}) d\mathbf{x} = \mathbb{E}[\eta(\mathbf{X}) \exp(-2i\pi \mathbf{k} \cdot \mathbf{X})], \quad (1.15)$$

où $\mathbf{k} \cdot \mathbf{x}$ est le produit scalaire euclidien de $\mathbf{k}$ avec $\mathbf{x}$. La simplification de cette intégrale par une intégrale unidimensionnelle repose sur le résultat énoncé par Weyl [Wey16], et donné comme suit,

$$\int_{[0,1]^p} g(\mathbf{x}) d\mathbf{x} = \lim_{T \rightarrow +\infty} \frac{1}{2T} \int_{-T}^T g(x_1(t), \dots, x_p(t)) dt,$$

où g est une fonction bornée et intégrable sur $[0,1]^p$, et $\forall i \in [1:p]$, $x_i(t) = \omega_i t - \lfloor \omega_i t \rfloor$, avec $\lfloor \cdot \rfloor$ la partie entière, est une courbe paramétrée.

Pour contourner les problèmes d'intégration sur un ensemble infini, $\mathbf{x}(t) = (x_i(t))_{i \in [1:p]}$ est remplacée par un signal $\mathbf{x}^*(t) = \mathbf{G}(\boldsymbol{\omega}, t)$, défini comme

$$x_i^*(t) = G_i(\sin(2\pi\omega_i t)), \quad i \in [1:p],$$

avec, pour tout $i \in [1:p]$, $\omega_i \in \mathbb{N}$, et G_i une fonction définie sur $[-1,1]$, à valeurs dans $[0,1]$. Ainsi, le $\mathbf{k}$ -ème coefficient de Fourier $c_{\mathbf{k}}(\eta)$ est approché par

$$c_{\mathbf{k}}(\eta) \simeq \int_0^1 \eta \circ \mathbf{x}^*(t) \exp(-2i\pi(\mathbf{k} \cdot \boldsymbol{\omega})t) dt. \quad (1.16)$$

En utilisant une grille de discrétisation à n points pour estimer (1.16), il vient

$$\hat{c}_{\mathbf{k}}(\eta) = \hat{c}_{\mathbf{k} \cdot \boldsymbol{\omega}}(\eta) = \frac{1}{n} \sum_{j=0}^{n-1} \eta \circ \mathbf{x}^*\left(\frac{j}{n}\right) \exp(-2i\pi(\mathbf{k} \cdot \boldsymbol{\omega})\frac{j}{n}).$$

Ainsi, nous venons de voir le caractère sympathique des bases de Fourier : la projection de toute fonction bornée et intégrable dans les bases de Fourier peut être approchée par des intégrales simples, elles-mêmes estimées sur une simple grille de discrétisation.

En utilisant ce résultat remarquable, nous allons pouvoir en déduire une estimation des indices de Sobol. En effet, la variance globale $V(Y)$ peut se réécrire en somme de coefficients de Fourier comme

$$V(Y) = \mathbb{E}(\eta(\mathbf{X})^2) - \mathbb{E}(\eta(\mathbf{X}))^2 = \sum_{\mathbf{k} \in \mathbb{Z}^p} |c_{\mathbf{k}}(\eta)|^2 - |c_0(\eta)|^2.$$

Elle peut donc s'estimer aisément. L'estimation de la décomposition de $V(Y)$ en variances partielles d'ordre 1, notée V_i , est aussi immédiate,

$$\hat{V}(Y) = \sum_{\mathbf{k} \cdot \boldsymbol{\omega}=1}^{n-1} |\hat{c}_{\mathbf{k} \cdot \boldsymbol{\omega}}(\eta)|^2$$

$$\hat{V}_i = \sum_{k_i \omega_i=1}^{n_i} |\hat{c}_{k_i \omega_i}(\eta)|^2$$

et

$$\hat{S}_i^{\text{FAST}} = \frac{\sum_{k_i \omega_i=1}^{n_i} |\hat{c}_{k_i \omega_i}(\eta)|^2}{\sum_{\mathbf{k} \cdot \boldsymbol{\omega}=1}^{n-1} |\hat{c}_{\mathbf{k} \cdot \boldsymbol{\omega}}(\eta)|^2},$$

où n_i est un ordre de troncature généralement peu élevé.

Cette méthode génère alors plusieurs questions autour de l'échantillonnage, et des choix optimaux de fonctions G_i et de fréquences ω_i , $i \in [1 : p]$. Ces questions font l'objet de travaux initiés par Cukier *et al.* [CLS78, CS75], et qui furent repris plus tard comme des variantes de la méthode FAST.

La méthode EFAST [STC99] permet de gérer les problèmes d'échantillonnage liés à l'estimation des coefficients de Fourier. En constatant que les fonctions G_i , $i \in [1 : p]$ peuvent subir une certaine forme de perturbation sans que la théorie FAST en soit affectée, on peut considérer de nouvelles courbes d'échantillonnage, et non plus suivant la courbe paramétrée $\mathbf{x}^*$ exclusivement. Cela permet donc un enrichissement de la grille d'estimation. La méthode RBD [TGM06] a été introduite dans le but de contourner les difficultés liées au choix des fréquences ω_i , $i \in [1 : p]$. Elle consiste principalement à considérer des permutations aléatoires des coordonnées de points $\mathbf{x}^*(j/n)$. Ces méthodes, plutôt heuristiques, font l'objet d'une étude rigoureuse et approfondie dans [Tis12, TP12b].

Note 1.2. La généralisation de l'indice $\hat{S}_i^{\text{FAST}}$ à un indice d'ordre supérieur est très similaire à ce que nous avons présenté, à la différence que des combinaisons linéaires de fréquences ω_i sont considérées [TP12b].

Note 1.3. Nous notons que l'échantillonnage, point essentiel de la méthode FAST, est réalisé sous l'hypothèse d'indépendance des entrées $(X_i)_{i \in [1:p]}$. Dans le cas de modèles à facteurs corrélés, cette méthode n'est plus valide en tant que telle. Dans le chapitre consacré à l'analyse de sensibilité pour modèles à entrées dépendantes (Chapitre 3), nous verrons qu'une technique d'échantillonnage est proposée afin de prendre en compte la dépendance.

Cependant, la décomposition en base de Fourier est un cas particulier de décomposition spectrale, dont la forme générale est donnée en (1.14). En

effet, nous pouvons considérer d'autres formes possibles de $\Phi_{\mathbf{k}}$, comme des bases orthonormales tensorisées, qui forment des bases adaptées à la décomposition ANOVA sur $L^2_{\mathbb{R}}$. Dans la suite, nous considérons un autre type de décomposition spectrale, celle en chaos polynomial, polynômes orthogonaux associés à des lois de probabilité.

1.3.3 Polynômes de chaos

Basée sur la théorie de Wiener [Wie38], qui a cherché à faire le lien entre processus aléatoires et théorie du chaos en mécanique, la théorie sur les polynômes de chaos se présente aujourd'hui par les travaux de Cameron et Martin [CM47] qui développèrent des bases orthogonales pour des espaces L^2 . En particulier, un des résultats remarquables est que toute variable de carré intégrable peut s'écrire sous la forme d'un développement en polynômes orthogonaux pour une mesure particulière. Similairement à la méthode FAST, notre objectif est de faire le lien entre chaos polynomial et analyse de sensibilité. En particulier, nous nous focalisons sur l'estimation des coefficients par régression telle qu'elle a été étudiée par Sudret [Sud08], et par Crestaux *et al.* [CLMM09].

Il est à noter que beaucoup de travaux affèrent aux polynômes de chaos en analyse de sensibilité. Pour plus de détails, on pourra notamment consulter [Jan12, DV07, Bla09] qui abordent différents aspects de l'analyse de sensibilité par l'utilisation de polynômes de chaos.

Les polynômes de chaos

Rappelons tout d'abord que nous nous plaçons dans l'espace de probabilité $(\Omega, \mathcal{A}, P)$. Notons Θ l'ensemble des fonctions mesurables de Ω à valeurs réelles, et $L^2(\Omega) = \{X \in \Theta, \int_{\Omega} X(\omega)^2 P(d\omega) < +\infty\} \subset \Theta$.

Soit $k \in \mathbb{N}$, et $s \in \mathbb{N}^*$. Définissons $\Gamma_k(\xi_{i_1}, \dots, \xi_{i_s})$ comme suit,

$$\Gamma_k(\xi_{i_1}, \dots, \xi_{i_s}) := \{H_{\alpha}^k(\xi_{i_1}, \dots, \xi_{i_s}), \alpha \in \Lambda_k^s\},$$

où $(\xi_{i_j})_{j \in [1:s]}$ sont s variables aléatoires indépendantes sur l'espace $(\Omega, \mathcal{A}, P)$, et H_{α}^k une famille de polynômes orthogonaux de dimension d et de degré au plus k telle que $\alpha = (\alpha_1, \dots, \alpha_d) \in \mathbb{N}^d$ est un d -uplet, pour $d \in \mathbb{N}$. Par convention, $\alpha = 0$ si $d = 0$. L'ensemble Λ_k^s est défini comme

$$\Lambda_k^s = \{\alpha \in \mathbb{N}^s, \sum_{j=1}^s \alpha_j = k\}.$$

Enfin, notons par $\mathcal{S}(\Gamma_k)$ l'espace fermé engendré par les $\Gamma_k(\xi_{i_1}, \dots, \xi_{i_s})$, aussi appelé chaos homogène d'ordre k . La décomposition de Wiener résulte de la décomposition suivante [CM47],

$$L^2(\Omega) = \bigoplus_{k=1}^{+\infty} \mathcal{S}(\Gamma_k).$$

Ainsi, toute variable aléatoire de carré intégrable X admet une représentation en polynômes de chaos [GS03],

$$X = \sum_{k \in \mathbb{N}} \sum_{\alpha \in \Lambda_k^s} \sum_{i_1, \dots, i_s} c_{i_1, \dots, i_s}^\alpha H_\alpha^k(\xi_{i_1}, \dots, \xi_{i_s}), \quad (1.17)$$

où, $H_\alpha^k(\xi_{i_1}, \dots, \xi_{i_s})$ est un polynôme de $\Gamma_k(\xi_{i_1}, \dots, \xi_{i_s})$, aussi appelé chaos polynomial d'ordre k . Pour des raisons de symétrie, cette expression peut se simplifier comme suit [GS03],

$$X = \sum_{k \in \mathbb{N}} \sum_{\alpha \in \Lambda_k^k} \sum_{i_1=1}^{+\infty} \dots \sum_{i_k=1}^{i_{k-1}} c_{i_1, \dots, i_k}^\alpha H_\alpha^k(\xi_{i_1}, \dots, \xi_{i_k}). \quad (1.18)$$

En pratique, nous remarquons que chaque chaos polynomial dépend d'un nombre infini de variables aléatoires $(\xi_i)_{i \in \mathbb{N}^*}$, résultant d'une représentation infinie. Or les dimensions infinies ne sont pas utilisables en pratique, ce qui amène à considérer un nombre fini N de variables aléatoires. L'ordre k du chaos polynomial est également tronqué à l'ordre P pour ne pas avoir une infinité de termes. Il en résulte que la représentation de X en polynômes de chaos admet au total $\binom{N+P}{P}$ composantes après cette double approximation, donnée par

$$X \simeq \sum_{k=0}^P \sum_{\alpha \in \Lambda_k^k} \sum_{i_1=1}^N \dots \sum_{i_k=1}^{i_{k-1}} c_{i_1, \dots, i_k}^\alpha H_\alpha^k(\xi_{i_1}, \dots, \xi_{i_k}). \quad (1.19)$$

La représentation en (1.19) suscite alors un certain nombre de questions :

1. A partir de quelle dimension N et quel ordre de troncature P obtient-on une bonne approximation (1.19) ?
2. Comment procède t'on au choix de distribution de $(\xi_i)_{i \in \mathbb{N}^*}$, et au choix des polynômes orthogonaux ?
3. Comment estimer les coefficients $c_{i_1, \dots, i_k}^\alpha$?

Nous verrons dans la partie suivante qu'en analyse de sensibilité, la dimension N correspond au nombre de facteurs dans le modèle, soit $N = p$. La dimension est donc fixée. L'ordre de troncature P est quant à lui relié à la densité de la variable aléatoire que l'on cherche à représenter. Néanmoins, plusieurs travaux préconisent un ordre peu élevé, et $P = 2$ ou $P = 3$ est généralement suffisant pour bien représenter en pratique une variable aléatoire X [Bou04]. Le second point est capital, car un mauvais choix de distributions

peut engendrer d'importantes erreurs d'estimation. En analyse de sensibilité, ce point est d'autant plus important que la densité des facteurs influence fortement la sensibilité que nous cherchons à évaluer. Pour un grand nombre de distributions paramétriques usuelles, nous connaissons précisément leur famille de polynômes orthogonaux associée [DV07, Bla09]. Par conséquent, une connaissance-ou une estimation- *a priori* de la distribution permettra une bonne représentation en polynômes de chaos de toute variable aléatoire. Nous renvoyons notamment aux travaux de DaVeiga [DV07] pour de telles considérations. Enfin, plusieurs méthodes d'estimation existent pour estimer les coefficients $c_{i_1, \dots, i_k}^\alpha$: méthodes intrusives [GS03], méthodes de projection [AS65] ou de régression [Sud08]. Dans la suite, nous nous focaliserons sur les méthodes de régression. Après avoir fait le lien entre les indices de Sobol et la représentation en polynômes de chaos, nous expliquerons comment la méthode de régression parvient à fournir une bonne estimation des coefficients.

Indices de Sobol et polynômes de chaos

Revenons désormais à notre modèle

$$Y = \eta(\mathbf{X}) = \eta(X_1, \dots, X_p).$$

Puisque $Y \in L^2(\Omega)$, et que $X_1, \dots, X_p$ sont des variables aléatoires indépendantes, nous pouvons représenter Y par une décomposition en polynômes de chaos, dépendant des p variables,

$$Y \simeq \sum_{k \in \mathbb{N}} \sum_{\alpha \in \Lambda_k^p} b_\alpha H_\alpha^k(X_1, \dots, X_p), \quad (1.20)$$

où $H_\alpha^k(X_1, \dots, X_p) \in \Gamma_k(X_1, \dots, X_p)$. De plus, nous supposons que

$$H_\alpha^k(X_1, \dots, X_p) = \prod_{j=1}^p H_{\alpha_j}(X_j), \quad \alpha = (\alpha_1, \dots, \alpha_p) \in \mathbb{N}^p.$$

où $(H_{\alpha_j})_{j \in [1:p], \alpha_j \in \mathbb{N}}$ est une famille unidimensionnelle de polynômes orthogonaux, avec $H_0 = 1$.

En réordonnant les termes de (1.20) selon les paramètres dont ils dépendent,

on obtient ainsi la décomposition suivante,

$$\begin{aligned}
 Y \simeq & b_0 + \sum_{i=1}^p \left(\sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^i} b_\alpha H_\alpha^k(X_i) \right) + \sum_{1 \leq i < j \leq p} \left(\sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^{ij}} b_\alpha H_\alpha^k(X_i, X_j) \right) \\
 & + \cdots + \left(\sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^p} b_\alpha H_\alpha^k(\mathbf{X}) \right),
 \end{aligned} \tag{1.21}$$

avec

$$\Lambda_k^u = \begin{cases} 0 & \text{si } u = \emptyset \\ \{\alpha \in \Lambda_k^p, \alpha_j = 0 \text{ si } j \notin u\} & \text{si } u \in S \setminus \{\emptyset, [1 : p]\} \\ \Lambda_k^p & \text{si } u = [1 : p]. \end{cases}$$

Nous pouvons faire un parallèle avec la décomposition de Hoeffding-Sobol (1.1). Nous savons que chaque composante de la décomposition (1.1) est unique, et que tous les termes sont mutuellement orthogonaux. Ici, comme les polynômes sont orthonormés, nous pouvons poser, pour tout $u \in S$,

$$\begin{aligned}
 \eta_\emptyset &= b_0 \\
 \eta_u(\mathbf{X}_u) &= \sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^u} b_\alpha H_\alpha^k(\mathbf{X}_u).
 \end{aligned}$$

Nous pouvons en déduire que chaque terme de la décomposition de Hoeffding-Sobol admet une expression en termes de polynômes de chaos. Et, par orthogonalité des polynômes, l'expression des indices de Sobol se résume à un rapport de coefficients,

$$S_u^{\text{PC}} = \frac{\sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^u} b_\alpha^2}{\sum_{u \in S \setminus \emptyset} \sum_{k \in \mathbb{N}^*} \sum_{\alpha \in \Lambda_k^u} b_\alpha^2}, \quad u \in S \setminus \{\emptyset\}.$$

La décomposition par polynômes de chaos, tout comme celle par bases de Fourier, constituent des méthodes attractives pour l'analyse de sensibilité. De manière plus générale, les approches spectrales permettent d'obtenir une décomposition fonctionnelle de η , et, par passage à la norme, permettent d'exprimer la variance comme la somme des coefficients liés à la base spectrale choisie.

Estimation des indices par régression

Comme évoqué en introduction de cette section, il existe plusieurs méthodes d'estimation des coefficients, notamment revues dans [Bla09]. Nous nous focalisons ici sur la technique de régression, employée entre autre par [Sud08,

[CLMM09](#)]. Cette méthode est très classique en statistique, mais est rappelée ici, car elle constitue une des technique d'estimation employée dans nos travaux contributifs. Tout d'abord, étant donné que la sommation infinie en k des coefficients n'est pas utilisable en pratique, nous considérons une troncature des sommes induites dans les indices à un ordre $P \in \mathbb{N}^*$, soit

$$S_u^{\text{PC}} \simeq \frac{\sum_{k \leq P} \sum_{\alpha \in \Lambda_k^u} b_\alpha^2}{\sum_{u \in S \setminus \emptyset} \sum_{k \leq P} \sum_{\alpha \in \Lambda_k^u} b_\alpha^2}, \quad u \in S \setminus \{\emptyset\}.$$

Soit $(\mathbf{x}^1, \dots, \mathbf{x}^n) \in \mathcal{M}_{n,p}(\mathbb{R})$, un n -échantillon d'observations issues de la distribution des entrées $\mathbf{X}$, et $\mathbf{y} = (y^1, \dots, y^n)$ l'ensemble des sorties du modèle calculées en les points $(\mathbf{x}^1, \dots, \mathbf{x}^n)$. Notons ici l'importance attribuée au plan d'expérience, qui est à spécifier dans ce contexte. Dans [\[Sud08\]](#), l'auteur relate une construction basée sur les racines de polynômes par points de quadrature.

L'estimation des coefficients de la décomposition en polynômes de chaos se déduit par moindres carrés comme

$$(\hat{b}_\alpha)_{\substack{\alpha \in \Lambda_k^u \\ k \leq P, u \in S}} = \underset{\substack{b_\alpha \in \mathbb{R} \\ \alpha \in \Lambda_k^u \\ k \leq P, u \in S}}{\text{Argmin}} \sum_{l=1}^n \left[y^l - \sum_{u \in S} \sum_{k \leq P} \sum_{\alpha \in \Lambda_k^u} b_\alpha H_\alpha^k(\mathbf{x}_u^l) \right]^2. \quad (1.22)$$

Si le problème est bien posé, soit $n > p$, nous retrouvons l'estimation des moindres carrés ordinaires. En posant $\mathbf{H}$ la matrice composée des $H_\alpha^k(\mathbf{x}^l)$, $l \in [1 : n]$, pour $\alpha \in \Lambda^M := \bigcup_{\substack{k \leq P \\ u \in S \setminus \{\emptyset\}}} \Lambda_k^u$, et $\mathbf{B} = (b_\alpha)_{\alpha \in \Lambda^M}$, la solution de (1.22) est donnée par

$$\hat{\mathbf{B}} = ({}^t \mathbf{H} \mathbf{H})^{-1} \cdot {}^t \mathbf{H} \mathbf{y}. \quad (1.23)$$

Si le problème est mal posé (soit $n < p$), une régression pénalisée peut alors être utilisée afin de procéder à une sélection plus ou moins drastique des variables dans le modèle [\[Bla09\]](#). La régression pénalisée fait l'objet de la Partie [II](#) de cette introduction.

Note 1.4. Comme relaté en premier lieu, le choix des polynômes dans la représentation en polynômes de chaos est crucial, car il reflète la distribution des variables d'entrées, sous-jacente à la quantification de sensibilité dans un modèle. Une bonne appréhension de la distribution est donc nécessaire. Il est aussi à noter que la décomposition en polynômes de chaos se rapproche de l'utilisation d'un méta-modèle pour substituer le modèle initial η , lorsque celui-ci est trop coûteux à calculer en pratique. En effet, au vu de (1.20), on peut considérer la régression linéaire suivante,

$$\mathbb{E}(Y|\beta) = f(\mathbf{X})\beta.$$

Et, en utilisant la formule de l'indice de Sobol (1.9), nous pouvons en déduire une estimation de la sensibilité plus économique numériquement. Pour gérer les erreurs d'observations induites par ce type d'approche, les processus gaussiens constituent une autre classe de méta-modèles. Entièrement caractérisés par leur moyenne et leur noyau de covariance, ils permettent un contrôle sur l'erreur de prédiction du méta-modèle, ainsi que sur l'erreur d'estimation des indices de sensibilité [OO04, MILR09]. Plus généralement, les RKHS [Aro09] permettent, sous un choix approprié de noyaux, de proposer une décomposition type ANOVA dans un cas indépendant [DGRC13].

Cependant, l'hypothèse d'indépendance des entrées est sous-jacente aux méthodes décrites ici. D'une part, les propriétés entraînant la décomposition de Hoeffding-Sobol ne sont permises que pour des distributions produit. L'expression analytique des composantes sous forme d'espérance conditionnelle n'est admissible que lorsqu'il est possible d'utiliser le théorème de Fubini. En outre, la construction théorique des indices de sensibilité, donnant lieu à un indice physiquement interprétable, repose sur cette hypothèse d'indépendance. D'autre part, les méthodes d'estimation explicitées ici ne prennent pas en compte la liaison qui peut exister entre les variables.

Pour des systèmes aléatoires décrivant des phénomènes physiques, biologiques ou économiques, il est parfois peu réaliste de supposer que les paramètres d'entrée sont uniquement caractérisés par leur distribution marginale. Dans l'étude de sensibilité de ces modèles, il est donc nécessaire de développer des mesures de sensibilité et/ou des techniques numériques pour prendre en compte la dépendance des facteurs d'entrée. Néanmoins, la dépendance entre des variables est une notion complexe, sur laquelle il existe une très importante littérature. Il semble donc judicieux, pour l'étude d'un système donné, de comprendre les dépendances mises en jeu pour être capable de mieux les appréhender. Dans le Chapitre 2, nous introduisons les différents concepts de dépendance usuels. En particulier, nous définissons et illustrons les copules pour définir la structure de dépendance d'un système. Le Chapitre 2 présente aussi des outils pratiques qui répondent aux besoins d'un utilisateur, c'est-à-dire qui permettent d'identifier les relations existantes entre les facteurs.

Chapitre 2

Dépendance en Statistique

Si l'on pense souvent corrélation, c'est-à-dire dépendance linéaire, lorsque l'on aborde la notion de dépendance dans un modèle, c'est parce que le coefficient de corrélation linéaire est ce qui a été défini en premier par Galton [Gal88], mais aussi et surtout parce que c'est la notion la plus simple à aborder. Cependant, la dépendance, qui n'est pas à confondre ici avec la notion de fonction de variables indépendantes, terminologie parfois utilisée en statistique, est un concept très riche, que nous allons développer ici. Nous privilégions ici les aspects pratiques de la dépendance. C'est pourquoi nous nous focaliserons sur les fonctions copule, qui non seulement représentent un outil puissant pour définir la structure de dépendance, mais permet aussi une approche pratique.

La Section 2.1 définit la notion de copules, et montre le côté attrayant de cet outil théorique. Nous donnons aussi quelques exemples illustrés de copules, en indiquant, pour chacune d'elle, leur intérêt pratique. Les Sections 2.2-2.4 ont des orientations pratiques. À travers un exemple simple, nous abordons plusieurs techniques pour gérer la dépendance d'un jeu de données en pratique.

2.1 Théorie des copules

Le mot copule, qui signifie littéralement "union, lien" (*Robert de la langue française*) a été employé pour la première fois en Mathématiques par Sklar, dans un théorème qui porte aujourd'hui son nom [Sk171], et qui est à l'origine de la théorie des copules en statistique. Dans ses travaux, l'idée de Sklar est bien d'unir plusieurs fonctions de distribution unidimensionnelles afin d'obtenir une fonction de distribution multivariée. Plus formellement, la fonction de copule permet de décrire la structure de dépendance entre plusieurs variables, puisqu'elle fait le lien entre la distribution jointe de ces variables avec l'ensemble de ses distributions marginales. Les copules connaissent aujourd'hui un intérêt croissant pour des applications en finance [KR08], en

géophysique [GF07], en mécanique [CS10], mais aussi en imagerie [Pou11]. Un tel engouement s'explique par diverses raisons :

1. Les copules présentent un bon moyen de traiter la dépendance dans un cadre très général ;
2. Elles ouvrent un gigantesque champ à la construction de distributions multivariées, avec notamment la possibilité de construire des structures spécifiques pour un problème donné ;
3. Les copules, comme nous le verrons, sont simples à comprendre, et offrent ainsi une utilisation facile en pratique.

Dans la suite, nous définissons formellement les fonctions copule, et en donnons quelques propriétés remarquables connues [Nel06, Leb13]. Nous donnons également plusieurs exemples illustrés de copules.

2.1.1 Définitions et propriétés

Bien qu'il soit usuel de présenter les copules bidimensionnelles, nous présentons ici la généralisation des copules à p dimensions, pour $p \geq 2$. Pour cela, nous définissons tout d'abord certaines notations.

Soit $\bar{\mathbb{R}} := \mathbb{R} \cup \{-\infty, +\infty\}$ la fermeture de l'espace des réels, et $\bar{\mathbb{R}}^p = \bar{\mathbb{R}} \times \cdots \times \bar{\mathbb{R}}$. Pour tout $\mathbf{x} = (x_1, \dots, x_p), \mathbf{y} = (y_1, \dots, y_p) \in \bar{\mathbb{R}}^p$, $\mathbf{x} \leq \mathbf{y}$ si $x_i \leq y_i$, $\forall i \in [1 : p]$. Notons par $[\mathbf{x}, \mathbf{y}]$ le produit cartésien $[x_1, y_1] \times \cdots \times [x_p, y_p]$ tel que $\mathbf{x} \leq \mathbf{y}$. Enfin, pour une fonction $K : \mathbb{R}^p \rightarrow \mathbb{R}$, le K-volume d'un espace $[\mathbf{x}, \mathbf{y}]$ se définit comme

$$V_K([\mathbf{x}, \mathbf{y}]) = \Delta_{x_p}^{y_p} \cdots \Delta_{x_1}^{y_1} K(\mathbf{t}),$$

où $\Delta_{x_i}^{y_i}$ est la i ème différence finie de K , c'est-à-dire ,

$$\Delta_{x_i}^{y_i} K(\mathbf{t}) = K(t_1, \dots, t_{i-1}, y_i, t_{i+1}, \dots, t_p) - K(t_1, \dots, t_{i-1}, x_i, t_{i+1}, \dots, t_p).$$

Définissons maintenant une fonction copule.

Définition 2.1. Une copule C p -dimensionnelle est une fonction vérifiant les propriétés suivantes,

1. C est une fonction définie sur $[0, 1]^p$ à valeurs dans $[0, 1]$;
2. Pour tout $\mathbf{u} \in [0, 1]^p$, $C(\mathbf{u}) = 0$ s'il existe au moins un $i \in [1 : p]$ avec $u_i = 0$;
3. Pour tout $i \in [1 : p]$ et tout $u_i \in [0, 1]$, $C(1, \dots, 1, u_i, 1, \dots, 1) = u_i$;
4. Pour tout $\mathbf{u}, \mathbf{v} \in [0, 1]^p$ tel que $\mathbf{u} \leq \mathbf{v}$,

$$V_C([\mathbf{u}, \mathbf{v}]) \geq 0.$$

Note 2.1. La condition définie au point 4 admet une écriture simple en petite dimension. Si $p = 2$, la condition 4 est donnée par, pour tout $(u_1, u_2), (v_1, v_2) \in [0, 1]^2$, avec $u_i \leq v_i, i = 1, 2$,

$$C(v_1, v_2) - C(v_1, u_2) - C(u_1, v_2) + C(u_1, u_2) \geq 0.$$

Pour bien comprendre ce que sont les copules, plaçons nous dans le cas simple $p = 2$, et dans l'espace $L^2_{\mathbb{R}}$ des variables aléatoires de carrés intégrables.

Soit maintenant un vecteur aléatoire $U = (U_1, U_2)$ tel que chacune des marginales suive une loi uniforme, soit $U_i \sim \mathcal{U}[0, 1]$, $i = 1, 2$. Considérons la fonction de répartition de U , soit $F(u_1, u_2) = P(U_1 \leq u_1, U_2 \leq u_2)$. Il est alors aisé de montrer que F est une copule :

1. Le support de U_i étant $[0, 1]$, on a bien $F : [0, 1]^2 \rightarrow [0, 1]$;
2. $\forall u_1, u_2 \in [0, 1], F(u_1, 0) = P(U_1 \leq u_1, U_2 \leq 0) = F(0, u_2) = P(U_1 \leq 0, U_2 \leq u_2) = 0$;
3. $\forall u_1, u_2 \in [0, 1] F(u_1, 1) = P(U_1 \leq u_1) = u_1$ et $F(1, u_2) = P(U_2 \leq u_2) = u_2$;
4. $F(v_1, v_2) - F(v_1, u_2) - F(u_1, v_2) + F(u_1, u_2) = P(u_1 \leq U_1 \leq v_1, u_2 \leq U_2 \leq v_2) \geq 0$ pour $u_i \leq v_i, i = 1, 2$.

Ainsi, une autre définition des copules peut être apportée, qui est énoncée dans la Définition 2.2,

Définition 2.2. Une copule C est une fonction de répartition dont les marges sont uniformes sur $[0, 1]$.

De plus, lorsque $X_i \sim F_i, i = 1, 2$, avec F_i une fonction de répartition continue, il est bien connu que $U_i = F_i(X_i)$ définit une variable uniforme [Dev86] (cf Lemme 2.1). Il est alors évident que $F(F_1(x_1), F_2(x_2))$ définit une probabilité bidimensionnelle dont les marges sont F_1 et F_2 . De plus,

$$F(F_1(x_1), F_2(\infty)) = F(F_1(x_1), 1) = F_1(x_1).$$

Les copules sont donc un outil très puissant pour construire des distributions multidimensionnelles dont les marges sont données. Néanmoins, Sklar a avancé un résultat encore plus intéressant, énoncé dans le Théorème 2.1 [Sk171].

Théorème 2.1. (Sklar) Soit F une fonction de répartition p -dimensionnelle dont les marges sont $F_1, \dots, F_p$. Alors il existe un copule C p -dimensionnelle telle que, pour tout $\mathbf{x} \in \mathbb{R}^p$,

$$F(x_1, \dots, x_p) = C(F_1(x_1), \dots, F_p(x_p)). \quad (2.1)$$

De plus, si $F_1, \dots, F_p$ sont des fonctions continues, la copule C est unique. Sinon, C est uniquement déterminée sur $F_1(\mathbb{R}) \times \dots \times F_p(\mathbb{R})$.

Réciproquement, si C est une copule p -dimensionnelle et $F_1, \dots, F_p$ sont des fonctions de répartition unidimensionnelles, alors F définie par (2.1) est une fonction de répartition p -dimensionnelle dont les marges sont $F_1, \dots, F_p$.

Plusieurs corollaires s'ensuivent naturellement.

Corollaire 2.1. Soient $F, F_1, \dots, F_p$ et C les fonctions définies au Théorème 2.1. Si les marges $F_1, \dots, F_p$ sont strictement croissantes, alors, pour tout $\mathbf{u} \in [0, 1]^p$,

$$C(u_1, \dots, u_p) = F(F_1^{-1}(u_1), \dots, F_p^{-1}(u_p)).$$

Corollaire 2.2. Soient $F, F_1, \dots, F_p$ et C les fonctions définies au Théorème 2.1. Si F admet une densité $p_{\mathbf{X}}$, alors, pour tout $\mathbf{x} \in \mathbb{R}^p$,

$$p_{\mathbf{X}}(x_1, \dots, x_p) = c(F_1(x_1), \dots, F_p(x_p)) \prod_{i=1}^p p_{X_i}(x_i), \quad (2.2)$$

où p_{X_i} , $i \in [1 : p]$ sont les densités marginales de $p_{\mathbf{X}}$, et c est la densité de copule de C définie comme, pour tout $\mathbf{u} \in [0, 1]^p$,

$$c(u_1, \dots, u_p) = \frac{\partial^p C}{\partial u_1 \dots \partial u_p}(u_1, \dots, u_p). \quad (2.3)$$

De plus,

- Pour tout $\mathbf{u} \in [0, 1]^p$, $c(\mathbf{u}) \geq 0$ par la condition 4 de la Définition 2.1 ;
- La densité de copule c est unique si $F_1, \dots, F_p$ sont des fonctions continues. Sinon, c est uniquement déterminée sur $F_1(\mathbb{R}) \times \dots \times F_p(\mathbb{R})$.

Par conséquent, si $X_1, \dots, X_p$ sont des variables aléatoires de densité jointe $p_{\mathbf{X}}$, et de marginales $p_{X_1}, \dots, p_{X_p}$, la copule permet donc de lier $p_{\mathbf{X}}$ au produit de densités $p_{X_1} \times \dots \times p_{X_p}$, c'est-à-dire la densité lorsque les variables sont supposées indépendantes. Ainsi, la densité de la copule indépendante est égale à 1, et la copule indépendante Π est donnée par la formule suivante,

$$\Pi(\mathbf{u}) = \prod_{i=1}^p u_i.$$

Les bornes de Fréchet sont aussi des fonctions caractéristiques : en effet, toute copule est bornée par les bornes de Fréchet. Ainsi, tout graphique de copules se représente entre les graphes des bornes de Fréchet-Hoeffding. Dans le cas $p = 2$, les bornes de Fréchet sont des copules. La borne supérieure représente la co-monotonie, c'est-à-dire une dépendance parfaitement positive, tandis que la borne inférieure correspond à l'anti-monotonie, soit une dépendance

parfaitement négative. Dans la suite, nous présentons d'autres copules qui se placent entre ces bornes. Même si le type de dépendance qu'elles représentent est moins évident que pour ces bornes, leur représentation graphique donne un aperçu de leur caractéristique.

2.1.2 Exemples de copules

Cette section a pour objectif de donner quelques illustrations de fonctions copule. Pour plus de clarté, nous nous placerons souvent dans le cas de copules bidimensionnelles.

En général, les copules sont groupées par familles dont les composantes présentent des propriétés de construction similaires. Ici, nous décidons de présenter la famille de copules elliptiques, car elle contient la copule gaussienne, qui, comme nous le verrons au Chapitre 3, est la plus utilisée pour traiter la dépendance en analyse de sensibilité. Nous présentons aussi quelques copules de la famille des copules archimédiennes, car elles regroupent un très grand nombre de copules étudiées. De plus, ces copules seront mentionnées dans le Chapitre 6 de la Partie III de ce manuscrit.

Copules elliptiques

Les copules elliptiques généralisent la distribution gaussienne, mais regroupent également des distributions utilisées pour les tests qui portent leur nom (t-distribution, Pearson, ...) [FKN90]. Rappelons d'abord ce qu'est une distribution elliptique.

Définition 2.3. *Un vecteur aléatoire $\mathbf{X}$ à valeurs dans $\mathbb{R}^p$ admet une distribution elliptique de paramètres $\mu \in \mathbb{R}^p$ et $\Sigma \in \mathcal{M}_{p,p}(\mathbb{R})$, symétrique et définie positive, si $\mathbf{X}$ admet la représentation stochastique suivante :*

$$\mathbf{X} \stackrel{\mathcal{L}}{=} \mu + R\mathbf{A}\mathbf{U},$$

où $\mathbf{A} \in \mathcal{M}_{p,k}(\mathbb{R})$ tel que $k = \text{Ran}\mathbf{X} \leq p$, et $\mathbf{A}^t\mathbf{A} = \Sigma$. $R \geq 0$ est une variable aléatoire et $\mathbf{U}$ est un vecteur aléatoire uniformément distribué sur

$$\Omega = \{\mathbf{x} \in \mathbb{R}^p, \sum_{i=1}^p x_i^2 = 1\}, \text{ et indépendant de } R.$$

Par conséquent, la fonction caractéristique $\psi(\mathbf{t})$ de $\mathbf{X}$ est de la forme,

$$\psi(\mathbf{x}) = \exp(i^t\mathbf{x}\mu) \cdot \varphi(^t\mathbf{x}\Sigma\mathbf{x}),$$

où φ est une fonction scalaire. Nous notons $\mathbf{X} \sim \mathcal{E}(\mu, \Sigma, \varphi)$.

De plus, comme la distribution de $\mathbf{X}$ est supposée continue, alors $\mathbf{X}$ admet une densité qui s'exprime comme,

$$p_{\mathbf{X}}(\mathbf{x}) = |\Sigma|^{-1/2} g(t(\mathbf{x} - \mu)\Sigma^{-1}(\mathbf{x} - \mu)),$$

où $g : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ et telle que $\int_0^{+\infty} y^{p/2-1} g(y) dy < \infty$, avec g relatif à la densité de $R \geq 0$. Dans ce cas, on peut noter $\mathbf{X} \sim \mathcal{E}(\mu, \Sigma, g)$.

Les distributions elliptiques ont les mêmes propriétés que la distribution gaussienne. En effet, les distributions marginales d'une distribution elliptiques sont elliptiques. Les distributions conditionnelles restent également elliptiques.

Pour définir la copule elliptique, nous considérons la distribution elliptique $\mathcal{E}(0, \Sigma, g)$, dont les marges sont identiques. Notons F_X la fonction de répartition marginale, et p_X la densité marginale elliptique de X à valeurs dans $\mathbb{R}$. La copule elliptique associée à $\mathcal{E}(0, \Sigma, g)$ se définit comme (Corollaire 2.1), pour tout $\mathbf{u} \in [0, 1]^p$,

$$C(u_1, \dots, u_p) = \frac{1}{|\Sigma|^{1/2}} \int_{-\infty}^{F_X^{-1}(u_1)} \dots \int_{-\infty}^{F_X^{-1}(u_p)} g(t(\mathbf{x} - \mu)\Sigma^{-1}(\mathbf{x} - \mu)) d\mathbf{x}.$$

La densité de copule c associée à C est de la forme

$$c(u_1, \dots, u_p) = \frac{p_{\mathbf{X}}(F_X^{-1}(u_1), \dots, F_X^{-1}(u_p))}{p_X(F_X^{-1}(u_1)) \times \dots \times p_X(F_X^{-1}(u_p))}.$$

Note 2.2. Dans le cas où $p = 2$, et $\Sigma = \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}$, $-1 < \rho < 1$, les fonctions de répartition et de densité associées à X_1 et X_2 sont données par [FFK02]

$$p_X(x) = \int_{x^2}^{+\infty} (y - x^2)^{-1/2} g(y) dy$$

$$F_X(x) = \frac{1}{2} + \int_{x^2}^{+\infty} \arcsin\left(\frac{x}{\sqrt{y}}\right) g(y) dy$$

Exemple 1 La distribution gaussienne est une distribution elliptique. Si $\mathbf{X} = (X_1, \dots, X_p)$ admet une distribution gaussienne $N(0, \Sigma)$, alors $\mathbf{X}$ est de distribution elliptique, et $g(\mathbf{x}) = \frac{1}{(2\pi)^{p/2}} \exp(-\mathbf{x}/2)$, $\mathbf{x} \in \mathbb{R}^p$. De plus, si

$$\Sigma_{ij} = \begin{cases} 1 & \text{si } i = j \\ \rho_{ij} & \text{sinon,} \end{cases}$$

alors la densité de copule gaussienne admet une expression analytique, donnée par

$$c(u_1, \dots, u_p) = \frac{1}{|\Sigma|^{1/2}} \exp \left(-\frac{1}{2} {}^t \mathbf{w} (\Sigma^{-1} - \mathbf{I}) \mathbf{w} \right), \quad \forall \mathbf{u} \in [0, 1]^p,$$

avec $\mathbf{w} = (\phi^{-1}(u_1), \dots, \phi^{-1}(u_p))$ et ϕ sont les fonctions de répartition marginales identiques de la fonction de répartition de $\mathbf{X}$.

La copule gaussienne est une copule très flexible, puisqu'elle peut tout aussi bien caractériser la dépendance positive, que la dépendance négative. La copule gaussienne a aussi pour particularité d'admettre une indépendance de queue, c'est-à-dire que si l'on s'éloigne suffisamment de la moyenne de la distribution, alors les marges agissent indépendamment les unes des autres [EMS02]. Précisons que la dépendance de queue, notion très utilisée en Finance, permet entre autre de décrire la probabilité limite qu'une marge excède un certain seuil sachant que les autres ont déjà dépassé ce même seuil [CS09]. Cela permet donc de décrire la dépendance au niveau des queues de la distribution, c'est-à-dire la distribution d'événements rares. Nous renvoyons le lecteur à [CS09, CS07, Sch02, TZ05] pour un approfondissement de cette notion, que nous ne détaillerons pas ici.

Exemple 2 La t -distribution multivariée (ou de Student) est une distribution elliptique. En effet, si $\mathbf{X}$ est une t -distribution de paramètres μ et Σ , la densité de $\mathbf{X}$ s'écrit comme,

$$p_{\mathbf{X}}(\mathbf{x}) = \frac{\Gamma(\frac{p+m}{2})}{(\pi m)^{p/2} \Gamma(\frac{m}{2})} |\Sigma|^{-1/2} \left[1 + \frac{1}{m} {}^t (\mathbf{x} - \mu) \Sigma^{-1} (\mathbf{x} - \mu) \right]^{-\frac{p+m}{2}},$$

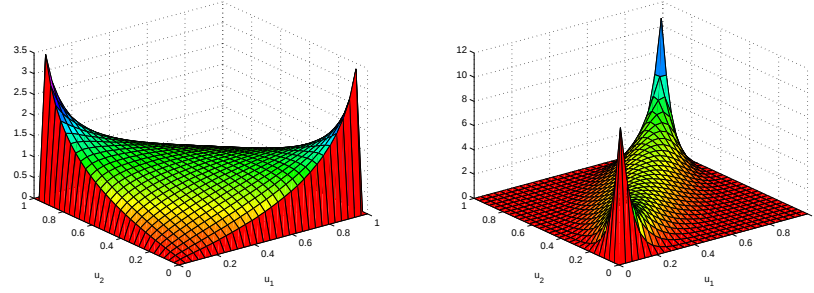
où Γ est la fonction Gamma, et $m > 0$. Nous notons $\mathbf{X} \sim t_m(\mu, \Sigma)$. La copule de Student est la copule associée à $t_m(0, \Sigma)$, avec $\Sigma_{ii} = 1$ pour tout $i \in [1 : p]$. Notons ϕ la fonction de répartition marginale de Student. L'expression de la copule de Student est donnée par, pour tout $\mathbf{u} \in [0, 1]^p$,

$$C(\mathbf{u}) = \frac{1}{|\Sigma|^{1/2}} \frac{\Gamma(\frac{p+m}{2})}{(\pi m)^{p/2} \Gamma(\frac{m}{2})} \int_{-\infty}^{\phi^{-1}(u_1)} \cdots \int_{-\infty}^{\phi^{-1}(u_p)} \left[1 + \frac{1}{m} {}^t \mathbf{x} \Sigma^{-1} \mathbf{x} \right]^{-\frac{p+m}{2}} d\mathbf{x}.$$

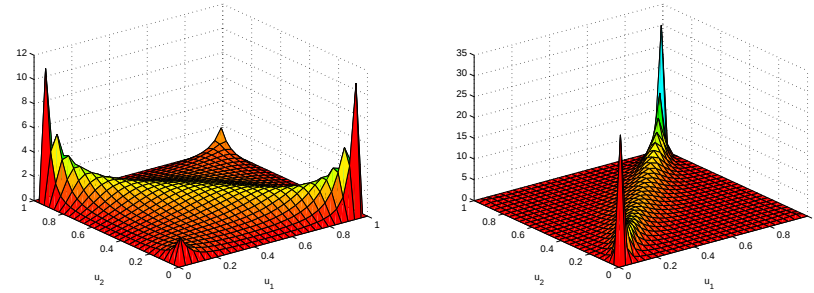
La densité de copule c associée est définie comme

$$c(\mathbf{u}) = \frac{\Gamma(\frac{p+m}{2}) \Gamma(\frac{m}{2})^{p-1}}{|\Sigma|^{1/2} \Gamma(\frac{1+m}{2})^p} \frac{\left[1 + \frac{1}{m} {}^t \mathbf{w} \Sigma^{-1} \mathbf{w} \right]^{-\frac{p+m}{2}}}{\prod_{i=1}^p \left[1 + \frac{1}{m} w_i^2 \right]^{-\frac{1+m}{2}}},$$

avec $\mathbf{w} = (\phi^{-1}(u_1), \dots, \phi^{-1}(u_p))$.



(a) Copule gaussienne avec $\rho = -0.5$ (b) Copule gaussienne avec $\rho = 0.9$



(c) Copule t-distribution avec $\rho = -0.5$ et $m = 1$

(d) Copule t-distribution avec $\rho = 0.9$ et $m = 1$

FIGURE 2.1 – Densité de copules elliptiques pour différents coefficients de corrélation

La copule de Student, de même que la copule gaussienne, inclut des dépendances positives ou négatives. Mais elle se distingue de la copule gaussienne par une dépendance de queue à droite (ou supérieure), quelque soit le coefficient de corrélation $\rho > -1$. Cette dépendance tend à augmenter lorsque le degré de liberté $m > 0$ diminue et le coefficient de corrélation ρ augmente. Cette différence est notable sur la Figure 2.1, où les copules gaussienne et de Student sont représentées pour différentes valeurs du coefficient de corrélation ρ .

Copules archimédiennes

Cette classe de copules regroupe un très grand nombre de copules connues dans la littérature. Elles sont spécifiées par une fonction φ . La définition d'une copule archimédienne est donnée en Définition 2.1.

Définition 2.1. Soit une fonction φ telle que :

1. $\varphi : [0, 1] \rightarrow \bar{\mathbb{R}}^+$;

2. $\varphi(1) = 0$;
 3. φ est continue et strictement décroissante, soit $\varphi'(u) < 0$, $\forall u \in [0, 1]$ si φ est dérivable ;
 4. φ est convexe, soit $\varphi''(u) > 0$, $\forall u \in [0, 1]$ si φ est deux fois dérivable ;
- Alors nous pouvons définir la fonction pseudo-inverse $\varphi^{[-1]}$ de φ , donnée par

$$\varphi^{[-1]}(t) = \begin{cases} \varphi^{-1}(t), & 0 \leq t \leq \varphi(0) \\ 0 & \varphi(0) \leq t \leq +\infty. \end{cases}$$

Dans ce cas, nous pouvons définir la copule archimédienne bidimensionnelle comme,

$$C(u_1, u_2) = \varphi^{[-1]}[\varphi(u_1) + \varphi(u_2)].$$

De plus, si $\varphi^{[-1]}$ est complètement monotone sur $\mathbb{R}^+$, c'est-à-dire $\forall t \in \mathbb{R}^+$, $\forall k \in \mathbb{N}$, $(-1)^k \frac{d^k}{dt^k} \varphi^{[-1]}(t) \geq 0$, alors $\varphi^{[-1]} = \varphi^{-1}$.

Pour définir la copule archimédienne p -dimensionnelle, nous supposons que φ vérifie 1.-4., et qu'elle est de plus complètement monotone. Ainsi,

$$C(u_1, \dots, u_p) = \varphi^{-1}[\varphi(u_1) + \dots + \varphi(u_p)], \quad p \geq 3. \quad (2.4)$$

Dans la suite, nous ne considérerons que des fonctions copule bivariées.

Outre les copules standards comme la copule de Frank ou la copule de Clayton que nous représentons en Figure 2.2, nous attirons l'attention sur la famille de copules de Ali-Mikhail-Hal (AMH) [AMH78], qui constitue une famille importante pour les apports de cette thèse. Cette famille a pour générateur,

$$\varphi_\theta(u) = \ln \left[\frac{1 - \theta(1 - u)}{u} \right], \quad \theta \in [-1, 1[.$$

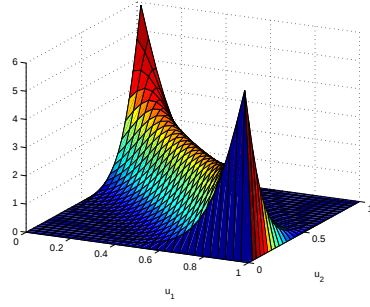
La fonction de copule s'écrit comme

$$C_\theta^{\text{AMH}}(u_1, u_2) = \frac{u_1 u_2}{1 - \theta(1 - u_1)(1 - u_2)}, \quad \theta \in [-1, 1[. \quad (2.5)$$

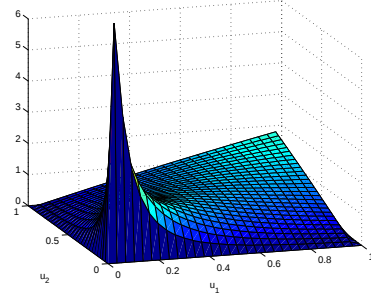
Néanmoins, nous pouvons remarquer que C définie en (2.5) peut s'écrire comme un développement en série [Sk171],

$$C_\theta^{\text{AMH}}(u_1, u_2) = u_1 u_2 \sum_{k \geq 0} [\theta(1 - u_1)(1 - u_2)]^k, \quad \theta \in [-1, 1[. \quad (2.6)$$

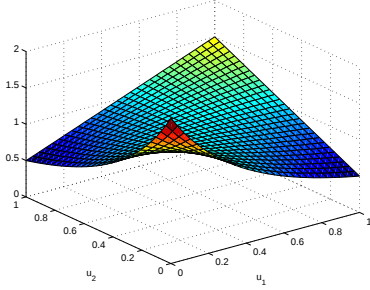
Ainsi, en tronquant C définie en (2.6) à un ordre k petit, on obtient de nouvelles familles de copules. Pour $k = 1$, nous obtenons la copule de Farlie-Gumbel-Morgensten (FGM) [Mor56], dont la densité de copule est donnée par



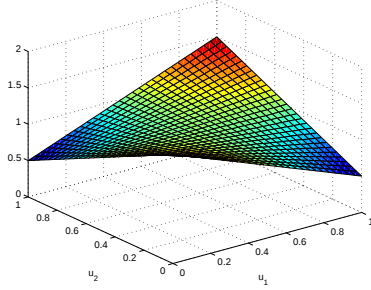
(a) Copule de Frank,
 $\varphi_\theta(t) = -\ln \frac{e^{-\theta t} - 1}{e^{-\theta} - 1}, \theta \in \mathbb{R}$



(b) Copule de Clayton,
 $\varphi(t) = \frac{1}{\theta}(t^{-\theta} - 1), \theta > 0$



(c) Copule AMH



(d) Copule FGM

FIGURE 2.2 – Densité de copules archimédiennes

$$c_\theta^{\text{FGM}}(u_1, u_2) = 1 + \theta(1 - 2u_1)(1 - 2u_2), \quad \theta \in [-1, 1].$$

Si $k = 2$, nous obtenons la copule FGM itérée de Lin, dont nous donnons directement la densité de copule, pour $\theta \in [-1, 1]$,

$$c_\theta^{\text{FGMI}}(u_1, u_2) = 1 + \theta(1 - 2u_1)(1 - 2u_2) + \theta^2(1 - 4u_1 + 3u_1^2)(1 - 4u_2 + 3u_2^2).$$

Notons que les deux approximations polynomiales de la copule AMH ne sont plus des copules archimédiennes. Cependant, elles sont plus populaires que la copule AMH de par leur forme simplifiée. Ces trois types de copules sont des perturbations de la copule indépendante, et pour $\theta = 0$, les copules AMH, FGM et FGMI coïncident avec la copule produit II. Néanmoins, comme nous pouvons le remarquer en Figure 2.2, ces copules sont adaptées à de faibles dépendances.

Note 2.3. Dans cette partie, nous avons voulu mettre en valeur la diversité de forme que peut prendre une fonction copule. De plus, un fait très

important est que la représentation de la dépendance pour un ensemble de variables peut être faite indépendamment du choix des distributions marginales. En effet, rappelons d'abord le lemme suivant [Dev86]

Lemme 2.1. *Soit F une fonction de répartition et F^{-1} sa fonction quantile définie comme*

$$\forall u \in [0, 1], \quad F^{-1}(u) = \inf\{x \in \mathbb{R} \cup \pm\infty : F(x) \geq u\}.$$

Si $U \sim \mathcal{U}[0, 1]$, alors $X = F^{-1}(U)$ est une variable aléatoire de distribution F . Si X a pour distribution F , et F continue, alors $F(X)$ est uniformément distribuée sur $[0, 1]$.

Supposons maintenant qu'on désire obtenir (X, Y) dont les marginales respectives sont G_1 et G_2 , mais dont la dépendance est assurée par une copule C de marges F_1 et F_2 continues.

Soit (Z_1, Z_2) un couple de variables aléatoires de distribution continue F telle que $F(Z_1, Z_2) = C(F_1(Z_1), F_2(Z_2))$. Nous avons $F_1(Z_1), F_2(Z_2)$ uniformément distribuées (Lemme 2.1), et de copule C . Il suffit maintenant de poser $X = G_1^{-1}(F_1(Z_1))$ et $Y = G_2^{-1}(F_2(Z_2))$ pour obtenir X, Y de marges respectives G_1 et G_2 , mais dont la dépendance est caractérisée par C . Ainsi, cette démarche permet d'obtenir un très large choix de modèles beaucoup plus flexibles.

Dans le cas où la copule C est elliptique, de telles distributions sont plus connues sous le nom de *distributions méta-elliptiques* [FFK02, JDHR10].

Dans cette première partie, nous avons évoqué certains des outils existants sur la théorie de la dépendance probabiliste. Dans la suite, nous nous posons des questions pratiques auxquelles un utilisateur voudra répondre face à un problème donné :

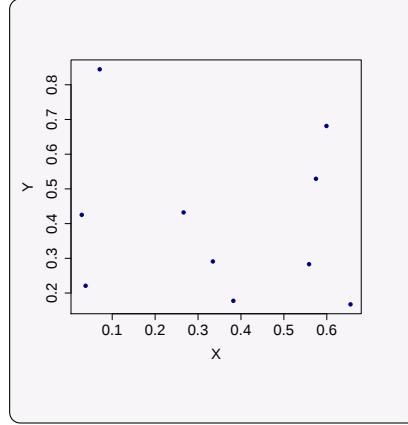
1. Existe-t'il de la dépendance dans les données/le modèle traité?
2. Si oui, avec quels outils peut-on l'identifier?
3. Comment simuler des observations issues d'une distribution dépendante?

Largement inspiré par les travaux de Genest & Favre [GF07], nous tentons de répondre à ces questions dans les Sections 2.2 et 2.3.

Pour cela, nous supposons que $P_{\mathbf{X}}$, la distribution de $\mathbf{X}$, est absolument continue par rapport à la mesure de Lebesgue. Ceci permet d'assurer le caractère diffus de la loi de $\mathbf{X}$, d'assurer la continuité de la fonction de répartition F de $\mathbf{X}$, et l'existence de la fonction de densité $p_{\mathbf{X}}$ de $\mathbf{X}$.

Nous supposons aussi que nous disposons de n observations $(x^i, y^i)_{i \in [1:n]}$ issues de la distribution de deux variables (X, Y) . Nous illustrons également les méthodes présentées avec un échantillon simulé de taille $n = 1000$ issu

de la distribution de la copule de Farlie-Gumbel-Morgenstern, avec une forte dépendance, $\theta = -1$. Cet échantillon a été généré via la commande `rcopula` du package *copula* du logiciel R. Nous donnons ici le graphe de (X, Y) pour quelques réalisations,



2.2 Détection et mesure de la dépendance en pratique

La première question à se poser face à un problème stochastique concerne la dépendance des variables. Peut-on raisonnablement considérer que les variables sont indépendantes ou, au contraire, qu'elles sont fortement dépendantes ? Dans beaucoup de situations, une connaissance *a priori* des experts sur le problème suffit à répondre à la question. Mais une telle conclusion peut être biaisée par la nécessité de simplifier l'étude qui s'ensuivra. Dans ce cas, on aura tendance à supposer les variables indépendantes. Dans d'autres circonstances, la méthode de construction d'un modèle rend l'hypothèse de dépendance inévitable. Néanmoins, on ignore comment et de combien les variables sont dépendantes. Le but de cette section est de fournir les outils pour répondre à cette question.

Coefficient de Pearson

Le premier outil, très classique mais néanmoins utile, est le coefficient de corrélation de Pearson [Pea96] qui permet de mesurer la dépendance linéaire entre deux variables,

$$\rho_n^{\text{Pear}} = \frac{\sum_{i=1}^n (x^i - \bar{x})(y^i - \bar{y})}{\sqrt{\sum_{i=1}^n (x^i - \bar{x})^2 \sum_{i=1}^n (y^i - \bar{y})^2}}, \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x^i, \quad \bar{y} = \frac{1}{n} \sum_{i=1}^n y^i.$$

ρ_n^{Pear} est l'estimateur empirique de la mesure de corrélation usuelle entre deux variables X et Y , donné par

$$\rho^{\text{Pear}} = \frac{\text{Cov}(X, Y)}{\sqrt{V(X)V(Y)}} = \frac{\mathbb{E}[(X - \mathbb{E}(X))(Y - \mathbb{E}(Y))]}{\sqrt{\mathbb{E}[(X - \mathbb{E}(X))^2]\mathbb{E}[(Y - \mathbb{E}(Y))^2]}}.$$

Néanmoins, ce coefficient, sous réserve d'exister (ce qui n'est pas vrai pour une distribution de Cauchy par exemple), est une mesure très limitante de la dépendance : en effet, si $\rho^{\text{Pear}} = 0$, cela signifie qu'il n'y a pas de dépendance linéaire, mais cela n'exclut pas une dépendance plus complexe.

Pour pallier ce problème, il est plus judicieux de s'intéresser aux rangs des observations. En effet, contrairement aux observations elles-mêmes, les rangs sont invariants par transformations croissantes, ce qui les rend plus en adéquation avec la caractérisation des copules, elles aussi invariantes par transformations monotones croissantes [Nel06]. Les rangs d'observations apportent également plus d'information que les observations elles-mêmes [Oak82, GF07]. Pour définir les mesures de dépendance basées sur les rangs, nous introduisons $(r^1, s^1), \dots, (r^n, s^n)$ les paires de rang associées à $(x^1, y^1), \dots, (x^n, y^n)$. Ainsi, si $x^1 < x^2 < \dots < x^n$, le rang r^i associé à x^i sera $r^i = i$.

Le ρ de Spearman

Le coefficient de corrélation sur les rangs, aussi connu sous le nom de rho de Spearman [Kru58], se définit simplement comme le coefficient de Pearson appliqué aux rangs. Ainsi, puisque $\bar{r} = \frac{1}{n} \sum_{i=1}^n r^i = \frac{n+1}{2} = \bar{s}$, nous avons

$$\rho_n^{\text{Spear}} = \frac{12}{n(n+1)(n-1)} \sum_{i=1}^n r^i s^i - 3 \frac{n+1}{n-1}.$$

Si F est la fonction de répartition jointe du couple (X, Y) telle que

$$F(x, y) = C(F_X(x), F_Y(y)),$$

la mesure ρ_n^{Spear} est un estimateur de ρ^{Spear} , qui peut être vu comme le coefficient de Pearson entre $U = F_X(X)$ et $V = F_Y(Y)$. Comme U et V sont uniformes sur $[0, 1]$ (Lemme 2.1), nous avons

$$\mathbb{E}(U) = \mathbb{E}(V) = \frac{1}{2}, \quad V(U) = V(V) = \frac{1}{12}.$$

Ainsi, ρ_n^{Spear} est l'estimateur de ρ^{Spear} , donné par

$$\begin{aligned}
 \rho^{\text{Spear}}(X, Y) = \rho^{\text{Pear}}(U, V) &= \frac{\mathbb{E}(UV) - \mathbb{E}(U)\mathbb{E}(V)}{\sqrt{V(U)V(V)}} \\
 &= 12 \left[\mathbb{E}(UV) - \frac{1}{4} \right] \\
 &= 12 \int_0^1 \int_0^1 uv \, dC(u, v) - 3 \\
 &= 12 \int_0^1 \int_0^1 [C(u, v) - uv] \, dudv.
 \end{aligned} \tag{2.7}$$

Ce coefficient est nettement plus général que celui de Pearson, en ce sens que si $\rho^{\text{Spear}} = \pm 1$, X est *fonctionnellement* dépendant de Y , c'est-à-dire que la copule C est une des bornes de Fréchet-Hoeffding. Si $\rho^{\text{Spear}} = 0$, X et Y sont indépendantes. Cette mesure est à l'origine d'un test permettant de tester la dépendance d'un couple [Ken38]. En effet, sous l'hypothèse nulle $H_0 : C = \Pi$ d'indépendance entre X et Y , la distribution de ρ_n^{Spear} est proche d'une distribution gaussienne de moyenne nulle et de variance $1/(n-1)$. Par conséquent, H_0 est rejetée à un niveau α si $\sqrt{1/(n-1)}|\rho_n^{\text{Spear}}| > z_{\alpha/2}$, où $z_{\alpha/2}$ est le quantile de la loi normale correspondant à $\alpha/2$.

Exemple

Le résultat du test de Spearman sur nos données simulées est le suivant, en utilisant la fonction `cor.test` du logiciel R.

```

> cor.test(X,Y,method="spearman")

Spearman's rank correlation rho

data:  X and Y
S = 232505090, p-value < 2.2e-16
alternative hypothesis: true rho is not equal to 0
sample estimates:
      rho
-0.3950319

```

Le τ de Kendall

Le tau de Kendall [Ken38] est un indice également basé sur les rangs, mais construit sur la concordance des observations. Il est aussi très utilisé pour mesurer la dépendance. Il se définit comme suit,

$$\tau_n = \frac{P_n - Q_n}{\binom{n}{2}} = \frac{4}{n(n-1)}P_n - 1,$$

où P_n est le nombre de paires concordantes, et Q_n le nombre de paires discordantes de l'échantillon. Pour $i \neq j$, (x^i, y^i) et (x^j, y^j) sont dites concordantes si $(x^i - x^j)(y^i - y^j) \geq 0$; elles sont dites discordantes sinon. De même que ρ^{Spear} , $\tau \in [-1, 1]$, et $\tau = 0$ si X et Y sont indépendantes. Cet estimateur est l'estimation d'une mesure d'association théorique τ , que l'on peut aussi exprimer comme suit. Soient (X, Y) et (X', Y') deux vecteurs i.i.d de distribution F , caractérisée par une fonction de copule C . Le tau de Kendall théorique est donné par

$$\begin{aligned}\tau &= P[(X - X')(Y - Y') \geq 0] - P[(X - X')(Y - Y') < 0] \\ &= 2P[(X - X')(Y - Y') \geq 0] - 1.\end{aligned}$$

Or,

$$P[(X - X')(Y - Y') \geq 0] = P(X \geq X', Y \geq Y') + P(X \leq X', Y \leq Y'),$$

et,

$$\begin{aligned}P(X \geq X', Y \geq Y') &= \iint_{\mathbb{R}^2} P(X' \leq x, Y' \leq y) \, dC(F_X(x), F_Y(y)) \\ &= \iint_{\mathbb{R}^2} C(F_X(x), F_Y(y)) \, dC(F_X(x), F_Y(y)).\end{aligned}$$

De même, on montre que $P(X \leq X', Y \leq Y') = P(X \geq X', Y \geq Y')$. Ainsi, par changement de variable, on obtient le τ théorique suivant [Nel06],

$$\tau = 4 \int_0^1 \int_0^1 C(u, v) \, dC(u, v) - 1. \quad (2.8)$$

Cette mesure est à l'origine d'un test permettant de tester la dépendance d'un couple [Ken38], qui est très similaire à celui donné pour le ρ de Spearman. Sous l'hypothèse nulle $H_0 : C = \Pi$ d'indépendance entre X et Y , la distribution de τ_n est proche d'une distribution gaussienne de moyenne nulle et de variance $2(2n+5)/[9n(n-1)]$. Par conséquent, H_0 est rejetée à un niveau α si $\sqrt{9n(n-1)/[2(2n+5)]}|\tau_n| > z_{\alpha/2}$, où $z_{\alpha/2}$ est le quantile de la loi normale correspondant à $\alpha/2$.

Exemple (suite)

Nous illustrons le test de Kendall sur nos données simulées,

```
> cor.test(X,Y,method="kendall")

Kendall's rank correlation tau

data:  X and Y
z = -12.5352, p-value < 2.2e-16
alternative hypothesis: true tau is not equal to 0
sample estimates:
      tau
-0.2647287
```

Au travers de notre illustration, nous constatons que les tests sur le ρ de Spearman et sur le τ de Kendall permettent de conclure que l'hypothèse d'indépendance est rejetée.

Note 2.4. Au vu des formes théoriques des mesures de Kendall et de Spearman, ces mesures peuvent être généralisées à d'autres mesures d'association à l'origine de tests de dépendance. Les travaux de Genest *et al.* [GR04] font état de ces mesures et de leur test associé.

Il est aussi à noter que ces mesures bivariées sont aussi généralisables au cas multivarié. En effet, les formes théoriques du ρ de Spearman (2.2), et du τ de Kendall (2.8) peuvent être simplement généralisées à p variables [JDHR10].

Outils graphiques

Outre la représentation en nuages de points des rangs $(r^i, s^i)_{i \in [1:n]}$, d'autres méthodes graphiques existent pour représenter la dépendance. Nous présentons ici le *Chi-plot*, proposé par Fisher & Switzer [FS01]. Basé sur le test d'indépendance en ANOVA classique, il repose sur l'estimation empirique des fonctions de répartition. La procédure est la suivante.

Pour tout $i \in [1 : n]$, calculer

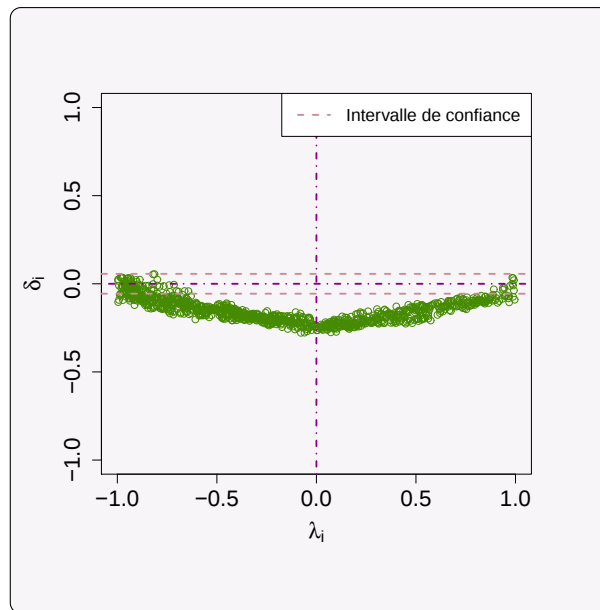
1. $H_i = \frac{1}{n-1} \sum_{j \neq i} \mathbb{I}(x^j \leq x^i, y^j \leq y^i);$
2. Les quantités $F_i = \frac{1}{n-1} \sum_{j \neq i} \mathbb{I}(x^j \leq x^i)$ et $G_i = \frac{1}{n-1} \sum_{j \neq i} \mathbb{I}(y^j \leq y^i).$
3. $\delta_i = \frac{H_i - F_i G_i}{\sqrt{F_i(1-F_i)G_i(1-G_i)}} \in [-1, 1];$
4. $\lambda_i = 4 \operatorname{sg}(\tilde{F}_i \tilde{G}_i) \max(\tilde{F}_i^2, \tilde{G}_i^2), \tilde{F}_i = F_i - 1/2, \tilde{G}_i = G_i - 1/2.$

La quantité δ_i mesure la distance entre la distribution qui porte la dépendance, et celle où les variables sont supposées indépendantes. Ainsi, sous l'hypothèse d'indépendance, $\delta_i = 0$. Similairement au coefficient de Spearman, y^i est une fonction croissante (décroissante) des x^i si $\delta_i = +1$ ($\delta_i = -1$).

La quantité λ_i mesure quant à elle la distance entre (x^i, y^i) et le centre du nuage de points. Dans un cas d'indépendance, la distribution asymptotique de λ_i est une distribution uniforme. Si $\lambda_i > 0$, les valeurs d'observations (x_j, y_j) , $j \neq i$, se comportent de manière analogue par rapport à la médiane. En revanche, si $\lambda_i < 0$, cela signifie que les observations ont un comportement contraire. Le *Chi-plot* est la représentation en nuage de points du couple (δ_i, λ_i) . Il a donc pour vocation de détecter visuellement une structure de dépendance qui peut s'avérer complexe. Pour plus de précisions, une région de confiance est aussi définie autour de la valeur nulle [FS01, MRL12].

Exemple (suite)

Nous représentons le *Chi-plot* de nos observations. Comme attendu, les données sont négativement éloignées de l'axe $\delta = 0$, signifiant que les données sont bien négativement associées quelque soit leur répartition. Notons également que les données sont à l'intérieur de la région de confiance aux extrémités, signifiant que les données x_j, y_j ayant une distribution fortement contraire ou fortement similaire sont indépendantes.



Le *Chi-plot* permet de visualiser la dépendance à l'aide d'une méthode très rapide à implémenter. Néanmoins, il présente quelques défauts concernant les valeurs aberrantes [FS01, GF07]. D'autres représentations graphiques plus complexes permettent de remédier à ce type de problème, et nous renvoyons à [GB03, GF07] pour plus de détails.

En utilisant ces outils, notre utilisateur a conclu que ses données présentaient une forme de dépendance. Naturellement, il a envie de savoir de quel(s)

modèle(s) théorique(s) ses données se rapprochent le plus. Dans la section suivante, nous présentons au préalable les techniques d'échantillonnage, répondant d'abord à la question 3 de l'utilisateur, puisque cela sera nécessaire pour définir un choix de modèle adapté (question 2).

2.3 Techniques d'échantillonnage de variables dépendantes

2.3.1 Méthode basée sur la corrélation de rangs

La méthode d'échantillonnage de Iman et Conover [IC82] a pour objectif d'induire une corrélation de rangs dans un échantillon d'observations des entrées. La technique repose essentiellement sur la décomposition de Cholesky. Pour décrire cette méthode, nous supposons que nous disposons d'une matrice d'observations $\mathbb{X}$ issues de $\mathbf{X} = (X_1, \dots, X_p)$ et d'une matrice de corrélation sur les rangs C dont la décomposition de Cholesky est donnée par $C = P^t P$. Par ailleurs, nous reprenons les notations antérieures et notons $\rho_n^{\text{Pear}}(A)$ et $\rho_n^{\text{Spear}}(A)$ la matrice de corrélation empirique de Pearson, et la matrice de corrélation empirique de Spearman de A , respectivement. Chaque étape de la procédure qui suit est illustrée par une simulation en R de $n = 500$ observations.

La première étape de ce procédé requiert le calcul de scores a_i , $i \in [1 : n]$. Pour cela, les auteurs proposent d'utiliser le score de van der Waerden [VdW52]. Ces scores sont en effet utilisés pour tester si deux populations ont même distribution, en se basant sur le critère de rang des observations. Ce sont les scores que nous avons utilisés dans notre illustration. Les p permutations indépendantes du score devraient permettre d'obtenir une matrice R décorrélée. Néanmoins, la résolution numérique engendre des erreurs qui ont pour conséquence que R peut ne pas être totalement décorrélée, d'où la nécessité de la décomposition de Cholesky sur la matrice de corrélations de R .

L'avantage de cette méthode est qu'elle est très intuitive, et simple à appliquer. Cependant, elle est d'un usage très limité, puisque certains modèles nécessitent une structure de dépendance spécifique, et plus complexe.

Dans la suite, nous revenons sur les copules, et donnons des méthodes numériques pour simuler une distribution de copules.

Procédure

Illustration

Soit une matrice
de corrélation
 $C = P^t P \in \mathcal{M}_{p,p}(\mathbb{R})$

```
C[1] C[2]
[1,]  1.0  0.5
[2,]  0.5  1.0
```

Soit une matrice
d'observations
 $\mathbb{X} \in \mathcal{M}_{n,p}(\mathbb{R})$

```
X[1] X[2]
[1,]  0.11 -1.26
[2,]  0.22  0.54
[3,] -0.13  1.63
[4,]  0.27 -0.29
```

Générer une matrice
 $R \in \mathcal{M}_{n,p}(\mathbb{R})$ de
 p permutations
indépendantes d'un
ensemble $a_i, i \in [1 : n]$

```
R[1] R[2]
[1,] -0.02 -0.31
[2,] -0.02 -0.15
[3,] -0.10  0.40
[4,]  0.46  1.21
```

Calculer $\rho_n^{\text{Pear}}(R) = T$

```
a[1] a[2]
[1,]  0.87  0.58
[2,] -1.01 -0.26
[3,] -0.92  1.01
[4,] -0.55  1.19
```

Décomposer $T = Q^t Q$
et poser $S = P Q^{-1}$

```
T[1] T[2] S[1] S[2]
1.00 0.01 1.00 0.00
0.01 1.00 0.49 0.87
```

Calculer $R^* = R^t S$.
Et, $\rho_n^{\text{Pear}}(R^*) = C$

```
Rs[1] Rs[2] 1 2
[1,] -0.02 -0.27 247 195
[2,] -0.02 -0.14 246 226
[3,] -0.10  0.30 231 301
[4,]  0.46  1.27 339 452
```

Etablir le rang de
chaque colonne de R^*

```
X[1] X[2]
[1,]  0.05 -0.28
[2,]  0.05 -0.15
[3,] -0.03  0.23
[4,]  0.52  1.22
```

Appliquer cet
ordre à $\mathbb{X}$.
Et, $\rho_n^{\text{Spear}}(\mathbb{X}) = C$

```
> cor(X,method="spearman")
      X[1]      X[2]
X[1] 1.000000 0.485698
X[2] 0.485698 1.000000
```

2.3.2 Méthode d'échantillonnage des copules

Si les copules sont des outils théoriques permettant de définir une structure de dépendance dans une distribution (cf Section 2.1), elles peuvent tout aussi bien être utilisées en pratique. D'une part, la fonction copule permet de calculer le τ de Kendall ou le ρ de Spearman théorique, permettant ainsi d'établir un indice capable de mesurer la dépendance (cf Section 2.2). Dans cette partie, nous montrons qu'il est également possible d'exploiter les copules pour générer des échantillons d'observations issues d'une distribution dont la dépendance est déterminée par une copule. Rappelons d'abord que

$$F(x_1, \dots, x_p) = C(F_1(x_1), \dots, F_p(x_p)),$$

où, comme rappelé en introduction de ce paragraphe, les fonctions de répartition $F, F_1, \dots, F_p$ sont supposées continues. Pour générer un échantillon issu de la distribution de $\mathbf{X} = (X_1, \dots, X_p)$, nous utilisons les remarques faites en Notes 2.3, et nous procédons en deux étapes :

1. simuler $\mathbf{U} = (U_1, \dots, U_p)$ de distribution C ;
2. poser $X_i = F_i^{-1}(U_i), i \in [1 : p]$.

La difficulté majeure réside dans le point 1. Dans la suite, nous proposons deux méthodes afin de simuler des observations issues de C . Par souci de clarté, nous nous restreindrons au cas bidimensionnel, soit $p = 2$.

Méthode des distributions

Cette méthode implique la connaissance de F et de ses marges F_1, F_2 . Elle est applicable dans le cas où il est plus simple de générer des observations avec F qu'avec C . La méthode se compose des étapes suivantes :

1. Simuler des réalisations du vecteur $\mathbf{X}$ de distribution F ;
2. Appliquer la transformation uniforme $\mathbf{U} = (F_1(X_1), F_2(X_2))$ sur l'échantillon. En effet, comme $\mathbf{U}$ est uniforme, nous pouvons reprendre le calcul fait en Section 2.1, puis par simples manipulations, obtenir que $C(u_1, u_2) = F(F_1^{-1}(u_1), F_2^{-1}(u_2))$, ce qui correspond à la caractérisation de la copule C donnée au Corollaire 2.1.

Par exemple, supposons que l'on veuille simuler des réalisations d'une distribution dont la copule est gaussienne, et dont les marges F_1 et F_2 sont exponentielle et de Student respectivement. Posons Φ la fonction de répartition jointe d'une loi gaussienne, et ϕ ses marges (supposées identiques). Alors, la méthode se déroule comme suit :

1. Simuler n réalisations $(\mathbf{y}^1, \dots, \mathbf{y}^n)$, avec $\mathbf{y}^l = (y_1^l, y_2^l), l \in [1 : n]$, issues de la distribution Φ ;
2. Calculer $u_1^l = \phi(y_1^l)$ et $u_2^l = \phi(y_2^l), l \in [1 : n]$. Ainsi, $\mathbf{u}^l = (u_1^l, u_2^l)$ est uniforme, $l \in [1 : n]$;

3. Calculer $x_1^l = F_1^{-1}(u_1^l)$ et $x_2^l = F_2^{-1}(u_2^l)$, $j \in [1 : n]$. Ainsi, l'échantillon $(\mathbf{x}^l)_{l \in [1:n]} = (x_1^l, x_2^l)_{l \in [1:n]}$ admet une dépendance gaussienne, mais les marginales sont respectivement exponentielle et de Student.

Méthode des distributions conditionnelles

Cette méthode implique aussi la connaissance des marges F_1 et F_2 . Néanmoins, la méthode des distributions conditionnelles permet de réduire la génération issue d'une distribution p -dimensionnelle à p générations de distribution unidimensionnelle [Dev86]. En effet, par le théorème de Bayes [Bar12], il vient, pour tout $(u_1, \dots, u_p) \in \mathbb{R}^p$

$$F(u_1, \dots, u_p) = F_1(u_1)F_2(u_2|u_1) \cdots F_p(u_p|u_1, \dots, u_{p-1}).$$

Dans le cas où $p = 2$, nous pouvons générer des observations de la distribution F par la génération d'observations issues de F_1 d'abord, puis par la génération issue de F_2 conditionné. La méthode des distributions conditionnelles repose sur ce principe. Elle est décrite ci-dessous.

1. Générer u_1 et u_2 issus de la loi uniforme $\mathcal{U}[0, 1]$;
2. Déterminer u_2 à l'aide de la copule conditionnelle $C_{2|1}$. Notons au passage que la copule conditionnelle $C_{2|1}$ est définie comme

$$C_{2|1}(u_1, u_2) = P(U_2 \leq u_2 | U_1 = u_1) = \partial_1 C(u_1, u_2).$$

Pour obtenir u_2 , il reste à définir le quantile de $C_{2|1}$, c'est-à-dire, par définition,

$$C_{2|1}^{-1}(p; u_1) = \inf_t \{t : C_{2|1}(u_1, t) = p\}.$$

Ainsi, il suffit de trouver u_2 tel que $C_{2|1}^{-1}(p; u_1) = u_2$.

Cette méthode est particulièrement attractive pour les copules pour lesquelles il est aisé de calculer la dérivée partielle. Par exemple, un échantillon issu de la copule de Clayton, ou encore la copule de Farlie-Gumbel-Morgenstern présentées à la Section 2.1.2 peut être généré via cette technique.

Exemple (suite)

La simulation de nos données (issues d'une distribution de copule FGM) aurait pu être faite ainsi :

$$\begin{aligned} C_{2|1}(u_1, u_2) &= \partial_1 C_\theta(u_1, u_2) \\ &= u_2[1 + \theta(1 - 2u_1)] + u_2^2[-\theta(1 - 2u_1)]. \end{aligned}$$

D'où

$$C_{2|1}^{-1}(u_1, u_2) = \frac{A - \sqrt{A^2 - 4(A-1)u_2}}{2(A-1)}, \quad A = 1 + \theta(1 - 2u_1).$$

Si $u_1, v_2 \sim \mathcal{U}[0, 1]$, alors

$$(u_1, u_2) = \left(u_1, \frac{A - \sqrt{A^2 - 4(A-1)v_2}}{2(A-1)} \right) \sim C_\theta^{FGM}.$$

Note 2.5. Ces deux techniques requièrent la connaissance des F_i . Néanmoins, si celles-ci ne sont pas connues, il est aisé de les estimer empiriquement :

$$F_{i,n}(x) = \frac{1}{n} \sum_{j=1}^n \mathbb{I}_{(x_i^j \leq x)}, \quad \forall i \in [1 : p].$$

Par le théorème de Glivenko-Cantelli, nous avons

$$\sup_t |F_{i,n}(t) - F_i(t)| \xrightarrow{n \rightarrow +\infty} 0 \text{ p.s.}$$

Notons qu'il n'existe pas toujours d'expression analytique de $C_{2|1}^{-1}$ pour appliquer la technique des distributions conditionnelles. Cependant, nous disposons de plusieurs méthodes numériques pour estimer cette quantité [Dev86].

La liste des différentes méthodes proposées n'est bien sûr pas exhaustive. Les logiciels statistiques (Matlab, R, ...) proposent des fonctionnalités pour les copules les plus standards. Néanmoins, il existe des copules non programmées et pour lesquelles les méthodes présentées ne sont pas très adéquates. Pour certaines copules, au vu de leur expression, une méthode spécifique doit être adoptée.

2.4 Choix de modèles

Le moyen le plus naturel pour choisir un modèle est de chercher l'adéquation entre les données dont on dispose, et les copules connues dans la littérature. Supposons que l'on dispose de f familles de copules $\mathcal{F}_1, \dots, \mathcal{F}_f$ dépendantes d'un paramètre θ , soit $\mathcal{F}_k = \{C_\theta^k, \theta \in \mathbb{R}^d, d \geq 1\}$, $k \in [1 : f]$. Pour comparer les copules théoriques aux données disponibles, nous devons, pour chaque famille $\mathcal{F}_k$,

1. Estimer le paramètre inconnu θ pour obtenir le représentant $C_{\hat{\theta}}^k$ de $\mathcal{F}_k$;
2. Générer un échantillon artificiel à partir de la copule $C_{\hat{\theta}}^k$;
3. Se servir d'un outil de comparaison et conclure.

La seconde étape a fait l'objet de la Section 2.3.2. Dans la suite, nous donnons des outils permettant de réaliser les points 1 et 3.

2.4.1 Estimation de θ

Estimation par le ρ de Spearman ou le τ de Kendall

Dans la Section 2.2, nous avons pu constater que les mesures basées sur les rangs admettaient une expression analytique dépendante de la copule C_θ . L'expression analytique de ces mesures pour un grand nombre de copules s'en déduit. De manière générale, si $\theta = h(\rho^{\text{Spear}})$, ou $\theta = g(\tau)$, nous disposons d'une estimation de θ en posant,

$$\hat{\theta} = h(\rho_n^{\text{Spear}}), \text{ ou } \hat{\theta} = g(\tau_n),$$

où les estimateurs ρ_n^{Spear} et τ_n , donnés à la Section 2.2, sont les estimateurs empiriques de ρ^{Spear} et τ .

Exemple (suite)

Comme C_θ est la copule FGM, alors

$$\begin{aligned} \rho^{\text{Spear}} &= 12 \int uv[1 + \theta(1-u)(1-v)]dudv - 3 \\ &= \frac{\theta}{3}, \end{aligned}$$

Donc $\theta = 3\rho^{\text{Spear}}$.

```
> rho_n<-cor(X,Y,method="spearman")
> rho_n
[1] -0.3950319
> theta<-rho_n*3
> theta
[1] -1.185096
```

Estimation par maximum de pseudo-vraisemblance

Lorsque $\theta \in \mathbb{R}^d$, $d \geq 2$, l'estimation par les mesures de dépendance n'est plus exploitable. Dans ce cas, la technique basée sur la vraisemblance est

une bonne alternative. Plus précisément, cette technique repose sur la maximisation de la vraisemblance basée sur les rangs, soit

$$\begin{aligned}\hat{\ell}(\theta) &= \sum_{i=1}^n \log \left[c_{\theta} \left(\hat{F}_1(x^i), \hat{F}_2(y^i) \right) \right] \\ \hat{F}_1(x^i) &= \frac{r^i}{n+1} \\ \hat{F}_2(y^i) &= \frac{s^i}{n+1},\end{aligned}$$

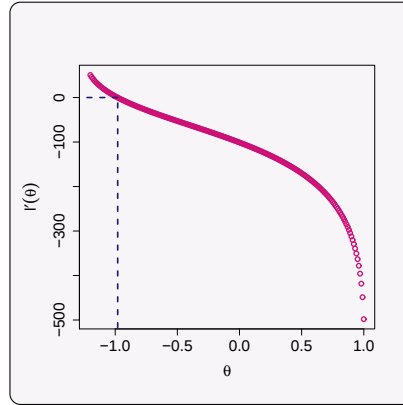
où, pour rappel, (r^i, s^i) est la paire de rangs associée à (x^i, y^i) , $i \in [1 : n]$. Dans ce cas, estimer θ revient à résoudre

$$\underset{\theta \in \mathbb{R}^d}{\text{Argmax}} \hat{\ell}(\theta).$$

Cependant, cette opération peut parfois s'avérer numériquement difficile à effectuer.

Exemple (suite)

A partir de nos observations (générées à partir d'une copule FGM où $\theta = -1$), nous représentons ici la dérivée de $\hat{\ell}(\theta)$ en fonction de θ , pour différentes valeurs de θ .



2.4.2 évaluation de la dépendance

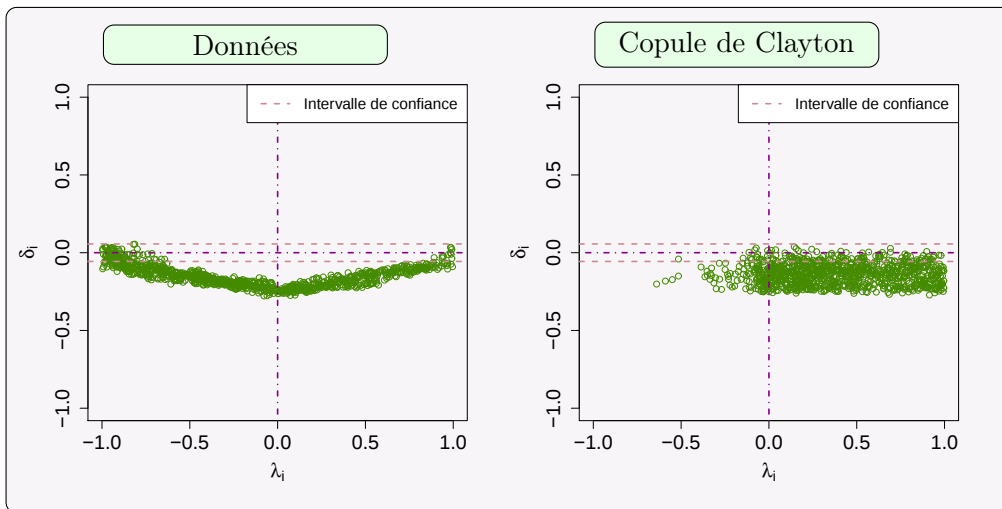
Une fois élu le "candidat" de chaque famille $\mathcal{F}_k$, c'est-à-dire que $C_{\hat{\theta}}^k$ est estimé pour tout $k \in [1 : f]$, Il suffit de simuler un n -échantillon issu de $C_{\hat{\theta}}^k$ (présentées en Section 2.3.2), et de le comparer avec les données réelles. La représentation la plus élémentaire est celle en nuage de points, qui permet de comparer les deux échantillons sur un même graphique. Cependant, à moins que les jeux de données soient très importants, ce qui n'est pas toujours

possible, une telle représentation peut s'avérer peu informative. Le *Chi-plot*, présenté page 64, permet de rendre compte de la structure de dépendance d'une paire de variables. Cet outil graphique peut être exploité pour définir le choix de la famille $\mathcal{F}_k$ la plus appropriée.

Des tests quantitatifs sont aussi disponibles pour comparer les différentes distributions. Néanmoins, peu de méthodes existent pour une identification précise. Plusieurs copules peuvent correspondre à un échantillon donné, auquel cas nous pouvons songer à une combinaison linéaire convexe de copules, qui est encore une copule [Nel06].

Exemple (suite)

Nous reprenons le *Chi-plot* obtenu à partir de nos observations (générées à partir d'une copule FGM), et nous le comparons à un échantillon issu de la copule de Clayton, illustrée à la Section 2.1.2. Nous observons que pour la copule de Clayton, les données ont une distribution similaire par rapport à la médiane ($\lambda_i > 0$), et qu'elles sont négativement associées ($\delta_i < 0$). Ainsi, nous pouvons en déduire que la dépendance de nos données n'est pas assimilable à celle de la copule de Clayton.



Chapitre 3

Analyse de sensibilité pour variables d'entrée dépendantes

Lorsqu'un système est modélisé par plusieurs facteurs d'entrée dépendants, l'analyse de sensibilité offre un cadre plus restrictif de solutions que pour le cas indépendant pour mesurer la contribution d'une ou plusieurs entrées dans le modèle. En effet, l'indice de Sobol présenté au Chapitre 1 est une mesure exacte et univoque de sensibilité dont les propriétés découlent de la décomposition de Hoeffding, construite sur l'hypothèse d'indépendance des variables d'entrée. Dans le cas où les entrées sont dépendantes, les propriétés d'orthogonalité de la décomposition fonctionnelle ANOVA ne sont donc plus vérifiées. Néanmoins, les indices de Sobol de premier ordre sont des mesures d'importance que l'on peut encore exploiter dans le but de hiérarchiser les entrées dépendantes [STCR04].

En revanche, les méthodes d'estimation de l'indice de Sobol doivent être adaptées pour rendre compte de la dépendance. Dans la Section 3.1, nous proposons d'exposer plusieurs travaux faisant état de méthodes d'échantillonnage pour l'estimation des indices de Sobol dans le cas dépendant. La première technique, étudiée par Kucherenko *et al.* [KTA12], revient sur la décomposition de la variance totale, initiée par [McK97] et mentionnée au Chapitre 1. Xu & Gertner [XG07] reprennent quant à eux la décomposition FAST développée à la Section 1.3.2 du Chapitre 1, et adaptent la technique d'échantillonnage du FAST au cas de facteurs corrélés. Da Veiga *et al.* [DV07, DVWG09] se concentrent sur une estimation non paramétrique de l'indice de Sobol, et en apportent les propriétés théoriques. Enfin, Mara *et al.* [MT12] contournent la dépendance en utilisant une technique d'orthogonalisation des variables pour se ramener, dans le cas très spécifique de distribution gaussienne, à une estimation classique des indices de Sobol.

Cependant, l'utilisation des indices de Sobol est parfois mise en doute, car la propriété de sommation des indices égale à 1 découle de la décomposi-

tion orthogonale de η , qui n'est plus vérifiée dans un contexte de variables dépendantes. L'interprétation des indices peut en être sérieusement affectée, car la dépendance entre deux facteurs peut être implicitement incluse dans leurs indices, et produire une information redondante. A titre d'exemple, considérons le modèle analytique suivant,

$$Y = X_1 + X_2, \quad \begin{pmatrix} X_1 \\ X_2 \end{pmatrix} \sim N(0, \Sigma), \quad \Sigma = \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}.$$

L'expression analytique des indices de Sobol est donnée par

$$\begin{aligned} S_1 &= \frac{(1 + \rho)^2}{2 + 2\rho} \\ S_2 &= \frac{(1 + \rho)^2}{2 + 2\rho} \\ S_{12} &= \frac{2\rho^2(1 + \rho)}{2 + 2\rho}. \end{aligned}$$

Maintenant, comparons les valeurs des indices pour $\rho = 0$ (c'est-à-dire pour des variables indépendantes), et $\rho = 0.9$ (c'est-à-dire pour des variables fortement dépendantes). Ces valeurs sont regroupées dans la Table 3.1

	$\rho = 0$	$\rho = 0.9$
S_1	0.5	0.95
S_2	0.5	0.95
S_{12}	0	0.81

TABLE 3.1 – Comparaison des indices de Sobol pour $\rho = 0$ et $\rho = 0.9$

Dans le cas indépendant, la contribution de X_1 est la même que celle de X_2 dans le modèle. Ce résultat est en concordance avec le modèle analytique posé, où les variances de chacune des deux variables sont les mêmes. En revanche, l'introduction d'une corrélation vient perturber notre interprétation, puisqu'il en ressort que les contributions de X_1 et X_2 sont aussi identiques, mais très élevées, et que l'interaction entre X_1 et X_2 est aussi très contributive à la variabilité du modèle. En conséquence, les indices ne sont plus sommables à 1. L'indice de sensibilité associé à un facteur prend en compte la sensibilité de l'autre, signifiant qu'une même information est prise en compte plusieurs fois.

Cet exemple nous montre l'importance de considérer les modèles à facteurs dépendants différemment de ce qui a été vu jusque là. Il peut donc être judicieux de considérer de nouvelles constructions d'indices qui tiennent compte

de la dépendance des facteurs d'entrée. Plusieurs solutions, que nous catégorisons en deux classes, ont été proposées dans la littérature.

1. Les indices basés sur la distribution de la sortie offrent ici un cadre nouveau à l'analyse de sensibilité. En effet, ces mesures sont construites à partir d'une distance métrique définie sur l'espace de probabilité des entrées, et non plus à partir des moments. Initiés par Chun *et al.* [CHT00], les mesures de sensibilité de Borgonovo *et al.* [Bor07, BCT11] reposent sur une version simplifiée de la distance de Wassertein. Ces travaux sont détaillés dans la Section 3.2.
2. L'analyse de la variance offre un cadre très général pour détecter et estimer les effets dans un modèle, donnant lieu à de puissants outils pour l'étude de la sensibilité. La décomposition ANOVA reste attractive pour l'analyse d'entrées dépendantes, mais nécessite cependant une adaptation pour rendre compte des corrélations dans un modèle. Plusieurs contributions, détaillées à la Section 3.3, proposent de définir une nouvelle décomposition ANOVA. En fonction de la décomposition choisie, des outils appropriés sont définis, tentant de rendre compte au mieux de la contribution des effets principaux ou d'interaction sur la sortie. Pour définir une décomposition, Bedford [Bed98] part du constat que les termes de la décomposition (1.1) se définissent comme des projections successives et suggère d'utiliser le procédé d'orthogonalisation de Gram-Schmidt sur une famille composée d'espérances conditionnelles. Xu & Gertner [XG08], suivis par Li *et al.* [LRY⁺10] proposent une autre décomposition, qui repose sur une reconstruction du modèle théorique η par un méta-modèle. Cette substitution permet ainsi d'en déduire un indice de sensibilité basé sur la variance.

Il est à noter que nos travaux contributifs reposent sur une décomposition ANOVA, mais nous expliquerons en quoi elle diffère des travaux relatés dans ce chapitre. Notons enfin que pour chacun des apports détaillés ici, les auteurs ont considéré la mesure de Lebesgue comme la mesure de référence sur $\mathbb{R}^p$.

3.1 Techniques d'échantillonnage

Le premier développement d'une méthode d'estimation des indices remonte aux travaux de Sobol [Sob93], qui, dans le cas où l'ensemble des entrées forme des groupes indépendants de variables dépendantes, suggère de quantifier la sensibilité de chaque groupe indépendant via l'indice de Sobol. Cette idée sera reprise et appliquée plus tard par Jacques *et al.* [JLD06], qui proposent de calculer ces indices par la technique de Monte Carlo. Néanmoins, cette approche a un inconvénient majeur : cet indice permet de quantifier la contribution d'un groupe de variables dans le modèle, mais pas la contribution individuelle de chaque variable. Ainsi, l'indice très élevé d'un en-

semble de variables ne permet pas de savoir si ceci est dû à un facteur en particulier, puisque cette technique ne permet pas de détecter la sensibilité individuelle dans un groupe donné. Cela est d'autant plus problématique lorsque les groupes les plus contributeurs sont constitués d'un grand nombre de variables. L'ensemble des travaux présentés maintenant ont pour objectif de quantifier la part individuelle de contribution des variables ou des effets d'interaction dans un modèle.

3.1.1 Décomposition de la variance totale et technique de Monte Carlo

Inspiré par les travaux de McKay [McK97] sur le *rapport de corrélation*, précisé à la Section 1.1 du Chapitre 1, Kucherenko *et al.* [KTA12] proposent de ré-exploiter la décomposition de la variance totale [WHH06] pour des modèles à entrées dépendantes. Pour cela, les auteurs considèrent un découpage de $\mathbf{X}$ en deux sous-ensembles de variables, $\mathbf{X} = (\mathbf{X}_u, \mathbf{X}_{-u})$, avec $u \in S \setminus \{\emptyset\}$. La décomposition de la variance totale s'exprime comme

$$V(Y) = V[\mathbb{E}(Y|\mathbf{X}_u)] + \mathbb{E}[V(Y|\mathbf{X}_u)].$$

En utilisant le même raisonnement que pour le *rapport de corrélation*, l'indice de sensibilité S_u associée à $\mathbf{X}_u$, qui correspond à l'indice *closed* notamment étudié dans [Jan12], est donné par

$$S_u = \frac{V[\mathbb{E}(Y|\mathbf{X}_u)]}{V(Y)}.$$

Pour l'estimation de cet indice, un algorithme de Monte Carlo est adopté. Afin d'améliorer la convergence numérique de l'estimateur de l'indice, on considère, pour un $u \in S$ donné, $\mathbf{X}' = (\mathbf{X}'_u, \mathbf{X}'_{-u})$, une copie indépendante de $\mathbf{X}$. Nous considérons également un vecteur $\mathbf{X}^*_{-u}$ issu de la densité conditionnelle $p_{\mathbf{X}|\mathbf{X}_u}(\mathbf{X}|\mathbf{X}_u)$. Posons $Y^* = \eta(\mathbf{X}_u, \mathbf{X}^*_{-u})$ et $Y' = \eta(\mathbf{X}')$. Ainsi, Y et Y' sont indépendants, et

$$\begin{aligned} V[\mathbb{E}(Y|\mathbf{X}_u)] &= \mathbb{E}[\mathbb{E}(Y|\mathbf{X}_u)^2] - \mathbb{E}(Y)^2 \\ &= \mathbb{E}[\mathbb{E}(Y|\mathbf{X}_u)\mathbb{E}(Y^*|\mathbf{X}_u)] - \mathbb{E}(Y)\mathbb{E}(Y') \\ &= \mathbb{E}[\mathbb{E}(Y\mathbb{E}(Y^*|\mathbf{X}_u)|\mathbf{X}_u)] - \mathbb{E}[\mathbb{E}(Y')Y] \\ &= \mathbb{E}[Y[\mathbb{E}(Y^*|\mathbf{X}_u) - \mathbb{E}(Y')]]. \end{aligned} \tag{3.1}$$

Prenons maintenant $(\mathbf{x}^l)_{l \in [1:n]}$, $(\mathbf{x}'^l)_{l \in [1:n]}$ les deux n -échantillons i.i.d issus de la distribution de $\mathbf{X}$. $(\mathbf{x}^{l*}_{-u})_{l \in [1:n]}$ est un autre n -échantillon i.i.d issu de la distribution de $\mathbf{X}^*_{-u}$, c'est-à-dire de la densité conditionnelle $p_{\mathbf{X}|\mathbf{X}_u}$.

Puis, posons $y^l = \eta(\mathbf{x}^l)$, $y'^l = \eta(\mathbf{x}'^l)$ et $y^{l*} = \eta(\mathbf{x}_u^l, \mathbf{x}^{l*}_{-u})$, $\forall l \in [1 : n]$.

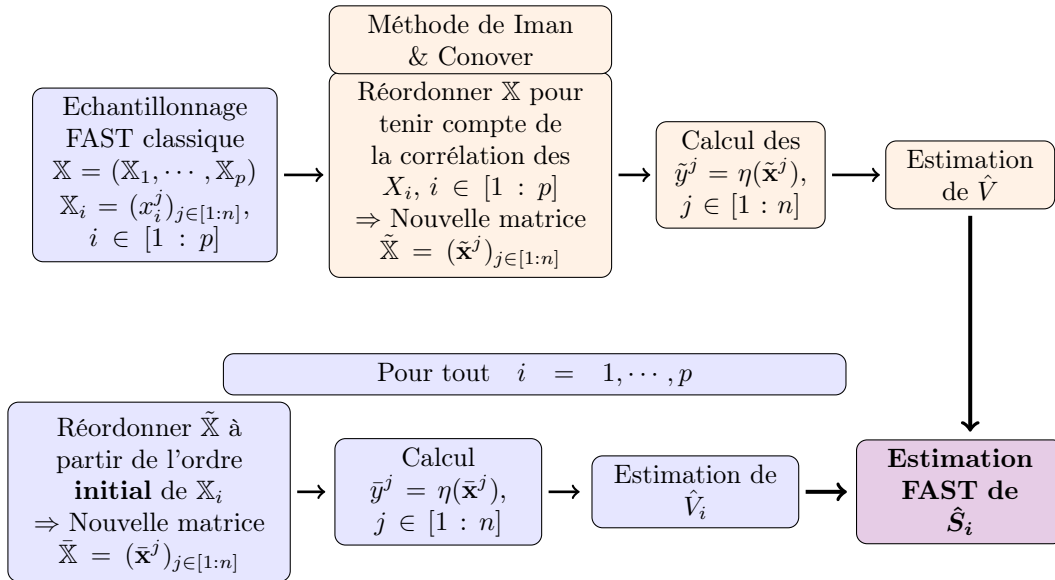
L'estimateur de S_u se déduit de (3.1) par

$$\hat{S}_u = \frac{\frac{1}{n} \sum_{l=1}^n y^l [y^{l*} - y^{l'}]}{D}, \quad D = \frac{1}{n} \sum_{l=1}^n (y^l)^2 - \left(\frac{1}{n} \sum_{l=1}^n y^l \right)^2.$$

Ainsi, cette technique d'échantillonnage se distingue d'une méthode construite pour le cas indépendant par la génération d'un échantillon issu de la densité conditionnelle $p_{\mathbf{X}|\mathbf{X}_u}$. Dès lors, cette étape est primordiale, mais difficile. En pratique, les auteurs se concentrent sur une densité gaussienne afin d'obtenir une expression analytique des densités conditionnelles, elles aussi gaussiennes.

3.1.2 Technique d'échantillonnage FAST

Comme présenté au Chapitre 1, la méthode FAST propose une technique rigoureuse pour reconstruire les indices de Sobol sous l'hypothèse d'indépendance des entrées $(X_i)_{i \in [1:p]}$. Xu & Gertner [XG07] proposent une technique d'échantillonnage afin d'adapter la méthode FAST au cas où les facteurs sont corrélés. En utilisant la méthode décrite par Iman & Conover [IC82] pour induire une corrélation des rangs sur la matrice d'observations des entrées (Section 2.3 du Chapitre 2), cette méthode permet d'intégrer une corrélation dans l'estimation des indices de sensibilité basée sur la décomposition FAST. Elle est décrite par le schéma ci-dessous.



La première étape consiste, comme expliqué à la Section 2.3 du Chapitre 2, à réordonner l'échantillon des observations indépendantes pour restituer la structure de corrélation des rangs. De cette manière la fréquence ω_i d'un

paramètre, qui caractérise la décomposition FAST, capture partiellement les variations des facteurs corrélés avec ce paramètre. Ainsi, la variance globale sera estimée, tenant compte de la dépendance des facteurs. Cependant, pour restituer la contribution propre de chaque variable X_i , les auteurs préconisent de réarranger la matrice corrélée suivant l'ordre initial des observations $(x_i^j)_{j \in [1:n]}$. Par exemple, considérons $p = 3$ et un réarrangement effectué selon X_1 . Donnons la ligne de la matrice corrélée $\tilde{\mathbb{X}}$ correspondant à l'observation x_1^1 ,

$$\begin{pmatrix} x_1^1 & x_2^k & x_3^j \end{pmatrix}, \quad k, j \in [1 : n]. \quad (3.2)$$

Dans ce cas, la nouvelle matrice $\bar{\mathbb{X}}$ aura pour première ligne celle donnée en (3.2), afin d'obtenir la distribution marginale de X_1 pour calculer $\hat{V}_1$.

Grâce à cette technique, les fréquences associées à une variable, sous-jacentes à la technique d'échantillonnage, capturent non seulement les variations de cette variable, mais aussi les variations des facteurs corrélés avec cette même variable.

3.1.3 Estimation par polynômes locaux

Pour estimer l'indice de Sobol de premier ordre, Da Veiga *et al.* [DVWG09, DV07] préconisent l'utilisation de la méthode des polynômes locaux [FG96]. Rappelons d'abord que l'indice de Sobol du premier ordre s'exprime comme, pour tout $i \in [1 : p]$,

$$S_i = \frac{V[\mathbb{E}(Y|X_i)]}{V(Y)}. \quad (3.3)$$

Soient $(y^l, x_i^l)_{l=1, \dots, n_1}$ et $(y^k, x_i^k)_{k=1, \dots, n_2}$ deux échantillons d'observations de (Y, X_i) de taille n_1 et n_2 respectivement. En posant

$$m(x) = \mathbb{E}(Y|X_i = x),$$

et en utilisant le développement de Taylor, la fonction m et ses dérivées peuvent être estimées via une estimation des moindres carrés pondérés, pour tout $k \in [1 : n_2]$,

$$\begin{pmatrix} \hat{m}(x_i^k) \\ \hat{m}'(x_i^k) \\ \dots \\ \frac{\hat{m}^{(q)}(x_i^k)}{q!} \end{pmatrix} = \underset{\substack{\beta_j \in \mathbb{R} \\ j \in [0:q]}}{\text{Argmin}} \sum_{l=1}^{n_1} \left(y^l - \sum_{j=0}^q \beta_j (x_i^l - x_i^k)^j \right)^2 \cdot K \left(\frac{x_i^l - x_i^k}{h} \right),$$

où la fonction noyau $K(\cdot)$ et le paramètre d'échelle h permettent d'inférer un poids aux observations x_i^l autour de x_i^k qui vont intervenir dans le calcul du

polynôme. Dès lors, la fonction m est estimée par le biais du premier échantillon, alors que sa prédiction est réalisée en les points du second échantillon. En pratique, on peut considérer un échantillon de taille N , et le découper en deux sous-échantillons de tailles respectives n_1 et n_2 tel que $N = n_1 + n_2$. Une autre façon est d'utiliser la technique du *leave-one-out* pour une meilleure approximation. Cela permet également de générer un échantillon de plus faible taille, ce qui peut être préférable si le code de calcul générant les observations est coûteux en temps de calcul.

Finalement, nous pouvons en déduire une estimation de (3.3) par estimation empirique des variances, soit,

$$\hat{S}_i = \frac{\frac{1}{n_2-1} \sum_{j=1}^{n_2} (\hat{m}(x_i^j) - \frac{1}{n_2} \sum_{j=1}^{n_2} \hat{m}(x_i^j))^2}{\hat{\sigma}_Y^2}, \quad \hat{\sigma}_Y^2 = \frac{1}{n_2-1} \sum_{j=1}^{n_2} (y^j - \frac{1}{n_2} \sum_{j=1}^{n_2} y^j)^2. \quad (3.4)$$

A partir des propriétés théoriques des estimateurs par polynômes locaux, les auteurs montrent que l'estimateur donné par (3.4) est asymptotiquement sans biais. L'utilisation d'une régression non paramétrique des indices de sensibilité permet ainsi de contourner la difficulté des variables corrélées en travaillant directement sur les distributions marginales, tout en assurant des résultats de convergence, qu'il est plus difficile d'établir pour les autres méthodes.

3.1.4 Procédé d'orthogonalisation des facteurs d'entrée

Dans une optique différente des procédures présentées jusque là, Mara *et al.* [MT12] cherchent à se ramener au cas indépendant en transformant les entrées dépendantes X_i , $i \in [1 : p]$ par orthogonalisation. Inspirée par les travaux de Bedford [Bed98], Mara *et al.* vont se servir de la technique de Gram-Schmidt [Mey00] pour orthogonaliser les entrées. Sous l'hypothèse que les moments conditionnels d'ordre un caractérisent complètement la dépendance entre les variables- par exemple si la structure de dépendance se définit totalement par une matrice de corrélation- les entrées sont ordonnées puis chacune est orthogonalisée par son espérance conditionnée par celles qui lui précèdent. Plus précisément, les nouvelles entrées orthogonales $\bar{X}_1, \dots, \bar{X}_p$, construites à partir de $X_1, \dots, X_p$, sont données par

$$\begin{aligned} \bar{X}_1 &= X_1 \\ \bar{X}_i &= X_i - \sum_{j < i} \mathbb{E}(X_i | \bar{X}_j), \quad i \geq 2. \end{aligned}$$

Sous l'hypothèse de moments conditionnels, le nouvel ensemble $\bar{\mathbf{X}} = (\bar{X}_1, \dots, \bar{X}_p)$ est indépendant, et les espérances conditionnelles peuvent être approchées

par des fonctions de régression unidimensionnelles [LCT07]. Lorsque les composantes de $\mathbf{X}$ ne suivent pas une distribution pour laquelle l'indépendance est équivalente à l'orthogonalité (comme une distribution gaussienne, par exemple), il est possible d'utiliser la transformée de Nataf pour rendre les variables indépendantes, sous condition que la structure de dépendance soit caractérisée par une copule elliptique. En effet, la transformation de Nataf permet de transformer des variables dépendantes dont la copule est elliptique en des variables gaussiennes centrées et réduites [Nat62, LD09]. Cette méthode peut donc être également utilisée ici pour rendre les facteurs indépendants.

Une fois les variables rendues orthogonales, on peut leur appliquer les indices de Sobol classiques, soit

$$\begin{aligned} S_1 &= \frac{V[\mathbb{E}(Y|X_1)]}{V(Y)} \\ S_{2-1} &= \frac{V[\mathbb{E}(Y|\bar{X}_2)]}{V(Y)} \\ &\vdots \\ S_{p-1 \dots p-1} &= \frac{V[\mathbb{E}(Y|\bar{X}_p)]}{V(Y)}. \end{aligned}$$

Ces derniers sont estimés par la méthode des polynômes de chaos, méthode décrite en Section 1.3 du Chapitre 1. Cependant, les indices de Sobol ne sont interprétables que pour les variables orthogonalisées, c'est-à-dire les variables conditionnées par les variables traitées antérieurement dans le procédé d'orthogonalisation. Ainsi, la méthode doit être réitérée $p!$ pour que l'ordre de traitement des variables dans le procédé n'influence pas le résultat final.

Points importants

Ces méthodes reposent finalement sur un procédé en deux étapes. La première a pour objectif de gérer et traiter la dépendance des entrées à travers les observations. Une fois ce traitement opéré, la seconde étape consiste à utiliser une mesure de sensibilité classiquement exploitée en analyse de sensibilité.

3.2 Mesures basées sur la distribution

Contrairement aux mesures basées sur les moments définies à la Section 1.1, les indices basés sur la distribution prennent en compte l'intégralité de la

distribution jointe des variables d'entrée. Ces nouvelles mesures de sensibilité génèrent donc de nouvelles questions quant à leur interprétation et à leur estimation.

3.2.1 Mesure d'importance basée sur les quantiles

Dans l'optique de réduire l'incertitude de la sortie liée aux changements distributionnels des entrées, Chun *et al.* [CHT00] ont proposé la mesure suivante

$$MD_i = \frac{\left(\int_0^1 [F_i^{-1}(q) - F_Y^{-1}(q)]^2 dq \right)^{1/2}}{\mathbb{E}(Y)}, \quad \mathbb{E}(Y) \neq 0,$$

où F_Y est la fonction de répartition de $Y = \eta(\mathbf{X})$, et F_i est la fonction de répartition de Y lorsque la variable X_i a subi un changement.

En utilisant la mesure MD_i , l'auteur veut quantifier l'écart moyen entre les fonctions quantiles de deux distributions. Ainsi, si un changement de X_i perturbe fortement la distribution jointe des facteurs d'entrée, l'indice MD_i est élevé, ce qui signifie que X_i est de forte contribution dans le modèle. Le changement le plus naturel, et celui opéré par les auteurs, consiste à fixer X_i à sa valeur nominale.

En pratique, la distribution empirique de F_Y , dénotée $F_{Y,n}$ se déduit naturellement comme

$$F_{Y,n}(y) = \frac{1}{n} \sum_{l=1}^n \mathbb{I}_{\{y^l \leq y\}},$$

où $(y^l, \mathbf{x}^l)_{l \in [1:n]}$ est un n -échantillon de Monte Carlo issu de la distribution de $(Y, \mathbf{X})$. De cette manière, il est aisé d'en déduire le quantile $F_{Y,n}^{-1}(q) = \inf_t \{t : F_{Y,n}(t) = q\}$. De la même manière, $F_{i,n}$ correspond à l'estimation empirique de F_i . L'estimation de l'indice MD_i est donnée par

$$MD_{i,n} = \frac{\sqrt{\frac{1}{n} \sum_{l=1}^n [F_{i,n}^{-1} - F_{Y,n}^{-1}]^2}}{\frac{1}{n} \sum_{l=1}^n y^l}.$$

Cet indice est conceptuellement simple à appréhender, et facile à calculer en pratique, même si la technique utilisée requiert une très grande taille d'échantillon. Cependant, la modification subie par X_i est le point central de la méthode, et peut donner lieu à de fortes variations de l'indice. Pour remédier à ce problème, Borgonovo *et al.* [Bor07, BT08, BCT11, BCT12] proposent un estimateur similaire, où X_i représente la valeur moyenne de l'ensemble des valeurs prises par X_i .

3.2.2 Mesure d'importance shiftée

Dans un cadre général, Borgonovo [Bor07] introduit la mesure suivante,

$$\begin{aligned}\delta_u &= \frac{1}{2} \mathbb{E}_{\mathbf{X}_u} [s(\mathbf{X}_u)] \\ &= \frac{1}{2} \int_{\mathbb{R}^{|u|}} p_{\mathbf{X}_u}(\mathbf{x}_u) \left[\int_{\mathbb{R}} |p_Y(y) - p_{Y|\mathbf{X}_u}(y)| dy \right] d\mathbf{x}_u, \quad \forall u \in S,\end{aligned}$$

où p_Y est la densité de Y , $p_{Y|\mathbf{X}_u}$ est la densité de Y conditionnellement à $\mathbf{X}_u$ et $p_{\mathbf{X}_u} = \int_{\mathbb{R}^{p-|u|}} p_{\mathbf{X}}(\mathbf{x}) d\mathbf{x}_{-u}$ est la fonction de densité marginale de $\mathbf{X}_u$. La fonction s est la fonction "shift", puisqu'elle caractérise l'écart entre deux distributions. Cette fonction shift permet de mesurer l'impact que peut avoir la modification d'un groupe de variables sur la distribution globale du modèle. L'indice δ_u quantifie donc l'écart moyen normé entre la distribution nominale d'un modèle et cette même distribution lorsque le modèle a subi une perturbation.

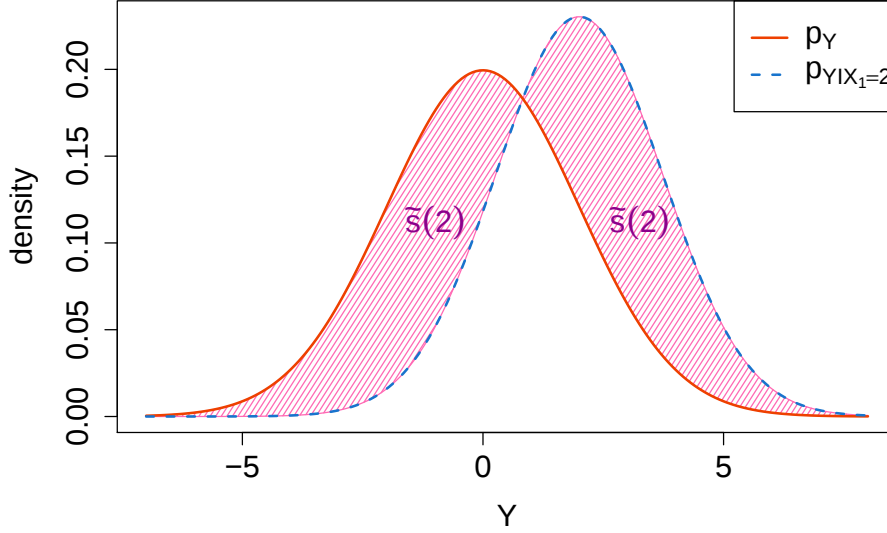
A titre d'illustration, on se donne le modèle additif suivant, où les entrées X_i sont indépendantes,

$$Y = \sum_{i=1}^4 X_i, \quad X_i \sim N(0, 1), \quad \forall i \in [1 : 4].$$

En supposant que l'on veuille connaître l'impact de la variable X_1 dans le modèle, nous pouvons d'ores et déjà calculer analytiquement et représenter graphiquement l'écart entre $p_Y(y)$ et $p_{Y|X_1}(y)$ lorsque X_1 est fixée à une certaine valeur. Par calcul, nous obtenons facilement que $Y \sim N(0, 4)$ et $Y|X_1 = x_1 \sim N(x_1, 3)$. La Figure 3.1 représente l'écart $\tilde{s}(x_1) := p_Y(y) - p_{Y|X_1=x_1}(y)$, lorsque $y \in [-10, 10]$ et pour $x_1 = 2$.

Les propriétés de δ_u , $\forall u \in S$ sont multiples. Comme l'indice de Sobol, $0 \leq \delta_u \leq 1$. $\delta_u = 0$ signifie l'indépendance de $\mathbf{X}_u$ et Y , alors que le maximum $\delta_u = 1$ est atteint lorsque $u = [1 : p]$. Borgonovo *et al.* montrent également que la fonction shift peut s'écrire en termes de fonctions de répartition. De plus, l'indice est invariant par transformation monotone du modèle [BCT11].

La méthode d'estimation se décompose en plusieurs étapes. A partir d'un échantillonnage Hypercube Latin (LHS) [IDZ80] des entrées $\mathbf{X}$ en une matrice $(\mathbb{X})_{ij} = x_j^i$ à n lignes et p colonnes, on en déduit un échantillon $(y^1, \dots, y^n)$ de la sortie du modèle Y . Une estimation de la densité p_Y est obtenue par la méthode classique à noyau de Parzen-Rosenblatt [Par62]. L'estimation de la densité conditionnelle $p_{Y|\mathbf{X}_u}$ s'obtient en fixant les colonnes correspondantes à $\mathbf{X}_u$, soit en fixant $(\mathbb{X})_{i \in [1:n], j \in u}$ à une valeur donnée, disons $(x_j^r)_{j \in u}$. Il suffit alors de ré-échantillonner les $p - |u|$ autres entrées du modèle sachant que $X_j = x_j^r$ pour tout $j \in u$, et d'en déduire un nouvel échantillon

FIGURE 3.1 – Représentation de $\tilde{s}(2)$

$(\tilde{y}^1, \dots, \tilde{y}^n)$ de Y . Pour chaque valeur $(x_j^r)_{j \in u}$, une estimation de la densité conditionnelle $p_{Y|\mathbf{X}_u}$ est faite par la technique des noyaux en utilisant les observations $(\tilde{y}^1, \dots, \tilde{y}^n)$. Plusieurs cas test, et applications concrètes ont été testés avec cette mesure d'importance [BT08, BCT11, BCT12].

Points importants

La définition de ces indices ne reposent pas sur les moments des variables, mais sur leur distribution jointe. Néanmoins, l'interprétation, au même titre que pour les indices de Sobol, reste problématique dans le cas de modèles à variables dépendantes car on ignore comment la distribution conditionnelle est affectée par la dépendance.

3.3 Indices basés sur la décomposition fonctionnelle

Les travaux qui vont être présentés maintenant ont une approche semblable à celle exposée au Chapitre 1, c'est-à-dire qu'ils reposent sur une décomposition du modèle. En ce sens, ces apports sont aussi ceux qui se rapprochent le plus des contributions développées aux Chapitres 6 et 7 de la Partie III. Cependant, pour chacun des travaux qui seront exposés, nous argumenterons

sur leurs principales faiblesses, et en quoi nous avons tenté d'y remédier à travers cette thèse.

3.3.1 Procédé d'orthogonalisation

Une autre manière d'appréhender les indices de Sobol est de considérer la famille des espérances conditionnelles de $\eta(\mathbf{X})$, définie par

$$G = \{g_u(\mathbf{X}_u) = \mathbb{E}(\eta(\mathbf{X})|\mathbf{X}_u), \forall u \in S\} \subset L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}),$$

puis de considérer une orthogonalisation de cette famille afin d'exprimer η de manière unique. Une façon très naturelle de procéder est d'utiliser la méthode d'orthogonalisation de Gram-Schmidt [Mey00].

Supposons d'abord que les variables $(X_i)_{i \in [1:p]}$ sont indépendantes. En établissant un ordre lexicographique des éléments de G c'est-à-dire

$$\begin{aligned} g_0 &= \mathbb{E}(\eta), \\ g_1(X_1) &= \mathbb{E}(\eta|X_1), \dots, g_p(X_p) = \mathbb{E}(\eta|X_p) \\ g_{p+1}(X_1, X_2) &= \mathbb{E}(\eta|X_1, X_2), \dots, g_{p(p+1)/2} = \mathbb{E}(\eta|X_{p-1}, X_p) \\ &\vdots \\ g_{2^p}(\mathbf{X}) &= \mathbb{E}(\eta(\mathbf{X})|\mathbf{X}) = \eta(\mathbf{X}), \end{aligned}$$

la méthode d'orthogonalisation de Gram-Schmidt permet de construire la famille libre $(\eta_u)_{u \in S}$ comme

$$\begin{aligned} \eta_u(\mathbf{X}_u) &= g_u(\mathbf{X}_u) - \sum_{v \subset u} \frac{\langle \eta_v, g_u \rangle}{\|\eta_v\|^2} \eta_v(\mathbf{X}_v) \\ &= g_u(\mathbf{X}_u) + \sum_{v \subset u} (-1)^{|u|-|v|} g_v(\mathbf{X}_v) \\ &= \mathbb{E}(\eta(\mathbf{X})|\mathbf{X}_u) + \sum_{v \subset u} (-1)^{|u|-|v|} \mathbb{E}(\eta(\mathbf{X})|\mathbf{X}_v), \quad \forall u \in S. \end{aligned}$$

On retrouve exactement les termes de la décomposition donnée en Section 1.1 du Chapitre 1. De plus, en notant que $\eta \in G$ constitue la dernière opération du procédé de Gram-Schmidt, on obtient la décomposition fonctionnelle ANOVA établie aux Sections 1.1-1.2.

Suivant cette idée, Bedford [Bed98], utilise le même procédé pour reconstruire η dans le cas où les entrées $\mathbf{X}$ sont potentiellement dépendantes. Reprenant le sous-espace G et l'ordre lexicographique de ses éléments, Bedford construit une famille libre $(\eta_u)_{u \in S}$. Ainsi, $\eta(\mathbf{X})$ se déduit à partir de cette nouvelle construction comme

$$\eta(\mathbf{X}) = \sum_{u \in S} \frac{\langle \eta_u, \eta \rangle}{\|\eta_u\|^2} \eta_u(\mathbf{X}_u).$$

L'orthogonalisation des variables permet ainsi d'obtenir une nouvelle forme de décomposition ANOVA de η comme une somme de projecteurs sur les espaces engendrés par $\mathbf{X}_u$, $u \in S$.

Par conséquent,

$$V[\eta(\mathbf{X})] = \sum_{u \in S} \frac{\langle \eta_u, \eta \rangle^2}{\|\eta_u\|^2},$$

où, pour tout $u \in S$, η_u est une combinaison linéaire des $g_v(\mathbf{X}_v) = \mathbb{E}(\eta|\mathbf{X}_v)$, avec $|v| \leq |u|$, tel que v est antérieur à u dans l'ordre lexicographique établi au départ.

La décomposition de la variance établie permet de quantifier la contribution d'un ensemble de variables dans le modèle. Pour quantifier le degré d'incertitude de η relativement à toutes les variables $\mathbf{X}_v$, $|v| \leq |u|$, nous pouvons poser

$$S_u^{\text{GS}} = \frac{\langle \eta_u, \eta \rangle^2}{\|\eta_u\|^2}.$$

En pratique, les espérances conditionnelles, qui ne sont pas simplifiables par rapport au cas indépendant, sont estimées par technique de Monte Carlo, et l'estimation peut s'avérer numériquement coûteuse. Néanmoins, le calcul d'intégrales par rapport à des distributions conditionnelles peut être simplifié dans le cas de dépendance par arbre. En effet, si les entrées admettent une telle structure de dépendance, les distributions impliquées dans le calcul des intégrales admettent la même structure, donnant lieu à des simulations immédiates [Coo97]. Cette démarche diffère de celle de Mara *et al.*, postérieure à celle-ci et présentée à la Section 3.1. En effet, Bedford considère ici une orthogonalisation par l'espérance conditionnelle de la sortie par les entrées préalablement ordonnées, alors que Mara *et al.* exploitent l'orthogonalisation des entrées par rapport à celles transformées antérieurement pour se ramener au cas indépendant. Cependant, ces deux procédés présentent un inconvénient majeur : ils dépendent de l'ordre dans lequel les quantités d'intérêt ont été orthogonalisées. Ainsi, les mesures de sensibilité peuvent être fortement affectées par la décision d'ordonnement des variables.

3.4 Reconstruction par méta-modélisation

Les travaux suivant découlent les uns des autres. Pour cette raison, nous proposons de les décrire chronologiquement dans une même et seule partie, en mettant l'accent sur les contributions personnelles de chaque papier.

Considérons d'abord une régression linéaire classique,

$$Y = \mathbb{E}(Y|\mathbf{X}) + \varepsilon, \quad \mathbb{E}(Y|\mathbf{X}) = \beta_0 + \sum_{i=1}^p \beta_i X_i. \quad (3.5)$$

Xu & Gertner [XG08] remarque que, sous l'hypothèse d'indépendance des entrées, la variance globale du modèle peut se décomposer linéairement, soit

$$V(Y) = \sum_{i=1}^p V(\beta_i X_i) + V(\varepsilon). \quad (3.6)$$

En utilisant la technique d'échantillonnage de l'hypercube latin pour obtenir un n -échantillon d'observations $(y^l, \mathbf{x}^l)_{l \in [1:n]}$, tel que $\mathbf{x}^l = (x_1^l, \dots, x_p^l)$, $\forall l \in [1:n]$, les coefficients $(\beta_i)_{i \in [1:p]}$ peuvent être obtenus par moindres carrés ordinaires, et les variances impliquées dans (3.6) peuvent être estimées empiriquement. Ainsi, la sensibilité de X_i est estimée par

$$\hat{S}_i = \frac{\hat{V}_i}{\hat{V}}, \quad \text{où } \hat{V}_i = \hat{\beta}_i \widehat{V(X_i)}, \quad \hat{V} = \widehat{V(Y)}.$$

Maintenant, nous supposons que (3.5) est toujours admise, mais que les facteurs d'entrée sont corrélés. La décomposition (3.6) n'est plus nécessairement vérifiée, car V_i inclut la dépendance avec les autres paramètres X_j , $j \neq i$. L'idée des auteurs est de décomposer la variance partielle V_i en une contribution prenant en compte les corrélations, notée V_i^C , et une contribution due aux variations non corrélées, notée V_i^U .

Pour cela, l'idée des auteurs est de poser, pour tout $i \in [1:p]$, la décomposition suivante,

$$X_i = \mathbb{E}(X_i|X_j, j \neq i) + Z_i, \quad \mathbb{E}(X_i|X_j, j \neq i) = \gamma_0 + \sum_{j \neq i} \gamma_j X_j. \quad (3.7)$$

Ainsi, $Z_i = X_i - \mathbb{E}(X_i|X_j, j \neq i)$ représente la part décorrélée de X_i , pour tout $i \in [1:p]$, car $\mathbb{E}(Z_i X_j) = 0$, pour tout $j \neq i$.

Pour en déduire la part de variance décorrélée associée à X_i , V_i^U , la régression suivante est étudiée,

$$Y = \mathbb{E}(Y|Z_i) + \varepsilon, \quad \mathbb{E}(Y|Z_i) = \alpha_0 + \alpha_i Z_i.$$

Puis, pour en déduire la variance partielle V_i , les auteurs se servent de la régression linéaire suivante,

$$Y = \mathbb{E}(Y|X_i) + e, \quad \mathbb{E}(Y|X_i) = \theta_0 + \theta_i X_i.$$

Ainsi, l'indice de sensibilité défini pour mesurer la contribution de X_i dans le modèle (3.5) est donné par

$$S_i^{\text{XG}} = \frac{V_i^U + V_i^C}{V}, \quad V = V(Y), \quad (3.8)$$

avec

$$V_i^U = V[\mathbb{E}(Y|Z_i)], \quad V_i = V[\mathbb{E}(Y|X_i)] \text{ et } V_i^C = V_i - V_i^U.$$

En pratique, l'estimateur de V_i se déduit comme

$$\widehat{V}_i = \frac{1}{n-1} \sum_{l=1}^n (y_{(i)}^l - \bar{y}_{(i)})^2, \quad \bar{y}_{(i)} = \frac{1}{n} \sum_{l=1}^n y_{(i)}^l,$$

avec $y_{(i)}^l = \widehat{\theta}_0 + \widehat{\theta}_i x_i^l$ lorsque le couple (θ_0, θ_i) est estimé par moindres carrés classiques. Les coefficients linéaires $(\gamma_0, (\gamma_j)_{j \neq i})$ dans (3.7) sont facilement déduits par moindres carrés. La quantité résiduelle du modèle (3.7), notée $\hat{Z}_i$ est définie comme

$$\hat{Z}_i = X_i - \left(\hat{\gamma}_0 + \sum_{j \neq i} \hat{\gamma}_j X_j \right).$$

A partir de cette estimation, la variance V_i^U est déduite comme,

$$\widehat{V}_i^U = \frac{1}{n-1} \sum_{l=1}^n (y_{(-i)}^l - \bar{y}_{(-i)})^2, \quad \bar{y}_{(-i)} = \frac{1}{n} \sum_{l=1}^n y_{(-i)}^l$$

$$y_{(-i)}^l = \widehat{\alpha}_0 + \widehat{\alpha}_i \hat{z}_i^l, \quad \hat{z}_i^l = x_i^l - \left(\hat{\gamma}_0 + \sum_{j \neq i} \hat{\gamma}_j x_j^l \right), \quad l \in [1 : n],$$

avec (α_0, α_i) estimé par moindres carrés.

Néanmoins, ce développement n'est applicable que si l'effet de chaque paramètre sur la réponse est linéaire, et que la régression linéaire donnée en (3.5) est une approximation valable du modèle. Si le modèle est fortement non-linéaire, une transformation du modèle pour le rendre linéaire est possible, mais cela engendre une modification de la variance, non contrôlée par la méthode présentée.

Pour remédier à ce problème, Li *et al.* [LRY⁺10] proposent une approche plus générale pour définir un indice de sensibilité global. Sous hypothèse que le modèle $Y = \eta(\mathbf{X})$ peut être approché par une *représentation du modèle à grande dimension* (HDMR) [LRR01, LRH⁺08], soit,

$$Y = \sum_{u \in S_d} \eta_u(\mathbf{X}_u) + \varepsilon, \quad \varepsilon \sim N(0, \sigma^2) \quad (3.9)$$

où S_d est un sous-ensemble de S tel que $d = |S_d| \ll |S| = 2^p$. De plus ε est décorrélée des fonctions η_u , $u \in S_d$. Par conséquent, la décomposition de la variance globale s'exprime comme

$$\begin{aligned} V(Y) &= V\left[\sum_{u \in S_d} \eta_u(\mathbf{X}_u) + \varepsilon\right] \\ &= \sum_{\substack{u \in S_d \\ u \neq \emptyset}} \left[V(\eta_u(\mathbf{X}_u)) + \sum_{v \neq u \in S_d} \text{Cov}(\eta_u(\mathbf{X}_u), \eta_v(\mathbf{X}_v)) \right] + \sigma^2. \end{aligned}$$

En supposant que σ^2 est négligeable, c'est-à-dire que la décomposition proposée (3.9) est une bonne approximation du modèle initial, la variance partielle associée à X_i , que l'on note comme précédemment V_i , connaît une décomposition similaire à celle donnée pour les méta-modèles linéaires proposés par Xu & Gertner [XG08]. Or, dans ce cas,

$$V_i^U = V[\eta_i(X_i)], \quad V_i^C = \sum_{v \neq i \in S_d} \text{Cov}[\eta_i(X_i), \eta_v(\mathbf{X}_v)], \quad \text{et} \quad V_i = V_i^U + V_i^C.$$

L'influence de X_i est quantifiée par une nouvelle définition de l'indice de sensibilité,

$$S_i^{\text{LR}} = \frac{V_i}{V(Y)} = \underbrace{\frac{V[\eta_i(X_i)]}{V(Y)}}_{S_i^U} + \underbrace{\frac{\text{Cov}[\eta_i(X_i), \sum_{v \neq i \in S_d} \eta_v(\mathbf{X}_v)]}{V(Y)}}_{S_i^C}. \quad (3.10)$$

La détermination des fonctions $(\eta_u)_{u \in S_d}$ se base sur la méthode RS-HDMR [LRH⁺08], méthode optimale de construction des composantes. Dans ce procédé, les composantes de plus haut ordre sont décomposées en des fonctions d'ordre plus petit de façon à obtenir

$$Y = \sum_{u \in S_k} f_u(\mathbf{X}_u) + \varepsilon, \quad S_k = \{u \in S_d, |u| \leq k\}, \quad k \sim 1, 2, 3.$$

Pour rendre le problème identifiable, les auteurs supposent que chaque composante f_u est centrée. Chaque fonction f_i est approchée par des splines cubiques tronquées $(B_r)_{r \in \mathbb{Z}}$ [KW70], et les fonctions f_{ij} , voire f_{ijk} , $i, j, k \in [1 : p]$, sont approchées par le produit tensoriel de bases de splines, soit

$$\begin{aligned} f_i(X_i) &\simeq \sum_{r=-1}^{m+1} \alpha_r^i B_r(X_i) \\ f_{ij}(X_i, X_j) &\simeq \sum_{p=-1}^{m+1} \sum_{r=-1}^{m+1} \alpha_{pr}^{ij} B_r(X_i) B_p(X_j) \\ f_{ijk}(X_i, X_j, X_k) &\simeq \sum_{q=-1}^{m+1} \sum_{p=-1}^{m+1} \sum_{r=-1}^{m+1} \alpha_{qpr}^{ijk} B_r(X_i) B_p(X_j) B_q(X_k). \end{aligned}$$

avec m le nombre de noeuds, et α_r^i , α_{pr}^{ij} , α_{qpr}^{ijk} sont les coefficients qu'il reste à estimer. Les auteurs ont recours à une procédure de *backfitting* pour estimer ces coefficients [HT84]. En quelques mots, le *backfitting* est une méthode de descente permettant d'estimer les composantes d'un modèle additif généralisé lorsqu'il existe des effets d'interaction entre les différentes composantes. A chaque itération de l'algorithme *backfitting*, il s'agit d'estimer chaque composante fonctionnelle f_u par critère des moindres carrés de la sortie observée par rapport à $\sum_{v \neq u} f_v$. La procédure est répétée jusqu'à ce que les compo-

santes convergent vers un minimum local. Il existe néanmoins des variantes à la technique RS-HDMR. Pour rendre cette méthode plus performante, Zuniga *et al.* [ZKS13] proposent de remplacer la technique Monte Carlo par une technique Quasi Monte Carlo. Les mêmes auteurs proposent de définir un choix optimal du degré m de polynômes et du nombre de points d'échantillonnage n . Ces techniques se basent sur la minimisation de la variabilité des paramètres.

Cette décomposition, beaucoup plus générale que la précédente, peut s'appliquer à n'importe quelle famille de méta-modèles. Dans un autre contexte, Caniou *et al.* [Can12] remplacent les bases de splines par des polynômes de chaos. Plus précisément, à partir d'une substitution de la fonction modèle par un polynôme de chaos, les auteurs proposent une décomposition par rapport aux marginales des entrées. A partir d'observations issues de la distribution de $\mathbf{X}$, les auteurs utilisent (3.10) pour estimer les indices de sensibilité basés sur la décomposition en chaos polynomial. Pour pallier les difficultés numériques liées à l'utilisation de polynômes de chaos, Zuniga *et al.* [ZKS13] rejoignent les travaux de Mara *et al.* [MT12] et proposent une décorrélation des entrées dans le cas particulier où la dépendance des entrées est modélisée par une copule gaussienne [LD09, Leb13]. Plus généralement, les mêmes auteurs reprennent les travaux de Soize & Ghanem [SG04], qui stipulent que les polynômes de chaos peuvent être caractérisés par une fonction copule. Même si le but premier des auteurs est d'estimer des modèles à entrées dépendantes, et non les termes d'une décomposition fonctionnelle, ces techniques peuvent néanmoins être exploitées ici.

Points importants

Les développements décrits dans cette section regroupent un certain nombre de décompositions ANOVA que l'on peut imaginer du modèle. Néanmoins, la décomposition de Bedford [Bed98] est dictée par l'ordre dans lequel les fonctions g_u , $u \in S$, ont été traitées dans le procédé d'orthogonalisation. Par conséquent, cela signifie qu'il existe $p!$ décompositions possibles. Lorsqu'il s'agit de reconstruction par méta-modèles [XG08, LRY⁺10], la décomposition est déterminée par le modèle de substitution de η , signifiant qu'elle pourra différer d'un méta-modèle à l'autre.

La contribution des travaux que nous présentons dès la troisième partie de ce manuscrit repose sur une autre forme de décomposition ANOVA que celle abordée ici. En effet, nous nous sommes inspirés de la décomposition initialement présentée par Stone [Sto94], puis reprise par Hooker [Hoo07], afin de proposer une décomposition de la fonction théorique η . Cette décomposition repose sur la projection de η sur un espace complet, dont la caractéristique principale est d'assurer une propriété d'orthogonalité dite *hiérarchique* de ses composantes. Cette contrainte d'orthogonalité permet d'assurer une unique décomposition ANOVA fonctionnelle de η . Il découle de cette décomposition une définition d'indices de sensibilité basés sur la décomposition de la variance globale du modèle. Pour pallier la difficulté numérique de l'estimation de tels indices, nous avons exploré l'approche par projections, comme cela sera détaillé dans le Chapitre 6. La méthode numérique qui en découle engendre des outils d'optimisation. Néanmoins, cette technique d'estimation, n'est utilisable que pour des modèles de variables dépendantes par paires. Pour traiter des modèles plus généraux, nous avons exploité l'estimation par méta-modèle pour proposer une construction empirique adaptée à la forme particulière des composantes de notre décomposition. Cette construction engendre cependant un fléau de la dimension, puisque le nombre de représentants à estimer croît de manière exponentielle, comme cela sera détaillé dans le Chapitre 7. Pour remédier à ce problème, nous nous sommes intéressés à des méthodes de sélection de variables. La Partie II constitue une introduction aux méthodes d'estimation d'indices présentées aux Chapitres 6 et 7, et a pour objectif de présenter les différentes techniques numériques utilisées en Partie III. Le Chapitre 4 fait état des techniques numériques de résolution de systèmes linéaires, exploités au Chapitre 6. Le Chapitre 5 introduit les aspects théoriques de la sélection de variables, et propose une résolution pratique des problèmes de grande dimension.

Deuxième partie

Méthodes d'optimisation et de
sélection de variables

Chapitre 4

Méthodes d'optimisation

Cette seconde partie a pour objectif d'exposer les méthodes d'optimisation numérique qui ont été exploitées pour l'article présenté au Chapitre 6.

D'une part, nous redonnons les techniques numériques classiques de résolution d'un système linéaire, qui constitue pour l'essentiel la partie estimation de notre contribution. D'autre part, la formule de Sherman-Morrison est présentée pour mettre en avant le gain numérique qu'elle permet d'obtenir sur une estimation d'espérance conditionnelle par polynômes locaux.

Ces méthodes étant très connues en analyse numérique, nous en ferons une description succincte ici.

4.1 Résolution de système linéaire

Dans cette section, nous rappelons les différentes techniques de résolution de système linéaire. En effet, beaucoup de problèmes statistiques conduisent à la résolution d'un système linéaire. Dans l'estimation des indices présentés, nous avons eu à résoudre un système linéaire fonctionnel. Pour cette raison, nous exposons ici brièvement les techniques classiques de résolution de systèmes linéaires, en précisant les motivations qui nous ont poussées à opter pour une méthode plutôt qu'une autre.

Soit le système linéaire suivant,

$$A\mathbf{x} = b, \tag{4.1}$$

où $A = (a_{ij})_{i,j \in [1:n]} \in \mathcal{M}_{n,n}(\mathbb{R})$, $b = {}^t(b_1, \dots, b_n) \in \mathbb{R}^n$ sont fixés, et $\mathbf{x} = {}^t(x_1, \dots, x_n) \in \mathbb{R}^n$ est le vecteur inconnu que l'on cherche à obtenir.

Pour résoudre (4.1), plusieurs techniques numériques existent, et s'exploitent selon les différentes propriétés que vérifient les quantités impliquées dans (4.1). Nous développons ici les méthodes classiques, classées en deux grandes

catégories : les méthodes directes et les méthodes itératives.

Méthodes directes

Considérons le cas où A est une matrice de plein rang, soit A est inversible. Dans ce cas, (4.1) admet une unique solution et $\mathbf{x} = A^{-1}b$. Cependant, cette solution ne s'obtient pas en calculant A^{-1} , puis en calculant $A^{-1}b$, car cela reviendrait à remplacer un *seul* système linéaire par n systèmes linéaires pour l'inversion de A , ce qui peut s'avérer être une opération très coûteuse en pratique [Cia98]. Lorsque A est supposée inversible, il faut donc avoir recours à des méthodes directes plus économiques.

Dans le cas où A est une matrice triangulaire inférieure (ou supérieure), la résolution du système $A\mathbf{x} = b$ est immédiate car,

$$\begin{cases} a_{11}x_1 = b_1 \\ a_{21}x_1 + a_{22}x_2 = b_2 \\ \vdots \\ a_{n1}x_1 + \cdots + a_{nn}x_n = b_n \end{cases}$$

Puisque $\prod_{i=1}^n a_{ii} = \det(A) \neq 0$ de par l'hypothèse d'inversibilité de A , il suffit de poser $x_1 = b_1/a_{11}$, puis d'intégrer x_1 dans la deuxième égalité pour en déduire x_2 , etc.,... pour en déduire la solution analytique $\mathbf{x}$ du système (4.1).

Si A n'est pas triangulaire, les méthodes directes utilisent les opérations élémentaires sur les matrices pour se ramener à des matrices triangulaires. La méthode de Gauss consiste à déterminer pas à pas une matrice M telle que MA soit triangulaire supérieure [Cia98]. Une variante de la méthode de Gauss est la méthode de Gauss-Jordan, qui s'emploie à trouver une matrice M telle que MA soit diagonale [Mey00]. Ces deux méthodes reposent sur la détermination d'éléments pivots, à partir desquels nous effectuons des opérations élémentaires pour obtenir M .

Si le choix des pivots correspond aux éléments diagonaux des matrices construites à chaque pas, alors la méthode de Gauss débouchera sur la factorisation $A = LU$, où L est une matrice triangulaire inférieure, et U une matrice triangulaire supérieure. Dans ce cas, la résolution du système (4.1) revient à résoudre en $(\mathbf{y}, \mathbf{x})$ les deux systèmes linéaires suivants,

$$L\mathbf{y} = b, \quad \text{et} \quad U\mathbf{x} = \mathbf{y}.$$

D'autres choix dans la méthode d'élimination de Gauss nous ramènent à la décomposition $A = QU$, où Q est une matrice unitaire c'est-à-dire $Q^*Q = I_n$.

Ainsi, les méthodes directes reposent essentiellement sur la méthode du pivot. Le coût de calcul de telles approches peut être très important pour des systèmes à très grande dimension. Les méthodes itératives que nous développons maintenant offrent généralement des coûts de calculs moindres.

Méthodes itératives

Les méthodes itératives sont des méthodes qui consistent à trouver une solution de (4.1) par approximations successives,

$$\mathbf{x}^{(k+1)} = B\mathbf{x}^{(k)} + \mathbf{z}, \quad k \in \mathbb{N}, \quad B \in \mathcal{M}_{n,n}(\mathbb{R}), \quad \mathbf{z} \in \mathbb{R}^n \quad (4.2)$$

Nous donnons ici l'algorithme décrivant la recherche d'une solution générale pour les méthodes de Jacobi, de Gauss-Seidel et de relaxation. En effet, ces trois méthodes sont toutes des méthodes itératives qui reposent sur la décomposition de A en une différence de matrices, $A = M - N$, où M est une matrice simple à inverser. L'Algorithme 1 décrit la démarche de résolution du système (4.2).

Algorithm 1: Algorithme générique de méthodes itératives

Input: A , $\mathbf{x}^{(0)}$ et $\varepsilon > 0$

Output: $\hat{x} = x^{(s)}$

Initialisation :

décomposition de $A = M - N$;

Poser $k = 1$;

Trouver $x^{(1)}$ tel que $Mx^{(1)} = Nx^{(0)} + b$;

while $(x^{(k-1)} - x^{(k)}) > \varepsilon$ **do**

 Résoudre $Mx^{(k+1)} = Nx^{(k)} + b$;

$k = k + 1$;

end

Posons

$$D = (a_{ii})_{i \in [1:n]}, \quad L = (-a_{ij})_{i > j \in [1:n]} \quad \text{et} \quad U = (-a_{ij})_{i < j \in [1:n]}.$$

Les méthodes de Jacobi, de Gauss-Seidel et de relaxation reposent toutes sur la technique décrite dans l'Algorithme 1, seule la décomposition de A diffère selon la méthode utilisée. La Table 4.1 résume les principales caractéristiques de chacune de ces méthodes.

L'avantage de la méthode de Gauss-Seidel par rapport à la méthode de Jacobi est qu'elle exige une capacité mémoire moindre. En effet, à chaque itération k , le calcul de $x_i^{(k)}$ de $\mathbf{x}^{(k)}$ nécessite les $n - i$ composantes $x_{i+1}^{(k-1)}, \dots, x_n^{(k-1)}$ de l'étape précédente, tandis que la méthode de Jacobi requiert la mémorisation de $n - 1$ éléments $x_j^{(k-1)}$, $j \neq i$. En ce sens, la méthode de Gauss-Seidel

Méthode	Décomposition $A = M - N$	Description d'une itération
Jacobi	$A = D - (L + U)$	$Dx^{(k+1)} = (L + U)x^{(k)} + b$
Gauss-Seidel	$A = (D - L) - U$	$(D - L)x^{(k+1)} = Ux^{(k)} + b$
Relaxation	$A = (\frac{D}{\omega} - L) - (\frac{1-\omega}{\omega}D + U),$ $\omega \neq 0$	$(\frac{D}{\omega} - L)x^{(k+1)} = (\frac{1-\omega}{\omega}D + U)x^{(k)} + b$

TABLE 4.1 – Caractéristiques des méthodes de Jacobi, Gauss-Seidel et de relaxation.

est plus économique. La méthode de relaxation coïncide avec la méthode de Gauss-Seidel pour $\omega = 1$. Néanmoins, comme nous allons le voir dans le Théorème 4.1, la vitesse de convergence dépend du rayon spectral de la matrice impliquée dans les approximations successives. Ainsi, l'injection de ω dans la méthode de relaxation contribue à augmenter la vitesse de convergence de la solution du système linéaire. Bien sûr, ce gain a un coût car la méthode de relaxation implique une procédure supplémentaire pour trouver la valeur optimale de ω [Cia98, Kre98]. Dans la suite, nous avons donc opté pour la méthode de Gauss-Seidel, qui offre un bon compromis entre simplicité et efficacité numérique. Nous énonçons cependant le Théorème 4.1 qui assure la convergence d'une méthode itérative en toute généralité.

Théorème 4.1. *Les assertions suivantes sont équivalentes :*

1. La méthode itérative $\mathbf{x}^{(k+1)} = B\mathbf{x}^{(k)} + \mathbf{z}$ est convergente ;
2. $\rho(B) < 1$, où ρ désigne le rayon spectral de B ;
3. $\|B\| < 1$ pour au moins une norme matricielle $\|\cdot\|$.

En particulier, nous énonçons des conditions suffisantes de convergence de la solution pour la méthode de Gauss-Seidel dans la Proposition 4.1.

Proposition 4.1. *Posons*

$$p := \max_{i \in [1:n]} p_i,$$

avec les quantités p_i définies récursivement comme

$$p_1 := \sum_{j=2}^n \left| \frac{a_{1j}}{a_{11}} \right|, \quad p_i := \sum_{j=1}^{i-1} \left| \frac{a_{ij}}{a_{ii}} \right| p_j + \sum_{j=i+1}^n \left| \frac{a_{ij}}{a_{ii}} \right|, \quad i \in [2:n]$$

Si

- $p < 1$ ou
- A est à diagonale dominante stricte, c'est-à-dire, pour tout $i \in [1 : n]$,
 $|a_{ii}| > \sum_{\substack{j=1 \\ j \neq i}}^n |a_{ji}|$ ou
- A est une matrice symétrique définie positive,

Alors, quelque soit $\mathbf{x}^{(0)}, b \in \mathbb{R}^n$, la suite $(\mathbf{x}^{(k)})_{k \in \mathbb{N}}$ de la méthode de Gauss-Seidel converge vers l'unique solution de $A\mathbf{x} = b$.

Note 4.1. La recherche d'une solution au problème linéaire $A\mathbf{x} = b$ peut aussi être vue comme la recherche de l'optimum,

$$\min_{\mathbf{x} \in \mathbb{R}^n} \|B\mathbf{x} - c\|^2, \quad A = {}^tBB, \quad b = {}^tBc$$

où $\|\cdot\|$ est la norme de la fonction euclidienne sur $\mathbb{R}^n$. Dans ce cadre, la méthode de gradient et ses variantes peuvent être exploitées, sous réserve que $A = {}^tBB$ soit définie positive, car alors le gradient de $f(\mathbf{x})$ est donné par $\nabla f(\mathbf{x}) = A\mathbf{x} - c$ [Kel87, Ber09].

4.2 Formule de Sherman-Morrison et polynômes locaux

La résolution du système linéaire par une méthode itérative implique l'estimation d'espérances conditionnelles. L'objectif de cette section est de brièvement rappeler la méthode d'estimation des espérances conditionnelles par la régression par polynômes locaux. Cependant, l'utilisation de cette méthode peut s'avérer très coûteuse, comme nous l'expliquerons par la suite. Nous utilisons la formule de Sherman-Morrisson pour réduire ce coût.

Nous supposons que nous désirons estimer la fonction

$$m(x) = \mathbb{E}(Y|X = x)$$

lorsque Y et X sont des variables aléatoires réelles. Nous supposons également que nous disposons d'un n -échantillon d'observations $(y^l, x^l)_{l=1, \dots, n}$ issues de la distribution de (Y, X) .

Pour estimer m , nous utilisons la méthode non paramétrique d'estimation par polynômes locaux [FG96]. L'estimation par polynômes locaux consiste à approcher $m(x)$ par une fonction polynomiale de degré q , et à estimer les

coefficients impliqués par la technique des moindres carrés pondérés,

$$\hat{M}(x) = \begin{pmatrix} \hat{m}(x) \\ \hat{m}'(x) \\ \dots \\ \frac{\hat{m}^{(q)}(x)}{q!} \end{pmatrix} = \underset{\substack{\beta_j \in \mathbb{R} \\ j \in [0:q]}}{\text{Argmin}} \sum_{l=1}^n [y^l - \sum_{j=0}^q \beta_j (x^l - x)^j]^2 K\left(\frac{x^l - x}{h}\right)$$

Plus précisément, l'estimation de $m(x)$ est donnée par

$$\hat{m}(x) = {}^t e_1 [{}^t \mathbb{X}(x) D(x) \mathbb{X}(x)]^{-1} \cdot {}^t \mathbb{X}(x) D(x) \mathbb{Y}$$

avec $\mathbb{Y}_l = y^l$, $l \in [1 : n]$, e_1 est le premier élément de la base canonique de $\mathbb{R}^{q+1}$, et

$$\mathbb{X}(x) = \begin{pmatrix} (x^1 - x)^0 & \dots & (x^1 - x)^q \\ \vdots & & \vdots \\ (x^n - x)^0 & \dots & (x^n - x)^q \end{pmatrix},$$

$$D(x) = \text{diag} \left\{ K\left(\frac{x^1 - x}{h}\right), \dots, K\left(\frac{x^n - x}{h}\right) \right\}.$$

$K(\cdot)$ est une fonction noyau, et h est un paramètre d'échelle. Le terme $K\left(\frac{x^l - x}{h}\right)$ est un terme de pondération qui permet de sélectionner les observations x^l assez proche de x pour intervenir dans le calcul du polynôme. Remarquons par ailleurs que si le degré du polynôme est établi à $q = 0$, l'estimateur $\hat{m}(x)$ coïncide avec l'estimateur de Nadaraya-Watson [Cyb09] donné par

$$\hat{m}^{\text{NW}}(x) = \frac{\sum_{l=1}^n y^l K\left(\frac{x^l - x}{h}\right)}{\sum_{l=1}^n K\left(\frac{x^l - x}{h}\right)}$$

Pour estimer et prédire m , la procédure de *leave-one-out* est utilisée, c'est-à-dire que pour toute observation x^l , on prédit $\hat{m}(x^l) = \hat{\mathbb{E}}(Y|X = x^l)$, espérance conditionnelle obtenue à partir de toutes les réalisations de X sauf la l -ième. Ainsi, pour tout $l \in [1 : n]$, nous devons calculer

$$\hat{m}(x^l) = {}^t e_1 [{}^t \mathbb{X}_{-l}(x^l) D_{-l}(x^l) \mathbb{X}_{-l}(x^l)]^{-1} \cdot {}^t \mathbb{X}_{-l}(x^l) D_{-l}(x^l) \mathbb{Y}_{-l}$$

où la l -ième ligne de chaque matrice a été retirée. Cela signifie que, pour chaque observation x^l , la matrice ${}^t \mathbb{X}_{-l}(x^l) D_{-l}(x^l) \mathbb{X}_{-l}(x^l)$ doit être inversée,

résultant de n inversions. Pour réduire le coût de calcul d'une telle opération, la formule de Sherman-Morrison-Woodbury [SM50] est adoptée. Pour $A \in \mathcal{M}_{n,n}(\mathbb{R})$ une matrice inversible, la formule de Sherman-Morrison-Woodbury est donnée par

$$(A + UCV)^{-1} = A^{-1} - A^{-1}U(C^{-1} + VA^{-1}U)^{-1}VA^{-1} \quad (4.3)$$

où $U \in \mathcal{M}_{n,k}(\mathbb{R})$, $C \in \mathcal{M}_{k,k}(\mathbb{R})$ est inversible et $V \in \mathcal{M}_{k,n}(\mathbb{R})$. Cette formule est dérivée de l'inversion des matrices par blocks, et exige donc que le complément de Schur $C^{-1} + VA^{-1}U$ soit inversible.

Dans l'estimation par polynômes locaux, posons $S_n(x) = {}^t\mathbb{X}(x)D(x)\mathbb{X}(x)$. $S_n(x)$ peut se réécrire comme

$$S_n(x) = \sum_{j=1}^n \Phi_j(x) {}^t\Phi_j(x) = \underbrace{\sum_{j \neq l} \Phi_j(x) {}^t\Phi_j(x)}_{S_{-l}(x)} + \Phi_l(x) {}^t\Phi_l(x)$$

avec

$$\Phi_l(x) = {}^t(K(\frac{x^l - x}{h})\mathbb{X}_l(x)), \quad \mathbb{X}_l(x) = (1 \cdots (x^l - x)^q), \quad \forall l \in [1 : n]$$

$$S_{-l}(x) = {}^t\mathbb{X}_{-l}(x)D_{-l}(x)\mathbb{X}_{-l}(x)$$

$S_{-l}(x)$ est la quantité à inverser en chaque point d'observations. Nous pouvons réécrire

$$S_{-l}(x) = S_n(x) - \Phi_l(x) {}^t\Phi_l(x), \quad \forall l \in [1 : n]$$

En appliquant la formule d'inversion (4.3), nous avons alors

$$S_{-l}^{-1}(x) = S_n^{-1}(x) + \frac{S_n^{-1}(x)\Phi_l(x) {}^t\Phi_l(x)S_n^{-1}(x)}{1 - {}^t\Phi_l(x)S_n^{-1}(x)\Phi_l(x)}, \quad \forall l \in [1 : n]$$

Sous réserve que $S_n(x)$ et ${}^t\Phi_l(x)S_n^{-1}(x)\Phi_l(x) \neq 1$ en les points d'observations $(\mathbf{x}^l)_{l \in [1:n]}$, la formule de Sherman-Morrison-Woodbury permet de réduire un problème d'inversion de n matrices à un problème d'inversion d'une seule matrice $S_n(x)$.

Chapitre 5

Méthodes de sélection de variables

5.1 Introduction

Cette partie est consacrée à la présentation des méthodes de sélection de variables, techniques exploitées dans le Chapitre 7 de la Partie III de ce manuscrit. Nous avons ici pour objectif de décrire quelques techniques utilisées en pratique lorsque nous sommes confrontés à la résolution d'un problème pour lequel une sélection parcimonieuse de composantes doit être réalisée. Cette sélection provient du fait que nous avons affaire à des modèles très coûteux en temps de calcul et/ou le nombre de termes à estimer est beaucoup plus important que le nombre d'observations. Pour cela, supposons que l'on dispose d'un système de fonctions $(\psi_k)_{k \in \mathbb{N}}$, et que nous désirons substituer le modèle d'intérêt η à la représentation suivante,

$$\sum_{k \in \Lambda} \beta_k \psi_k, \quad \Lambda \subset \mathbb{N}.$$

La théorie de l'approximation se classe en deux catégories [DeV98] :

- L'approximation linéaire consiste à projeter η dans l'espace $\mathcal{H}_m = \text{Span}\{\psi_k, k \in \Lambda_m := [1 : m]\}$, c'est-à-dire que η est approchée par une combinaison linéaire de m termes fixés du système $(\psi_k)_k$.
- L'approximation non linéaire consiste à projeter η dans l'espace $\mathcal{G}_m = \text{Span}\{\psi_k, k \in \Lambda_m \subset \mathbb{N}, |\Lambda_m| \leq m\}$. Contrairement à $\mathcal{H}_m$, $\mathcal{G}_m$ est un espace non linéaire, d'où l'appellation d'approximation non linéaire. Dans ce cas, on parle d'une approximation de η en m termes. Face à une telle approche, il est naturel de s'interroger sur la manière de choisir la "meilleure" approximation en m termes, c'est-à-dire comment sélectionner un sous-ensemble Λ_m de $\mathbb{N}$ tel que η soit bien approchée par une combinaison linéaire de $(\psi_k)_{k \in \Lambda_m}$. C'est la question abordée ici.

Nous travaillons ici avec le modèle de régression suivant

$$y^l = \eta(x^l) + \sigma\epsilon^l, \quad l \in [1 : n]. \quad (5.1)$$

où $y^l, x^l \in \mathbb{R}$ et $\mathbb{E}(\epsilon^l) = 0$, $\mathbb{E}((\epsilon^l)^2) = 1$. Puisque nous ne disposons pas d'un système $(\psi_k)_k$ de taille infinie en pratique, nous supposons que nous avons à notre disposition un dictionnaire de fonctions $\mathcal{D} = \{\psi_1, \dots, \psi_P\}$ de taille finie $P \gg m$. Dans ce contexte, les méthodes d'approximation non linéaires peuvent être vues comme des techniques de sélection de variables, où l'on cherche à réaliser une sélection parcimonieuse de composantes $(\psi_k)_{k \in [1:P]}$ pour l'estimation du modèle. Beaucoup de méthodes numériques existent pour résoudre ce problème dans divers domaines scientifiques [MZ93, LC09]. Les principales motivations de tels développements proviennent du fait que certaines fonctions du système considéré peuvent fournir de l'information redondante ou non nécessairement pertinente. Dans le cas où $P \gg n$, l'objectif est de retenir un très faible nombre de composantes informatives, et ainsi, se ramener à un problème bien défini, c'est-à-dire $n > p$.

Les procédures de sélection de variables classiques sont des méthodes de régularisation,

$$\min_{\beta_k \in \mathbb{R}} \frac{1}{n} \sum_{l=1}^n [y^l - \sum_{k \in [1:P]} \beta_k \psi_k(x^l)]^2 + \lambda J(\beta_1, \dots, \beta_P), \quad (5.2)$$

où $J(\cdot)$ est à valeurs positives en $(\beta_1, \dots, \beta_P)$, et $\lambda \geq 0$ est un paramètre de régularisation. Sous forme matricielle, (5.2) est équivalent à

$$\min_{\beta \in \mathbb{R}^P} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 + \lambda J(\beta), \quad (5.3)$$

avec $(\mathbb{Y})_l = y^l$, $(\mathbb{X})_{lj} = \psi_j(x^l)$, $(\beta)_j = \beta_j$ pour tout $l \in [1 : n]$, $j \in [1 : P]$ et $\|z\|_n^2 = \frac{1}{n} \sum_{l=1}^n (z^l)^2$.

Dès lors, la question de sélection réside dans le choix de la fonction J . Dans la suite, nous abordons les deux formes de pénalité les plus usuelles. Elles feront partie de notre contribution.

Pénalisation ℓ_0

L'approche la plus intuitive pour procéder à une sélection parcimonieuse est d'utiliser la régularisation ℓ_0 , c'est-à-dire

$$J(\beta) = \|\beta\|_0, \quad \|\beta\|_0 = \sum_{k \in [1:P]} \mathbb{I}(\beta_k \neq 0).$$

Dans [Mey12], les auteurs montrent que la résolution du problème (5.3) avec $J(\beta) = \|\beta\|_0$ revient à calculer une solution $\hat{\beta}$ pour toutes les collections possibles de modèles, puis à choisir le "meilleur" modèle parmi l'ensemble des collections estimées. Ceci peut être montré d'un point de vue mathématique. Soit $\mathcal{S}_u$, $u \subset [1 : P]$, l'espace vectoriel engendré par la base canonique $\{e_j, j \in u\}$ de $\mathbb{R}^{|u|}$. On a,

$$\inf_{\beta} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 + \lambda \|\beta\|_0 = \inf_u \inf_{\beta \in \mathcal{S}_u} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 + \lambda |u|. \quad (5.4)$$

Ce développement suggère donc de calculer une solution $\hat{\beta}$ pour tous les éléments possibles de $[1 : P]$. Cependant, cette démarche nécessite 2^P estimations, ce qui rend le problème de minimisation (5.4) très difficile à résoudre quand P est élevé. Pour remédier à ce problème, les méthodes dites *greedy* sont des méthodes d'approximation qui offrent de bonnes stratégies pour gérer la pénalisation ℓ_0 . Ces méthodes seront étudiées dans la Section 5.2.

Pénalisation ℓ_1

Cependant, la pénalisation ℓ_0 n'est pas convexe, et souffre d'instabilités statistiques [Bre95, Tib96]. En remarquant que $\|\beta\|_0 = \lim_{q \rightarrow 0} \|\beta\|_q^q$, avec $\|\beta\|_q^q = \sum_k |\beta_k|^q$, on peut envisager de remplacer la pénalisation ℓ_0 par une pénalisation ℓ_q , pour q proche de 0. Soit considérer

$$\min_{\beta} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 + \lambda \|\beta\|_q^q, \quad q > 0. \quad (5.5)$$

étant donné que $q = 1$ est la valeur la plus proche de 0 tel que le problème de minimisation (5.5) soit convexe, Tibshirani [Tib96] propose d'étudier (5.5) avec la pénalisation ℓ_1 , minimisation aussi connue sous le nom de régression Lasso, dont l'estimateur est donné par

$$\hat{\beta}(\lambda) \in \underset{\beta \in \mathbb{R}^P}{\text{Argmin}} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 + \lambda \|\beta\|_1, \quad \|\beta\|_1 = \sum_{k \in [1:P]} |\beta_k|. \quad (5.6)$$

Dans le cas où $\mathbb{X}$ est une matrice orthogonale, c'est-à-dire ${}^t\mathbb{X}\mathbb{X} = nI_P$, l'estimateur Lasso peut être explicitement calculé par les conditions KKT [Mey12, BvdG11]. Pour tout $i \in [1 : P]$, on a,

$$\hat{\beta}_i(\lambda) = \begin{cases} \tilde{\beta}_i - \lambda/2 & \text{si } \tilde{\beta}_i > \lambda/2 \\ \tilde{\beta}_i + \lambda/2 & \text{si } \tilde{\beta}_i < -\lambda/2 \\ 0 & \text{sinon,} \end{cases}$$

où $\tilde{\beta} = \underset{\beta \in \mathbb{R}^P}{\operatorname{Argmin}} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2 = {}^t\mathbb{X}\mathbb{Y}$ est l'estimateur des moindres carrés ordinaires.

Notons tout d'abord que la pénalisation ℓ_1 provoque un seuillage doux des estimations des coefficients obtenus par la méthode des moindres carrés ordinaires. En particulier, les coefficients β_j tels que $|\tilde{\beta}_j| \leq \lambda/2$ sont estimés à 0, donnant lieu à une sélection de variables qui dépend du paramètre λ . Plus λ sera grand, plus le nombre de coefficients de régression nuls sera important. Au contraire, si $\lambda = 0$, il n'y a pas de sélection, et la régression Lasso coïncide exactement avec les moindres carrés ordinaires. De plus, notons que cette solution permet de voir que le chemin de régularisation $\{\hat{\beta}(\lambda), \lambda \in \mathbb{R}^+\}$ est continu et linéaire par morceaux. Cette remarque permet d'en déduire que l'ensemble des solutions de la pénalisation ℓ_1 ne nécessite que le calcul d'une collection restreinte de ces solutions aux points de discontinuité du chemin de régularisation [Mey12].

Dans le cas où $\mathbb{X}$ n'est pas une matrice orthogonale, l'ensemble des solutions de (5.6) n'admet plus une solution explicite. En revanche, pour un paramètre $\lambda \geq 0$ fixé, il est aisé de montrer que (5.6) admet une unique ou une infinité de solutions [Tib12]. L'ensemble des solutions lasso $\{\hat{\beta}(\lambda)\}$ satisfait les conditions d'optimalité KKT suivantes [Tib12],

$$\begin{aligned} {}^t\mathbb{X}(\mathbb{Y} - \mathbb{X}\hat{\beta}(\lambda)) &= \lambda\gamma, \\ \gamma_i &\in \begin{cases} \operatorname{sign}(\hat{\beta}_i(\lambda)) & \text{si } \hat{\beta}_i(\lambda) \neq 0 \\ [-1, 1] & \text{si } \hat{\beta}_i(\lambda) = 0 \end{cases} \quad i \in [1 : P], \end{aligned} \quad (5.7)$$

où $\operatorname{sign}(a) = \pm 1$ est le signe de a . Pour $\lambda > 0$, on pose $E = \{i \in [1 : P], \left| {}^t\mathbb{X}(:, i)(\mathbb{Y} - \mathbb{X}\hat{\beta}(\lambda)) \right| = \lambda\}$, et $s = \operatorname{sign}({}^t\mathbb{X}(:, E)(\mathbb{Y} - \mathbb{X}\hat{\beta}(\lambda)))$, où $\mathbb{X}(:, i)$ dénote la i ème colonne de $\mathbb{X}$, et $\mathbb{X}(:, E) = (\mathbb{X}(:, i))_{i \in E}$. À partir des conditions KKT (5.7), une forme plus explicite de l'ensemble des solutions $\hat{\beta}(\lambda)$ est donnée par

$$b \in \ker(\mathbb{X}(:, E)), \quad \operatorname{sign}(\hat{\beta}_i(\lambda)) \cdot [\mathbb{X}(:, E)^+(\mathbb{Y} - ({}^t\mathbb{X}(:, E))^+ \lambda s + b_i)] \geq 0, \quad (5.8)$$

avec

$$\ker(\mathbb{X}(:, E)) = \{b \in \mathbb{R}^{|E|}, \mathbb{X}(:, E)b = 0\}, \quad \mathbb{X}(:, E)^+ = ({}^t\mathbb{X}(:, E)\mathbb{X}(:, E))^- \cdot {}^t\mathbb{X}(:, E),$$

où A^- est le pseudo inverse de Moore-Penrose de A . Par conséquent, pour $\lambda > 0$ fixé, la solution lasso est unique si et seulement si $\ker(\mathbb{X}(:, E)) = \{0\}$. Dans ce cas, la solution est donnée par l'expression suivante [Tib12],

$$\hat{\beta}_{-E} = 0, \quad \hat{\beta}_E = ({}^t\mathbb{X}(:, E)\mathbb{X}(:, E))^{-1}({}^t\mathbb{X}(:, E)\mathbb{Y} - \lambda s), \quad (5.9)$$

où $\hat{\beta}_E$ (respectivement $\hat{\beta}_{-E}$) est un sous-vecteur de $\hat{\beta}$ dont l'indexation des composantes est dans E (resp. $-E$). Toutefois, des conditions moins restrictives existent pour assurer l'unicité de la solution lasso. Par exemple, si les entrées de $\mathbb{X}$ sont distribuées par rapport à une mesure continue, l'unicité est également assurée presque sûrement [Tib12].

La forme de la solution (5.9) permet de voir que le chemin de régularisation $\hat{\beta}(\lambda)$ est linéaire et continue par morceaux comme fonction de λ . Ainsi, sous les conditions assurant l'unicité de la solution lasso pour un $\lambda > 0$ fixé, le chemin de régularisation $\{\hat{\beta}(\lambda), \lambda \geq 0\}$ est l'ensemble des solutions de la pénalisation ℓ_1 . Il suffit donc de trouver les valeurs de λ aux points de discontinuité du chemin pour en déduire l'ensemble des solutions lasso. Nous précisons ici les propriétés des solutions lasso dans le but de comprendre les approches numériques auxquelles nous allons nous intéresser à la Section 5.3 de cette partie.

Pour plus de précisions sur le sujet, le lecteur pourra se référer aux ouvrages très complets de Bühlmann *et al.* [BvdG11] et de Hastie *et al.* [HTF09]. Nous renvoyons également le lecteur aux travaux de Blatman et Touzani [Bla09, Tou11], qui exploitent la pénalisation ℓ_1 dans le cadre de l'analyse de sensibilité et lorsque le modèle théorique est remplacé par un méta-modèle impliquant l'estimation paramétrique de coefficients.

5.2 L'approximation *greedy*

Les algorithmes *greedy* permettent une résolution numérique de (5.3) avec $J(\beta) = \|\beta\|_0$. À partir d'un dictionnaire $\mathcal{D} = \{\psi_1, \dots, \psi_P\}$, l'objectif de la méthode *greedy* est de sélectionner un nombre restreint d'éléments de $\mathcal{D}$ pour réduire la dimension du problème [Tem11]. Introduits dans le cadre de l'estimation statistique [BCDD08, LC09], les algorithmes *greedy* sont aussi très utilisés en traitement du signal [ADM97], en réseau de neurones [LBW95] afin de produire des modèles plus facilement interprétables. Ainsi, il existe de nombreuses versions d'algorithmes. Les versions les plus classiques sont les algorithmes *backward* et *forward*. Ces deux méthodes sont récursives, et ont pour but de sélectionner (pour la méthode *forward*) ou d'éliminer (pour la méthode *backward*) séquentiellement les prédicteurs qui ont le moins d'impact sur l'ajustement de la régression.

L'inconvénient majeur à utiliser l'une ou l'autre de ces techniques est qu'elles sont basées sur une minimisation locale et qu'elles ignorent l'optimalité globale. Plus précisément, à chaque pas d'un des deux algorithmes, un prédicteur n'est sélectionné (ou éliminé) que si son ajout (ou son retrait) dans le modèle permet de réduire l'erreur de prédiction par rapport à l'erreur estimée avant son ajout (ou son retrait). Cependant, cette démarche n'assure pas de corriger les erreurs de sélection (ou d'élimination) de prédicteurs faites antérieurement. L'idée qui sera développée en Section 5.2.3 et au Chapitre 7 de la Partie III est de mixer les deux algorithmes. Cette approche a été proposée par Zhang [Zha11].

Les Sections 5.2.1, 5.2.2 et 5.2.3 détaillent les algorithmes *forward*, *backward* et *forward-backward*.

5.2.1 L'approche *forward*

Comme dans [CB00], nous notons Ξ l'ensemble des indices de colonnes de $\mathbb{X}$, et Γ un sous-ensemble de Ξ . La notation $\mathbb{X}(:, \Gamma)$ est utilisé pour désigner une sous-matrice de $\mathbb{X}$ dont les indices de colonnes sont dans Γ . $\mathbf{Z}_\Gamma$ dénote le sous-vecteur de $\mathbf{Z}$ dont les indices appartiennent à Γ . Soit $R(\Gamma)$ le résidu défini comme

$$R(\Gamma) = \min_{\beta_\Gamma} \|\mathbb{Y} - \mathbb{X}(:, \Gamma)\beta_\Gamma\|_n^2.$$

À partir de $\mathbb{X}(:, \emptyset) := 0$, l'algorithme *forward* consiste à successivement ajouter une colonne à $\mathbb{X}$ de manière à ce que le nouveau résidu $R(\Gamma_{\text{new}})$ soit plus petit que celui du pas précédent. La procédure s'arrête dès que l'erreur du modèle est raisonnablement petite. L'Algorithme 2 décrit la technique *forward*.

Algorithm 2: Algorithmme *forward*

Input: $\mathbb{X}$, $\mathbb{Y}$ et $\varepsilon > 0$
Output: $\Gamma^{(s)}$, $\hat{\beta}_{\Gamma^{(s)}}$
Initialisation : $\Gamma^{(0)} = \emptyset$, $R^{(0)} = \|\mathbb{Y}\|_{\mathbf{n}}^2$;
 $k = 1$;
 $j^{(1)} = \underset{j}{\operatorname{Argmin}} R(\Gamma^{(0)} \cup \{j\})$;
 $R^{(1)} = R(\Gamma^{(0)} \cup \{j^{(1)}\})$;
while $(R^{(k-1)} - R^{(k)}) > \varepsilon$ **do**
 $k = k + 1$;
 $j^{(k)} = \underset{j}{\operatorname{Argmin}} R(\Gamma^{(k-1)} \cup \{j\})$;
 $\Gamma^{(k)} = \Gamma^{(k-1)} \cup \{j^{(k)}\}$;
 $R^{(k)} = R(\Gamma^{(k)})$;
end
Solution :
 $\Gamma^{(s)}$;
 $\hat{\beta}_{\Gamma^{(s)}} = \underset{\beta_{\Gamma^{(s)}}}{\operatorname{Argmin}} \|\mathbb{Y} - \mathbb{X}(:, \Gamma^{(s)})\beta_{\Gamma^{(s)}}\|_{\mathbf{n}}^2$.

5.2.2 L'approche *backward*

A partir du modèle complet, la méthode *backward* consiste à éliminer pas à pas le prédicteur le moins informatif jusqu'à atteindre un critère d'arrêt donné. Le critère d'arrêt de l'algorithme peut être l'obtention d'une erreur minimale du modèle, ou un nombre fixé r de prédicteurs à retenir. L'Algorithme 3 présente une description de cette stratégie, stoppée lorsque la sélection de r prédicteurs est atteinte.

5.2.3 L'approche *forward-backward*

Nous reprenons ici les notations utilisées dans les Sections 5.2.1 et 5.2.2 pour décrire la méthode *forward-backward*, abrégée *FoBa* [Zha11].

A partir de $\mathbb{X}(:, \emptyset) := 0$, la méthode *FoBa* consiste à ajouter une colonne à $\mathbb{X}$ à chaque pas, exactement comme dans l'Algorithme 2. Cependant, un prédicteur de $\mathbb{X}(:, \Gamma_{\text{new}})$ est consécutivement éliminé s'il est jugé non informatif en comparaison avec les autres prédicteurs. L'Algorithme 4 donne précisément les critères de sélection/d'élimination de prédicteurs dans le modèle.

Par minimisation du résidu courant, un prédicteur entre dans le modèle (étape *forward*). Consécutivement, si en enlevant un prédicteur déjà présent

Algorithm 3: Algorithme *backward***Input:** $\mathbb{X}, \mathbb{Y}$ et $r \in \mathbb{N}^*$ **Output:** $\Gamma^{(r)}, \hat{\beta}_{\Gamma^{(r)}}$ *Initialisation* : $k = P, \Gamma^{(P)} = \{1, \dots, P\}, R^{(P)} = R(\Gamma^{(P)})$;**while** $k > r$ **do**

$$\begin{aligned} & k = k - 1; \\ & j^{(k)} = \underset{j \in \Gamma^{(k+1)}}{\text{Argmin}} R(\Gamma^{(k+1)} \setminus \{j\}); \\ & \Gamma^{(k)} = \Gamma^{(k+1)} \setminus \{j^{(k)}\}; \\ & R^{(k)} = R(\Gamma^{(k)}); \end{aligned}$$
end*Solution* :
$$\begin{aligned} & \Gamma^{(r)}; \\ & \hat{\beta}_{\Gamma^{(r)}} = \underset{\beta_{\Gamma^{(r)}}}{\text{Argmin}} \|\mathbb{Y} - \mathbb{X}(:, \Gamma^{(r)})\beta_{\Gamma^{(r)}}\|_{\text{n}}^2. \end{aligned}$$

du modèle, le résidu est significativement réduit en comparaison avec le résidu courant, un prédicteur sera éliminé du modèle (étape *backward*). Le critère d'arrêt ε équivaut au paramètre de régularisation λ dans des modèles de pénalisation convexe. Ainsi, si ε est très petit, l'algorithme *FoBa* aura tendance à sélectionner un grand nombre de fonctions. Au contraire, la méthode est agressive, c'est-à-dire qu'elle opère une sélection drastique, si ε est élevé. Le second critère δ est un scalaire de $]0, 1]$. Si $\delta > 1/2$, l'algorithme aura tendance à facilement réduire la dimension du modèle. Si $\delta < 1/2$, le contraire s'opérera. Ces deux paramètres peuvent être obtenus par validation croisée, bien que cette technique représente un coût numérique qui est à prendre en considération dans des modèles de très grande dimension [CB00, Zha11].

L'avantage de ces différentes approches est qu'elles sont très simples à implémenter, et très peu coûteuses en temps de calcul. Ces procédures permettent une sélection très parcimonieuse du nombre de prédicteurs informatifs, en estimant simultanément les coefficients associés. Cependant, ces méthodes souffrent d'instabilité numérique et le manque de résultats théoriques sur leur convergence en font des outils à utiliser avec précaution. Néanmoins, nous montrerons qu'ils peuvent apporter des résultats satisfaisants en pratique. Dans la Section 5.3, nous proposons une alternative via la technique

LARS adaptée à la régression Lasso.

Algorithm 4: Algorithme *FoBa*

Input: $\mathbb{X}, \mathbb{Y}, \varepsilon > 0$ et $\delta \in]0, 1]$
Output: $\Gamma^{(s)}, \hat{\beta}_{\Gamma^{(s)}}$

Initialisation : $\Gamma^{(0)} = \emptyset, R^{(0)} = \|\mathbb{Y}\|_n^2$;
 $k = 1$;
 $j^{(1)} = \underset{j}{\operatorname{Argmin}} R(\Gamma^{(0)} \cup \{j\})$;
 $R^{(1)} = R(\Gamma^{(0)} \cup \{j^{(1)}\})$;

while $(R^{(k-1)} - R^{(k)}) > \varepsilon$ **do**
 $\Gamma^{(k)} = \Gamma^{(k-1)} \cup \{j^{(k)}\}$;
 // Etape *backward*
 $m^{(k)} = \underset{m}{\operatorname{Argmin}} R(\Gamma^{(k)} \setminus \{m\})$;
 soit $d^+ = R^{(k-1)} - R^{(k)}$ et $d^- = R(\Gamma^{(k)} \setminus \{m^{(k)}\}) - R^{(k)}$;
 if $d^- \leq \delta d^+$ **then**
 $\Gamma^{(k)} = \Gamma^{(k)} \setminus \{m^{(k)}\}$
 end
 // Etape *forward*
 $k = k + 1$;
 $j^{(k)} = \underset{j}{\operatorname{Argmin}} R(\Gamma^{(k-1)} \cup \{j\})$;
 $R^{(k)} = R(\Gamma^{(k-1)} \cup \{j^{(k)}\})$;
end

Solution :
 $\Gamma^{(s)}$;
 $\hat{\beta}_{\Gamma^{(s)}} = \operatorname{Argmin}_{\beta \in \mathbb{R}^P} \|\mathbb{Y} - \mathbb{X}(:, \Gamma^{(s)})\beta\|_n^2$.

5.3 Algorithmes pour régression Lasso

Comme précédemment évoqué, les méthodes de sélection via la pénalisation ℓ_0 peuvent être extrêmement variables de par leur caractère discret. Pour remédier à ce problème, la pénalisation ℓ_1 , initiée par Tibshirani [Tib96] est un bon compromis entre une sélection drastique de prédicteurs, et une régression ridge, connue pour apporter une solution stable de par son caractère continu [HTF09]. Pour estimer les paramètres Lasso, plusieurs méthodes pratiques ont été développées. Tibshirani [Tib96] a d'abord montré que la régression lasso (5.6) est équivalente à

$$\hat{\beta}(\lambda) \in \operatorname{Argmin}_{\beta \in \mathbb{R}^P} \|\mathbb{Y} - \mathbb{X}\beta\|_n^2, \quad \sum_{k \in [1:P]} |\beta_k| \leq t, \quad t \geq 0 \quad (5.10)$$

où il existe une relation bijective entre λ et t . Par conséquent, pour une valeur fixée de t , des techniques de quadrature peuvent être utilisées pour résoudre (5.10) [AS65]. Le paramètre t peut être estimé par validation croisée, ou estimation du risque, comme relaté dans [Tib96].

Cette technique a cependant été supplantée par d'autres algorithmes, qui s'avèrent être plus efficaces. Une liste d'algorithmes non exhaustive a notamment été réalisée par Schmidt [Sch05], auquel le lecteur pourra se référer. Parmi ces techniques, l'algorithme LARS, développé par Efron *et al.* [EHJT04], s'inspire des propriétés des estimateurs relatés plus haut pour proposer un algorithme efficace.

Dans la suite, nous introduisons d'abord l'algorithme *forward stagewise regression* [Wei80], qui constitue le point de départ de l'algorithme LARS. Nous évoquerons les motivations de Efron *et al.* à s'inspirer de cette méthode pour construire leur algorithme. Nous expliciterons ensuite la technique LARS et sa variante pour l'estimation des paramètres Lasso.

5.3.1 Algorithme *forward stagewise regression*

Ce procédé itératif permet d'identifier et d'estimer les prédictors les plus corrélés avec le résidu du modèle. Le principe de cette technique itérative est d'identifier à chaque pas le prédicteur le plus corrélé avec le résidu, puis d'en modifier sa direction jusqu'à obtenir un second prédicteur qui sera à son tour modifié. L'algorithme *forward stagewise regression* est décrit via l'Algorithme 5. Cette simple méthode constitue l'idée principale de l'algorithme LARS.

Algorithm 5: Algorithme *forward Stagewise Regression*

Input: $\mathbb{X}$, $\mathbb{Y}$ et $\varepsilon > 0$

Initialisation : $\hat{\beta}^{(0)} = 0$, $\hat{\mu}^{(0)} = \mathbb{X}\hat{\beta}^{(0)}$ et $\hat{c}^{(0)} = {}^t\mathbb{X}\mathbb{Y}$;

$k = 1$;

while $\hat{c}^{(k-1)} \neq 0$ **do**

$j^{(k)} = \underset{j}{\text{Argmax}} |\hat{c}_j^{(k-1)}|$;

$\hat{\beta}^{(k)} = \hat{\beta}^{(k-1)} + \varepsilon \cdot \text{sign}(\hat{c}_{j^{(k)}}^{(k-1)}) e_{j^{(k)}}$, avec

$(e_{j^{(k)}})_i = \begin{cases} 1 & \text{si } i = j^{(k)} \\ 0 & \text{sinon} \end{cases}$;

$\hat{\mu}^{(k)} = \mathbb{X}\hat{\beta}^{(k)}$;

$\hat{c}^{(k)} = {}^t\mathbb{X}(\mathbb{Y} - \hat{\mu}^{(k)})$;

$k = k + 1$;

end

5.3.2 Méthode LARS adaptée au Lasso

L'algorithme LARS est une version stylisée du *forward stagewise regression*. De même que dans l'Algorithme 5, la méthode LARS a pour but d'identifier à chaque pas le prédictor le plus corrélé avec le résidu. Mais contrairement à l'Algorithme 5, LARS établit une direction équiangulaire entre les différents prédictors retenus afin d'intégrer un nouveau prédictor dans l'ensemble actif des prédictors (c'est-à-dire l'ensemble des prédictors identifiés). L'algorithme commence à $\lambda = +\infty$, qui correspond à la solution triviale $\hat{\beta}(\lambda) = 0 \in \mathbb{R}^P$. Puis, à chaque itération k de l'algorithme, λ décroît progressivement jusqu'à ce qu'un $\lambda_k < \lambda$ soit sélectionné, et une nouvelle solution $\hat{\beta}(\lambda_k)$ est calculée successivement. La procédure se poursuit jusqu'à ce que $\lambda = 0$, c'est-à-dire jusqu'à ce que la solution des moindres carrés ordinaires soit obtenue. A chaque pas, le choix de λ_k est établi par le critère de corrélation maximale avec le résidu.

La modification de LARS pour la régression Lasso se base sur la caractérisation de la solution donnée en (5.8), et s'inspire de la méthode d'homotopie proposée par Osborne *et al.* [Os00]. En outre, elle repose sur le fait que, dans une régression Lasso, le signe des coefficients non nuls doit concorder avec le signe de corrélation des prédictors avec le résidu actif. Si ce n'est pas le cas, le prédictor sélectionné est retiré de l'ensemble actif des prédictors. L'Algorithme 6 décrit la technique LARS adaptée à la régression Lasso. Au préalable, nous définissons plusieurs quantités.

Notons $\hat{\mu} = \mathbb{X}\hat{\beta}$. Soit Γ un sous-ensemble de $\{1, \dots, P\}$, et $s_k = \pm 1$ le signe de $k \in \Gamma$. Définissons également,

$$X_\Gamma := (s_k \mathbb{X}_k \dots)_{k \in \Gamma}, \quad G_\Gamma = {}^t X_\Gamma X_\Gamma, \quad \text{and } A_\Gamma = ({}^t \mathbb{I}_\Gamma G_\Gamma^{-1} \mathbb{I}_\Gamma)^{-1/2},$$

où $\mathbb{I}_\Gamma$ est un vecteur de 1 de taille $|\Gamma|$. Aussi u_Γ est la direction équiangulaire définie comme

$$u_\Gamma = X_\Gamma w_\Gamma, \quad w_\Gamma = A_\Gamma G_\Gamma^{-1} \mathbb{I}_\Gamma.$$

Enfin, e_Γ désigne la matrice de signes de taille $P \times |\Gamma|$,

$$e_\Gamma := (\dots e_k \dots)_{k \in \Gamma}, \quad \text{où } (e_k)_i = \begin{cases} s_k & \text{si } i = k \in \Gamma \\ 0 & \text{sinon.} \end{cases}$$

En partant du principe que pour un $\lambda > 0$ donné, la minimisation lasso (5.6) admette une unique solution lasso $\hat{\beta}(\lambda)$, Efron *et al.* ont montré que la solution obtenue par l'algorithme LARS coïncide avec $\hat{\beta}(\lambda)$ [EHJT04, Tib12]. Ainsi, cet algorithme se révèle très performant pour reconstruire le chemin de régularisation $\{\hat{\beta}(\lambda), \lambda \in \mathbb{R}^+\}$. Il reste toutefois à déterminer le "meilleur" λ pour obtenir une unique solution de (5.6). Pour cela, nous nous sommes

Algorithm 6: Algorithme LARS adapté au Lasso**Input:** $\mathbb{X}, \mathbb{Y}$

Initialisation : $\hat{\beta}^{(0)} = 0$, $\hat{\mu}^{(0)} = \mathbb{X}\hat{\beta}^{(0)} = 0$, $\hat{c}^{(0)} = {}^t\mathbb{X}\mathbb{Y}$ et $\Gamma^{(0)} = \{\emptyset\}$
 l'ensemble actif;

 $k = 1$;**while** $k < \min\{P, n\}$ **do**1. Calculer les corrélations $\hat{C} = \max_j |\hat{c}_j^{(k-1)}|$; $\Gamma = \{j : |\hat{c}_j^{(k-1)}| = \hat{C}\}$; $s_j^{(k)} = \text{sign}(\hat{c}_j^{(k-1)})$, $j \in \Gamma$;Calculer X_Γ , A_Γ , u_Γ et e_Γ avec $s_j^{(k)}$ et Γ ;

2. Calculer une première direction // LARS classique

$$\hat{\gamma} = \min_{j \notin \Gamma} \left\{ \frac{\hat{C} - \hat{c}_j^{(k-1)}}{A_\Gamma - a_j}, \frac{\hat{C} + \hat{c}_j^{(k-1)}}{A_\Gamma + a_j} \right\};$$

$$\hat{j} = \underset{j \notin \Gamma}{\text{Argmin}} \left\{ \frac{\hat{C} - \hat{c}_j^{(k-1)}}{A_\Gamma - a_j}, \frac{\hat{C} + \hat{c}_j^{(k-1)}}{A_\Gamma + a_j} \right\} \text{ avec } a_j := ({}^t\mathbb{X} \cdot u_\Gamma)_j;$$

3. Calculer une autre direction // modification Lasso

$$\tilde{\gamma} = \begin{cases} \min_{\gamma_j > 0} \{\gamma_j\}, & \gamma_j = -\frac{\hat{\beta}_j^{(k-1)}}{d_j}, \quad d_j = \begin{cases} s_j^{(k)} w_{\Gamma_j} & \text{si } j \in \Gamma \\ 0 & \text{sinon} \end{cases} \\ +\infty & \text{si } \gamma_j \leq 0 \quad \forall j \end{cases};$$

$$\tilde{j} = \underset{j}{\text{Argmin}} \min_{\gamma_j > 0} \{\gamma_j\};$$

if $\tilde{\gamma} < \hat{\gamma}$ **then**

$$\Gamma^{(k)} = \Gamma^{(k-1)} \setminus \{\tilde{j}\};$$

$$\hat{\beta}^{(k)} = \hat{\beta}^{(k-1)} + \tilde{\gamma} e_\Gamma G_\Gamma^{-1} \mathbb{I}_\Gamma;$$

$$\hat{\mu}^{(k)} = \mathbb{X}\hat{\beta}^{(k)};$$

else

$$\Gamma^{(k)} = \Gamma^{(k-1)} \cup \{\hat{j}\};$$

$$\hat{\beta}^{(k)} = \hat{\beta}^{(k-1)} + \hat{\gamma} e_\Gamma G_\Gamma^{-1} \mathbb{I}_\Gamma;$$

$$\hat{\mu}^{(k)} = \mathbb{X}\hat{\beta}^{(k)};$$

end

$$\hat{c}^{(k)} = {}^t\mathbb{X}(\mathbb{Y} - \hat{\mu}^{(k)});$$

$$k = k + 1;$$

end

intéressés au critère de sélection proposé par Efron *et al.*. Imitant le C_p de Mallows [Sap06], les auteurs proposent de calculer, à chaque pas k de l'algorithme, le critère suivant,

$$E_k = \left\| \mathbb{Y} - \mathbb{X}\hat{\beta}(\lambda_k) \right\|_n^2 / \hat{\sigma}^2 - n + 2k$$

où λ_k est le paramètre de régularisation choisie à l'itération k , et $\hat{\sigma}^2$ est la variance estimée à partir du modèle des moindres carrés ordinaires. Le paramètre "optimal" $\hat{\lambda} = \lambda(\hat{k})$ est choisi de telle sorte que $\hat{k} = \underset{k}{\operatorname{Argmin}} E_k$.

Ainsi, la solution lasso retenue est $\hat{\beta}(\hat{\lambda})$.

Troisième partie

Indices généralisés pour variables dépendantes

Introduction to Chapter 6

The Hoeffding decomposition studied in Chapter 1 of Part I is a unique representation of the model function η as a sum of increasing dimension functions. These components represent the main effects and the interaction effects that describe how input variables interact with each other in the model. Under the hypothesis of independent inputs, i.e. $P_{\mathbf{X}} = \otimes_{i=1}^p P_{X_i}$, the functional decomposition is unique under the assumption that summands are mutually orthogonal. In sensitivity analysis, the Hoeffding decomposition leads to an exact variability quantification that highlights the most influent effects in the model. For these reasons, the ANOVA decomposition is an attractive formula, and has received much attention over the past decades in sensitivity analysis when inputs' independence are no more assumed. In a context of regression and density estimation, Stone proposes a generalized decomposition of η and studies the theoretical properties of such an approach [Sto94]. To get a well-defined decomposition, he imposes orthogonality conditions on the ANOVA components. Following the idea, Hooker unifies this work to treat high-dimensional models with dependent variables for machine learning [Hoo07].

Inspired by these papers, we revisit the generalized functional decomposition to quantify the variability in the model. The objective is to bring an exact and unambiguous definition of a variance-based sensitivity index under general assumptions on the joint distribution of the input variables. This contribution also aims at unifying the classical Sobol indices, deduced from the Hoeffding decomposition, with generalized sensitivity indices in presence of correlations among inputs. That is, if the joint distribution is characterized by the marginal distributions, our generalized indices are the classical Sobol ones. At last, the purpose of the forthcoming work is to propose a numerical method to estimate the new defined sensitivity indices for low-dimensional models.

Chapter 6 consists of the published article [CGP12]. The article is organized as follows.

Under weak conditions on the joint distribution function, we come back on the generalized functional decomposition proposed by Stone. We define ap-

appropriate Hilbert spaces, and we show, by completeness of the spaces, the uniqueness of the decomposition. Following the track of the Sobol indices construction (Chapter 1 of Part I), we define generalized sensitivity indices based on the global variance decomposition. Thus, as inputs are not necessarily independent, covariance components appear in the indices. When we come back to an independent framework, covariance terms disappear, and we obtain the Sobol indices. However, these results are only applicable under assumptions made on the dependence structure of the inputs. By using the concepts defined in Chapter 2 of Part I, we will try to clarify these hypothesis in terms of copulas.

At last, the last part aims at using the components' orthogonality properties to propose an estimation method. The procedure performs the resolution of a functional linear system involving projection operators. This approach is illustrated by several test cases.

Chapter 6

Generalized Hoeffding-Sobol decomposition for dependent variables - application to sensitivity analysis

Abstract: In this paper, we consider a regression model built on dependent variables. This regression modelizes an input output relationship. Under boundedness type assumptions on the joint density function of the input variables, we show that a generalized Hoeffding-Sobol decomposition is available. This leads to new indices measuring the sensitivity of the output with respect to the input variables. We also study and discuss the estimation of these new indices.

Keywords: Sensitivity index; Hoeffding decomposition; dependent variables; Sobol decomposition.

6.1 Introduction

Sensitivity analysis (SA) aims to identify the variables that most contribute to the variability into a non linear regression model. Global SA is a stochastic approach whose objective is to determine a global criterion based on the density of the joint probability distribution function of the output and the inputs of the regression model. The most usual quantification is the variance-based method, widely studied in SA literature. Hoeffding decomposition [Hoe48] (see also Owen [Owe94]) states that the variance of the output can be uniquely decomposed into summands of increasing dimensions under orthogonality constraints. Following this approach, Sobol [Sob93] introduces variability measures, the so called *Sobol sensitivity indices*. These indices quantify the contribution of each input on the system.

Different methods have been exploited to estimate Sobol indices. The Monte Carlo algorithm was proposed by Sobol [Sob01], and has been later improved by the Quasi Monte Carlo technique, performed by Owen [Owe05]. FAST methods are also widely used to estimate Sobol indices. Introduced earlier by Cukier *et al.* [CS75] [CLS78], they are well known to reduce the computational cost of multidimensional integrals thanks to Fourier transformations. Later, Tarantola *et al.* [TGM06] adapted the Random Balance Designs (RBD) to FAST method for SA (see also recent advances on the subject by Tissot *et al.* [TP12a]).

However, these indices are constructed on the hypothesis that the input variables are independent, which seems unrealistic for many real life phenomena. In the literature, only a few methods and estimation procedures have been proposed to handle models with dependent inputs. Several authors have proposed sampling techniques to compute marginal contribution of inputs to the outcome variance (see the introduction in Mara and references therein [MT12]). As underlined in Mara *et al.* [MT12], if inputs are not independent, the amount of the response variance due to a given factor may be influenced by its dependence to other inputs. Therefore, classical Sobol indices and FAST approaches for dependent variables are difficult to interpret (see, for example, Da Veiga's illustration [DV07] p.133). Xu and Gertner [XG08] proposed to decompose the partial variance of an input into a correlated part and an uncorrelated one. Such an approach allows to exhibit inputs that have an impact on the output only through their strong correlation with other incomes. However, they only investigated linear models with linear dependences.

Later, Li *et al.* [LRY⁺10] extended this approach to more general models, using the concept of High Dimensional Model Representation (HDMR [LRH⁺08]). HDMR is based on a hierarchy of component functions of increasing dimensions (truncation of Sobol decomposition in the case of independent variables). The component functions are then approximated by expansions in terms of some suitable basis functions (e.g., polynomials, splines, ...). This meta-modeling approach allows the splitting of the response variance into a correlative contribution and a structural one of a set of inputs. Mara *et al.* [MT12] proposed to decorrelate the inputs with the Gram-Schmidt procedure, and then to perform the ANOVA-HDMR of Li *et al.* [LRY⁺10] on these new inputs. The obtained indices can be interpreted as fully, partially correlated and independent contributions of the inputs to the output. Nevertheless, this method does not provide a unique orthogonal set of inputs as it depends on the order of the inputs in the original set. Thus, a large number of sets has to be generated for the interpretation of resulting indices. As a different approach, Borgonovo *et al.* [Bor07, BCT11] initiated the construction of a new generalized moment free sensitivity index. Based on geometrical consideration, these indices measure the shift area between the outcome density and this same density conditionally to a parameter. Thanks

to the properties of these new indices, a methodology is given to obtain them analytically through test cases. Recently, Kucherenko *et al.* [KTA12] proposed to use first order and total sensitivity indices based on the classical decomposition of total variance. These new indices are estimated by crude Monte Carlo method on conditional expectation through several numerical examples.

Notice that none of these works has given an exact and unambiguous definition of the functional ANOVA for correlated inputs as the one provided by Hoeffding-Sobol decomposition when inputs are independent. Consequently, the exact form of the model has neither been exploited to provide a general variance-based sensitivity measures in the dependent frame.

In a pioneering work, Hooker [Hoo07], inspired by Stone [Sto94], shed new lights on hierarchically orthogonal function decomposition.

We revisit here the work of Hooker and Stone. We obtain hierarchical functional decomposition under a general assumption on the inputs distribution. Furthermore, we also show the uniqueness of the decomposition leading to the definition of new sensitivity indices. Under suitable conditions on the joint distribution function of the input variables, we give a hierarchically orthogonal functional decomposition (HOFD) of the model. The summands of this decomposition are functions depending only on a subset of input variables and are hierarchically uncorrelated. This means that two of these components are orthogonal whenever all the variables involved in one of the summands also appear in the other. This decomposition leads to the construction of generalized sensitivity indices well tailored to perform global SA when the input variables are dependent. In the case of independent inputs, this decomposition is nothing more than the Hoeffding one. Furthermore, our generalized sensitivity indices are in this case the classical Sobol ones.

Recently, Li *et al.* [LR12] proposed a numerical method to compute the decomposition components given by Hooker. Using a constrained minimization of the squared error, the variational counterpart of the decomposition also given in [Hoo07], they estimate the decomposition with an extended basis by using a continuous descent technique.

Here, we propose an estimation method performed by solving linear system involving suitable projection operators. We will focus on the particular case where the inputs are independent pairs of dependent variables (IPDV). Firstly, in the simplest case of a single pair of dependent variables, the HOFD may be obtained by solving a functional linear system of equations (see Procedure 1). In the more general IPDV case, the HOFD is then obtained in two steps (see Procedure 2). The first step is a classical Hoeffding-Sobol decomposition of the output on the input pairs, as developed in Jacques *et al.* [JLD06]. The second step is the HOFDs of all the pairs. In practical situations, the non parametric regression function of the model is generally not exactly known. In this case, one can only have at hand some realizations of the model and have to estimate, with this information, the HOFD. Here, we

study this statistical problem in the IPDV case. We build estimators of the generalized sensitivity indices and study numerically their properties. One of the main conclusion is that the generalized indices have a total normalized sum. This is not true for classical Sobol indices in the frame of dependent variables.

The paper is organized as follows. In Section 6.2, we give and discuss general results on the HOFD. The main result is Theorem 6.1. We show here that a HOFD is available under a boundedness type assumption (C.2) on the density of the joint distribution function of the inputs. Further, we introduce the generalized indices. In Section 6.3, we give examples of multivariate distributions to which Theorem 6.1 applies. We also state a sufficient condition for (C.2) and necessary and sufficient conditions in the IDPV case. Section 6.4 is devoted to the estimation procedures of the components of the HOFD and of the new sensitivity indices. Section 6.5 presents numerical applications. Through three toy functions, we estimate generalized indices and compare their performances with the analytical values. In the first two examples, we compare the sensitivity indices estimations to the true values to show the relevance of our new indices. The last example is a more realistic model. To prevent an industrial site from inundation, the overflow of a river is modeled by a set of eight inputs, in which some pairs are linearly correlated. The goal of this last application is to detect the most influential variables taking into account the dependence of the variables in the physical model. In Section 6.6, we give conclusions and discuss future work. Technical proofs and further details are postponed to the Appendix A.

6.2 Generalized Hoeffding decomposition-application to SA

To begin with, let introduce some notation. We briefly recall the usual functional ANOVA decomposition, and Sobol indices. We then state a generalization of this decomposition, allowing to deal with correlated inputs.

6.2.1 Notation and first assumptions

We denote by $\subset$ the strict inclusion, that is $A \subset B \Rightarrow A \cap B \neq B$, whereas we use $\subseteq$ when equality is possible.

Let $(\Omega, \mathcal{A}, P)$ be a probability space and let Y be the output of a deterministic model η . Suppose that η is a measurable function of a random vector $\mathbf{X} = (X_1, \dots, X_p) \in \mathbb{R}^p$, $p \geq 1$ and let $P_{\mathbf{X}}$ be the pushforward measure of P by $\mathbf{X}$,

$$\begin{aligned} Y : \quad (\Omega, \mathcal{A}, P) &\rightarrow (\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}) \rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R})) \\ \omega &\mapsto \mathbf{X}(\omega) \mapsto \eta(\mathbf{X}(\omega)). \end{aligned}$$

Let ν be a σ -finite measure on $(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p))$. Assume that $P_{\mathbf{X}} \ll \nu$ and let $p_{\mathbf{X}}$ be the density of $P_{\mathbf{X}}$ with respect to ν , that is $p_{\mathbf{X}} = \frac{dP_{\mathbf{X}}}{d\nu}$.

Also, assume that $\eta \in L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. The associated inner product of this Hilbert space is:

$$\langle h_1, h_2 \rangle = \int h_1(\mathbf{x})h_2(\mathbf{x})p_{\mathbf{X}}d\nu(\mathbf{x}) = \mathbb{E}(h_1(\mathbf{X})h_2(\mathbf{X})).$$

Here $\mathbb{E}(\cdot)$ denotes the expectation. The corresponding norm will be classically denoted by $\|\cdot\|$. Further, $V(\cdot) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))^2]$ denotes the variance, and $\text{Cov}(\cdot, *) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))(* - \mathbb{E}(*))]$ the covariance.

Let $[1 : p] := \{1, \dots, p\}$ and S be the collection of all subsets of $[1 : p]$. Define $S^- := S \setminus [1 : p]$ as the collection of all subsets of $[1 : p]$ except $[1 : p]$ itself. The cardinality of a set u is denoted by $\#(u)$ or $|u|$ according to the context.

Further, let $\mathbf{X}_{\mathbf{u}} := (X_l)_{l \in u}$, $u \in S \setminus \{\emptyset\}$. We introduce the subspaces of $L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$ $(H_u)_{u \in S}$, $(H_u^0)_{u \in S}$ and H^0 . H_u is the set of all measurable and square integrable functions depending only on $\mathbf{X}_{\mathbf{u}}$. $H_{\emptyset}$ is the set of constants and is identical to $H_{\emptyset}^0$. H_u^0 , $u \in S \setminus \emptyset$, and H^0 are defined as follows:

$$H_u^0 = \{h_u(\mathbf{X}_{\mathbf{u}}) \in H_u, \langle h_u, h_v \rangle = 0, \forall v \subset u, \forall h_v \in H_v^0\},$$

$$H^0 = \left\{ h(\mathbf{X}) = \sum_{u \in S} h_u(\mathbf{X}_{\mathbf{u}}), h_u \in H_u^0 \right\}.$$

At this stage, we do not make assumptions on the support of $\mathbf{X}$.

6.2.2 Sobol sensitivity indices

In this section, we recall the classical Hoeffding-Sobol decomposition, and the Sobol sensitivity indices if the inputs are independent, that is when $P_{\mathbf{X}} = P_{X_1} \otimes \dots \otimes P_{X_p}$.

The usual presentation is done when $\mathbf{X} \sim \mathcal{U}([0, 1]^p)$ [Sob93], but the Hoeffding decomposition remains true in the more general case of independent variables [VDV98].

Let $\mathbf{x} = (x_1, \dots, x_p) \in \mathbb{R}^p$. The decomposition consists in writing $\eta(\mathbf{x}) = \eta(x_1, \dots, x_p)$ as the sum of increasing dimension functions:

$$\begin{aligned} \eta(\mathbf{x}) &= \eta_{\emptyset} + \sum_{i=1}^p \eta_i(x_i) + \sum_{1 \leq i < j \leq p} \eta_{i,j}(x_i, x_j) + \dots + \eta_{1, \dots, p}(\mathbf{x}) \\ &= \sum_{u \subseteq \{1, \dots, p\}} \eta_u(\mathbf{x}_{\mathbf{u}}). \end{aligned} \tag{6.1}$$

The expansion (6.1) exists and is unique under the assumption

$$\int \eta_u(\mathbf{x}_u) dP_{X_i} = 0 \quad \forall i \in u, \forall u \subseteq \{1 \cdots p\}.$$

Equation (6.1) tells us that the model function $Y = \eta(\mathbf{X})$ can be expanded in a functional ANOVA. The independence of the inputs and the orthogonality properties ensure the global variance decomposition of the output as $V(Y) = \sum_{u \in S \setminus \emptyset} V(\eta_u(\mathbf{X}_u))$. Moreover, by integration,

$$\eta_\emptyset = \mathbb{E}(Y), \quad \eta_u = \mathbb{E}(Y|\mathbf{X}_u) - \sum_{v \subset u} \eta_v, \quad |u| \geq 1. \quad (6.2)$$

Hence, the contribution of a group of variables $\mathbf{X}_u$ in the model can be quantified in the fluctuations of Y . The Sobol indices are defined by:

$$S_u = \frac{V(\eta_u)}{V(Y)} = \frac{V[\mathbb{E}(Y|\mathbf{X}_u)] - \sum_{v \subset u} (-1)^{|u|-|v|} V[\mathbb{E}(Y|\mathbf{X}_v)]}{V(Y)}, \quad u \subseteq [1 : p]. \quad (6.3)$$

Furthermore, Sobol indices are summed to 1.

However, the main assumption is that the input parameters are independent. This is unrealistic in many cases. The use of the previous expressions is not excluded in case of inputs' dependence, but could lead to an unobvious or a wrong interpretation. Moreover, used methods to estimate them may mislead final results because most of these methods are built on the assumption of independence. For these reasons, the objective of the upcoming work is to show that the construction of sensitivity indices under dependence condition can be done into a mathematical frame.

In the next section, we propose a generalization of the Hoeffding decomposition under suitable conditions on the joint distribution function of the inputs. This decomposition consists of summands of increasing dimension, like in Hoeffding one. But this time, the components are hierarchically orthogonal instead of being mutually orthogonal. The hierarchical orthogonality will be mathematically defined further, and the obtained decomposition will be denoted by HOFD, as mentioned in the introduction. Thus, the global variance of the output could be decomposed as a sum of covariance terms depending on the summands of the HOFD. It leads to the construction of generalized sensitivity indices summed to 1 to perform well tailored SA in case of dependence.

6.2.3 Generalized decomposition for dependent inputs

We no more assume that $P_{\mathbf{X}}$ is a product measure. Nevertheless, we assume:

$$\begin{array}{c}
 P_{\mathbf{X}} << \nu \\
 \text{where} \\
 \nu(dx) = \nu_1(dx_1) \otimes \cdots \otimes \nu_p(dx_p)
 \end{array}
 \tag{C.1}$$

Our main assumption is:

$$\exists 0 < M \leq 1, \forall u \subseteq [1 : p], \quad p_{\mathbf{X}} \geq M \cdot p_{\mathbf{X}_u} p_{\mathbf{X}_{-u}} \quad \nu\text{-a.e.}
 \tag{C.2}$$

where $-u$ denotes the complement set of u in $[1 : p]$. $p_{\mathbf{X}_u}$ and $p_{\mathbf{X}_{-u}}$ are respectively the marginal densities of $\mathbf{X}_u$ and $\mathbf{X}_{-u}$.

A sufficient condition for (C.2) will be given later in Proposition 6.2. It will give a better understanding of (C.2). However, (C.2) may pave the way for another type of dependence. Indeed, for $p = 2$ and $M = 1$, it implies the positive quadrant dependence given by Lehmann [Leh66]. Furthermore, the reinterpretation with copulas, given in Proposition 6.3, shows that (C.2), despite its unobvious form, holds for a wide class of copulas.

The section is organized as follows: a preliminary lemma gives the main result to show that H^0 is a complete space. This is a reminder of Lemma 3.1 studied in [Sto94]. It ensures the existence and the uniqueness of the projection of η onto H^0 , as given in Theorem 3.1 of [Sto94]. The generalized decomposition of η is finally obtained by adding a residual orthogonal to every summand, as suggested in [Hoo07].

To begin with, let us state some definitions. In the usual ANOVA context, a model is said to be hierarchical if for every term involving some inputs, all lower-order terms involving a subset of these inputs also appear in the model. Correspondingly, a hierarchical collection T of subsets of $[1 : p]$ is defined as follows:

Definition 6.1. *A collection $T \subset S$ is hierarchical if for $u \in T$ and v a subset of u , one has $v \in T$.*

Using this definition, let state the following result:

Lemma 6.1. *Let $T \subset S$ be hierarchical. Suppose that (C.1) and (C.2) hold. Set $\delta = 1 - \sqrt{1 - M} \in]0, 1]$. Then, for any $h_u \in H_u^0$, $u \in T$, we have:*

$$\mathbb{E} \left[\left(\sum_{u \in T} h_u(\mathbf{X}) \right)^2 \right] \geq \delta^{\#(T)-1} \sum_{u \in T} \mathbb{E}[h_u^2(\mathbf{X})].
 \tag{6.4}$$

Lemma 6.1 is one of the key tools to show the hierarchical decomposition given in Theorem 6.1. To be self-contained, we give the proofs of Lemma 6.1 and Theorem 6.1 in the Appendix A.

Theorem 6.1. *Let η be any function in $L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. Then, under (C.1) and (C.2), there exist functions $\eta_0, \eta_1, \dots, \eta_{1,\dots,p} \in H_0 \times H_1^0 \times \dots \times H_{1,\dots,p}^0$ such that the following equality holds:*

$$\begin{aligned} \eta(\mathbf{X}) &= \eta_0 + \sum_i \eta_i(X_i) + \sum_{i,j} \eta_{ij}(X_i, X_j) + \dots + \eta_{1,\dots,p}(\mathbf{X}) \\ &= \sum_{u \in S} \eta_u(\mathbf{X}_u). \end{aligned} \quad (6.5)$$

Moreover, this decomposition is unique.

Notice that in the case where the input variables $X_1, \dots, X_p$ are independent, $\delta = 1$ and Inequality (6.4) of Lemma 6.1 becomes an equality. Indeed, in this case, this equality is directly obtained by orthogonality of the summands, and the HOFD turns out to be the classical Sobol decomposition.

6.2.4 Generalized sensitivity indices

As stated in Theorem 6.1, under (C.1) and (C.2), the output Y of the model can be uniquely decomposed as a sum of hierarchically orthogonal terms. Thus, the global variance has a simplified decomposition into a sum of covariance terms. From this fact, we can define generalized sensitivity indices.

Definition 6.2. *The sensitivity index S_u of order $|u|$ measuring the contribution of $\mathbf{X}_u$ into the model is given by:*

$$S_u = \frac{V(\eta_u(\mathbf{X}_u)) + \sum_{u \cap v \neq \emptyset, u \neq v} \text{Cov}(\eta_u(\mathbf{X}_u), \eta_v(\mathbf{X}_v))}{V(Y)} \quad (6.6)$$

More specifically, the first order sensitivity index S_i is given by:

$$S_i = \frac{V(\eta_i(X_i)) + \sum_{\substack{v \neq \emptyset \\ i \notin v}} \text{Cov}(\eta_i(X_i), \eta_v(\mathbf{X}_v))}{V(Y)} \quad (6.7)$$

An obvious consequence is given in Proposition 6.1 (see proof in the Appendix A):

Proposition 6.1. *Under (C.1) and (C.2), the sensitivity indices S_u defined previously sum to 1, i.e.*

$$\sum_{u \in S \setminus \{\emptyset\}} S_u = 1. \quad (6.8)$$

Discussion

First, we note that these indices are very similar to the ones proposed in Li *et al.* [LRY⁺10]. Nevertheless, their approach is different as they are mainly interested with the estimation of these quantities. For this purpose, they first approximate the model output by component functions. Further, they deduce a decomposition of the global variance. Here, indices originate from the HOFD of any regular function, but the assumptions (C.1) and (C.2) are required. Although the problem takes another route here, it should be noted that there exist several other methods to evaluate the importance of variables. Among them, the PLS [FF93], the Lasso regression [Tib96], or, more generally, the LARS method [EHJT04] are performed for shrinkage and variable selection. More closely related, the COSSO regression [LZ06] is a Lasso on SS-ANOVA [Gu02] components to select a sparse number of effects. In [LMM11], the prediction quality of linear models after parameters selection by sensitivity analysis is compared to Lasso regression. It highlights the complex relationship between sensitivity analysis and the prediction mean squared error.

In terms of interpretation, we notice that the covariance terms included in these indices allow to take into account the input dependence. Thus, they allow to measure the influence of a variable on the model, especially when a part of its variability is embedded into the one of other dependent terms. The form of the sensitivity indices allows for distinguishing the full contribution of a variable and its contribution into another correlated income. Also, if inputs are independent, the summands η_u are mutually orthogonal, so $\text{Cov}(\eta_u, \eta_v) = 0$, $u \neq v$, and we recover the well known Sobol indices. Hence, these new sensitivity indices can be seen as a generalization of Sobol indices.

6.3 Examples of distribution function

This section is devoted to examples of distribution function satisfying (C.1) and (C.2). The first hypothesis only implies that the reference measure is a product measure, whereas the second is trickier to obtain.

In the first part, we give a sufficient condition to get (C.2) for any number p of input variables. The second part deals with the case $p = 2$, for which we give equivalences of (C.2) in terms of copulas.

6.3.1 Boundedness of the inputs density function

The difficulty of Condition (C.2) is that the inequality has to be true for any splitting of the set $(X_1, \dots, X_p)$ into two disjoint blocks.

First, let give a distribution for which the condition (C.2) is not satisfied.

Let ν be the Lebesgue measure, and let $\mathbf{X} \sim N_p(0, \Sigma)$, with $\Sigma_{ii} = \sigma_i^2$, and $\Sigma_{ij} = \rho_{ij}$ if $j \neq i$. For a better understanding, we assume that, for a given

$\mathbf{u} \in S$, the vector $\mathbf{X}$ is reordered as $\mathbf{X} = (\mathbf{X}_{\mathbf{u}}, \mathbf{X}_{-\mathbf{u}})$. Thus, Σ may be written as

$$\Sigma = \begin{pmatrix} \Sigma_{\mathbf{u}} & \Omega_{\mathbf{u}, -\mathbf{u}} \\ \Omega_{\mathbf{u}, -\mathbf{u}} & \Sigma_{-\mathbf{u}} \end{pmatrix},$$

where $(\Omega_{\mathbf{u}, -\mathbf{u}})_{i,j} = \rho_{ij}$, for $i \in \mathbf{u}$, $j \in -\mathbf{u}$. In this example, it is straightforward to compute $p_{\mathbf{X}_{\mathbf{u}}}$ and $p_{\mathbf{X}_{-\mathbf{u}}}$, and,

$$\frac{p_{\mathbf{X}}(\mathbf{x})}{p_{\mathbf{X}_{\mathbf{u}}}(\mathbf{x}_{\mathbf{u}}) \cdot p_{\mathbf{X}_{-\mathbf{u}}}(\mathbf{x}_{-\mathbf{u}})} = \frac{|\Sigma|^{-1/2}}{(|\Sigma_{\mathbf{u}}| |\Sigma_{-\mathbf{u}}|)^{-1/2}} \cdot e^{-\frac{1}{2} [\mathbf{x} \Sigma^{-1} \mathbf{x} - \mathbf{x}_{\mathbf{u}} \Sigma_{\mathbf{u}}^{-1} \mathbf{x}_{\mathbf{u}} - \mathbf{x}_{-\mathbf{u}} \Sigma_{-\mathbf{u}}^{-1} \mathbf{x}_{-\mathbf{u}}]}.$$

Hence, as $\mathbf{x} \rightarrow +\infty$, $\frac{p_{\mathbf{X}}}{p_{\mathbf{X}_{\mathbf{u}}} \cdot p_{\mathbf{X}_{-\mathbf{u}}}} \rightarrow 0$, and (C.2) is not satisfied.

Now, we give a sufficient condition for (C.2) to hold in Proposition 6.2. This proposition is basically the condition given by Stone in the case of the Lebesgue measure on a compact rectangle. The proof can be found in [Sto94] page 132.

Proposition 6.2. *Assume that there exist $M_1, M_2 > 0$ with*

$$\boxed{M_1 \leq p_{\mathbf{X}} \leq M_2} \quad (\text{C.3})$$

Then, Condition (C.2) holds.

Let give now an example where (C.3) is satisfied.

Example 6.1. *Let ν be the multidimensional gaussian distribution $N_p(m, \Sigma)$ with*

$$m = \begin{pmatrix} m_1 \\ \vdots \\ m_p \end{pmatrix}, \quad \Sigma = \begin{pmatrix} \sigma_1^2 & \cdots & 0 \\ & \ddots & \\ 0 & \cdots & \sigma_p^2 \end{pmatrix}.$$

Assume that $P_{\mathbf{X}}$ is the Gaussian mixture $\alpha \cdot N_p(m, \Sigma) + (1 - \alpha) \cdot N_p(\mu, \Omega)$, $\alpha \in]0, 1[$ with

$$\mu = \begin{pmatrix} \mu_1 \\ \vdots \\ \mu_p \end{pmatrix}, \quad \Omega = \begin{pmatrix} \varphi_1^2 & \rho_{12} & \cdots & \rho_{1p} \\ & \cdots & & \\ \rho_{1p} & \cdots & \cdots & \varphi_p^2 \end{pmatrix}.$$

Then, (C.3) holds iff the matrix $(\Omega^{-1} - \Sigma^{-1})$ is positive definite.

In the next section, we will see that (C.2) has a copula version when $p = 2$. Thus, we establish two equivalent conditions to (C.2). We will give some examples of distribution satisfying one of these conditions.

6.3.2 Examples of distribution of two inputs

Here, we consider the simpler case $p = 2$. Also, until Section 6.4, we will assume that ν is absolutely continuous with respect to Lebesgue measure. The structure of dependence of X_1 and X_2 can be modeled by copulas. Copulas [Nel06] give a relationship between a joint distribution and its marginals. Sklar's theorem [Sk171] ensures that for any distribution function $F(x_1, x_2)$ with marginal distributions $F_1(x_1)$ and $F_2(x_2)$, F has the copula representation,

$$F(x_1, x_2) = C(F_1(x_1), F_2(x_2)),$$

where the measurable function C is unique whenever F_1 and F_2 are absolutely continuous.

The next corollary gives in the absolutely continuous case the relationship between a joint density and its marginal:

Corollary 6.1. *In terms of copulas, the joint density of $\mathbf{X}$ is given by:*

$$p_X(x_1, x_2) = c(F_1(x_1), F_2(x_2))p_{X_1}(x_1)p_{X_2}(x_2). \quad (6.9)$$

Furthermore,

$$c(u, v) = \frac{\partial^2 C}{\partial u \partial v}(u, v), \quad (u, v) \in [0, 1]^2. \quad (6.10)$$

Now, Condition (C.2) may be rephrased in terms of copulas:

Proposition 6.3. *For a two-dimensional model, the three following conditions are equivalent:*

$$1. \quad \boxed{\exists 0 < M < 1, p_X \geq M \cdot p_{X_1}p_{X_2} \quad \nu\text{-a.e.}} \quad (C.4)$$

$$2. \quad \boxed{\exists 0 < M < 1, c(u, v) \geq M, \quad \forall (u, v) \in [0, 1]^2} \quad (C.5)$$

$$3. \quad \boxed{\exists 0 < M < 1, a \text{ copula } \tilde{C}, C(u, v) = Muv + (1 - M)\tilde{C}(u, v)} \quad (C.6)$$

The proof of Proposition 6.3 is postponed to the Appendix A. Hence, the generalized Hoeffding decomposition holds for a wide class of examples. The Ali-Mikhail-Hal [AMH78] and the Frank copulas belong to this class.

Example 6.2.

- The Ali-Mikhail-Hal copula satisfies (C.5), with

$$c_\theta(u, v) = \frac{(1 - \theta)[1 - \theta(1 - u)(1 - v)] + 2\theta uv}{[1 - \theta(1 - u)(1 - v)]^3} \geq 1 - |\theta|, \quad \theta \in]-1, 1].$$

It should also be noted that the Morgenstern and the Lin's iterated Morgenstern copulas also satisfy (C.5).

- The Frank copula is a Archimedian copula, and satisfies (C.5). It is characterized by the generator:

$$\varphi(x) = \log \left(\frac{e^{-\theta x} - 1}{e^{-\theta} - 1} \right), \quad \theta \in \mathbb{R} \setminus \{0\},$$

and

$$C(u, v) = \varphi^{-1}[\varphi(u) + \varphi(v)], \quad u, v \in [0, 1]. \quad (6.11)$$

Here, $c(u, v) \geq -\theta(e^{-\theta} - 1)e^{-2\theta}$ if $\theta > 0$, $c(u, v) \geq -\theta(e^{-\theta} - 1)$ elsewhere.

Other examples of copulas from the Archimedian class also satisfy (C.4)-(C.6) by an intermediate proposition. Details are given in Appendix A. Leaving the class of copulas, we now directly work with the joint density function. Proposition 6.4 gives a general form of distribution for our framework:

Proposition 6.4. *If $p_{\mathbf{X}}$ has the form*

$$p_X(x_1, x_2) = \alpha \cdot f_{X_1}(x_1)f_{X_2}(x_2) + (1 - \alpha) \cdot g_X(x_1, x_2), \quad \alpha \in]0, 1[, \quad (6.12)$$

where f_{X_1}, f_{X_2} are univariate density functions, and g_X is any density function (with respect to ν) with marginals f_{X_1} and f_{X_2} , then $p_{\mathbf{X}}$ satisfies (C.5).

The proof is straightforward.

The remaining part of the paper is devoted to the estimation of the HOFD components. In the next section, we will assume that the inputs are independent pairs of dependent variables (abbreviated in IPDV).

The simplest case of a single pair of dependent variables is first discussed. Then, the more general IPDV case is studied. In all cases, first and second order indices are defined to measure the contribution of each pair of dependent variables and each of its components in the model. Indices of order greater than one involving variables from different pairs will not be studied here.

6.4 Estimation

Here, we investigate a different approach from [LR12]. The method relies on the property of hierarchical orthogonality ($H_u^0 \perp H_v^0, \forall v \subset u$), and on projection operator onto H_u^0 , denoted by $P_{H_u^0}$, for $u \in S$. The idea is to project

the output onto $H_u^0, \forall u$, to get the HOFD components $(\eta_u)_u$ as a solution of a functional linear system. It then consists in solving the system numerically. This section is first devoted to the HOFD terms computation with this method in two dimensional models. At last, we extend the procedure to the more general IPDV case.

6.4.1 Models of $p = 2$ input variables

This part is devoted to the simple case of bidimensional models $Y = \eta(X_1, X_2)$. Assuming that Conditions (C.1) and (C.2) both hold, we proceed as follows:

Procedure 1

1. HOFD of the output:

$$Y = \eta_\emptyset + \eta_1(X_1) + \eta_2(X_2) + \eta_{12}(X_1, X_2). \quad (6.13)$$

2. Projection of $Y = \eta(\mathbf{X})$ on $H_u^0, \forall u \subseteq \{1, 2\}$. As $H_u^0 \perp H_v^0, \forall v \subset u$, we obtain:

$$\begin{pmatrix} I & 0 & 0 & 0 \\ 0 & I & P_{H_1^0} & 0 \\ 0 & P_{H_2^0} & I & 0 \\ 0 & 0 & 0 & I \end{pmatrix} \begin{pmatrix} \eta_\emptyset \\ \eta_1 \\ \eta_2 \\ \eta_{12} \end{pmatrix} = \begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix}. \quad (6.14)$$

3. Computation of the right-hand side vector of (6.14):

$$\begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix} = \begin{pmatrix} \mathbb{E}(\eta) \\ \mathbb{E}(\eta|X_1) - \mathbb{E}(\eta) \\ \mathbb{E}(\eta|X_2) - \mathbb{E}(\eta) \\ \eta - \mathbb{E}(\eta|X_1) - \mathbb{E}(\eta|X_2) + \mathbb{E}(\eta) \end{pmatrix}. \quad (6.15)$$

In this frame, we have:

Proposition 6.5. Let η be any function of $L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. Then, under (C.1) and (C.2), the linear system

$$(\mathcal{S}) \begin{pmatrix} I & 0 & 0 & 0 \\ 0 & I & P_{H_1^0} & 0 \\ 0 & P_{H_2^0} & I & 0 \\ 0 & 0 & 0 & I \end{pmatrix} \begin{pmatrix} h_\emptyset \\ h_1 \\ h_2 \\ h_{12} \end{pmatrix} = \begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix}, \quad (6.16)$$

admits in $h = (h_\emptyset, \dots, h_{12}) \in H_\emptyset \times \dots \times H_{12}^0$ the unique solution $h^* = (\eta_\emptyset, \eta_1, \eta_2, \eta_{12})$.

4. *Reduction of the system (6.14).* As the constant term η_0 corresponds to the mean value of η , and the last term η_{12} can be deduced from the others, the dimension of the system (6.16) can even be reduced to:

$$\begin{pmatrix} I & P_{H_1^0} \\ P_{H_2^0} & I \end{pmatrix} \begin{pmatrix} \eta_1 \\ \eta_2 \end{pmatrix} = \begin{pmatrix} \mathbb{E}(\eta|X_1) - \mathbb{E}(\eta) \\ \mathbb{E}(\eta|X_2) - \mathbb{E}(\eta) \end{pmatrix}. \quad (6.17)$$

For the next step, we will denote by $A_2\Delta = B$ the system (6.17), where A_2 is the left-hand side matrix of projection operators of (6.17), $\Delta = {}^t(\eta_1 \ \eta_2)$, and B is the right-hand side vector of (6.17).

5. *Estimation procedure:* Suppose that we get a sample of n observations $(Y_l, \mathbf{X}_l)_{l=1, \dots, n}$ from the distribution of $(Y, \mathbf{X})$.
- The numerical resolution of (6.17) is achieved by an iterative Gauss Seidel algorithm [Kre98] which consists first in decomposing A_2 as a sum of lower triangular (L_2) and strictly upper triangular (U_2) matrices.

Further, the technique uses an iterative scheme to compute Δ . At step $k+1$, we have:

$$\Delta^{(k+1)} := \begin{pmatrix} \Delta_1^{(k+1)} \\ \Delta_2^{(k+1)} \end{pmatrix} = L_2^{-1}(B - U_2 \cdot \Delta^{(k)}) \quad (6.18)$$

Using expression of A_2 , we get:

$$\Delta^{(k+1)} = \begin{pmatrix} \mathbb{E}(Y - \Delta_2^{(k)}|X_1) - \mathbb{E}(Y - \Delta_2^{(k)}) \\ \mathbb{E}(Y - \Delta_1^{(k+1)}|X_2) - \mathbb{E}(Y - \Delta_1^{(k+1)}) \end{pmatrix}. \quad (6.19)$$

The iterative scheme (6.19) requires to estimate conditional expectations. As studied in Da Veiga et al [DVWG09], we propose to estimate these quantities by local polynomial regression at each point of observation $(Y_l, \mathbf{X}_l)$. Then, we use the leave-one-out technique to set the learning sample and the test sample. Moreover, as the local polynomial method can be summed up to a generalized least squares (see Fan and Gijbels [FG96]), the Sherman-Morrison formula [SM50] is applied to reduce the computational time.

We stop the procedure when $\|\Delta^{(k+1)} - \Delta^{(k)}\| \leq \varepsilon$, for a small positive ε .

Once (η_1, η_2) have been estimated, we deduce an estimation of η_{12} by subtraction.

- We use empirical variance and covariance estimation to estimate sensitivity indices S_1 , S_2 and S_{12} .

6. *Convergence of the algorithm:* now, we hope that the Gauss Seidel algorithm converges to the true solution. Looking back at (6.14), we see that we only have to consider $P_{H_1^0}$ (respectively $P_{H_2^0}$) restricted to H_2^0 (respectively to H_1^0).

Under this restriction, let us define the associated norm operator as:

$$\|P_{H_i^0}\|^2 := \sup_{\substack{\mathbb{E}(U^2)=1 \\ U \in H_j^0}} \mathbb{E}[P_{H_i^0}(U)^2], \quad i, j = 1, 2, \quad j \neq i.$$

As explained in [DP73], Gauss Seidel algorithm converges to the true solution Δ if A_2 is strictly diagonally dominant, which is implied by:

$$\|P_{H_i^0}\| < \|I\| = 1, \quad i = 1, 2. \quad (6.20)$$

As $P_{H_i^0}(U) = \mathbb{E}(U|X_i) - \mathbb{E}(U)$ by (6.15), the Jensen inequality [Jen06] is applied. Take $U \in H_1^0$:

$$\begin{aligned} \|P_{H_2^0}\| &= \sup_{\substack{\mathbb{E}(U^2)=1 \\ U \in H_1^0}} \mathbb{E}[(\mathbb{E}(U|X_2) - \mathbb{E}(U))^2] \\ &\leq \sup_{\substack{\mathbb{E}(U^2)=1 \\ U \in H_1^0}} \mathbb{E}[\mathbb{E}(U^2|X_2)] = 1 \quad \text{as } U \in H_1^0. \end{aligned}$$

The same holds for $\|P_{H_2^0}\|$. Thus $\|P_{H_2^0}\| < 1$ holds if U (function of X_j) is not X_i -measurable. Hence, the condition of convergence holds if X_1 is not a measurable function of X_2 .

6.4.2 Generalized IPDV models

Assume that the number of inputs is even, so $p = 2k$, $k \geq 2$. We note each group of dependent variables as $\mathbf{X}^{(i)} := (X_1^{(i)}, X_2^{(i)})$, $i = 1, \dots, k$. By rearrangement, we may assume that:

$$\mathbf{X} = (\underbrace{X_1, X_2}_{\mathbf{X}^{(1)}}, \dots, \underbrace{X_{2k-1}, X_{2k}}_{\mathbf{X}^{(k)}}).$$

If p is odd, one of the pairs is reduced to a single input. SA for IPDV models has already been treated in [JLD06]. Indeed, they proposed therein to estimate usual sensitivity indices on groups of variables via a Monte Carlo method. Thus, they have interpreted the influence of every group of variables on the global variance. Here, we will go further by trying to measure the influence of each variable on the output, but also the effects of the independent pairs.

To begin with, as a slight generalization of [Sob93] and used in [JLD06], let apply the Sobol decomposition on independent groups of dependent variables,

$$\eta(\mathbf{X}) = \eta_0 + \eta_1(\mathbf{X}^{(1)}) + \cdots + \eta_k(\mathbf{X}^{(k)}) + \sum_{|u|=2}^k \eta_u(\mathbf{X}^{(u)}),$$

where for $u = \{u_1, \dots, u_t\}$ and $t = |u|$, we set $\mathbf{X}^{(u)} = (\mathbf{X}^{(u_1)}, \dots, \mathbf{X}^{(u_t)})$. Furthermore, $\langle \eta_u, \eta_v \rangle = 0, \forall u \neq v$.

Under the assumptions discussed in the previous section, we can apply the HOFD on each component η_i , that is,

$$\eta_i(\mathbf{X}^{(i)}) = \eta_i(X_1^{(i)}, X_2^{(i)}) = \varphi_{i0} + \varphi_{i,1}(X_1^{(i)}) + \varphi_{i,2}(X_2^{(i)}) + \varphi_{i,12}(\mathbf{X}^{(i)}),$$

with $\langle \varphi_{i,u}, \varphi_{i,v} \rangle = 0, \forall v \subset u \subseteq \{1, 2\}$. In this way, let define some new generalized indices for IPDV models:

Definition 6.3. For $i = 1, \dots, k$, the sensitivity index measuring the respective contribution of $X_j^{(i)}$ ($j = 1, 2$) and $(X_1^{(i)}, X_2^{(i)})$ on the output is:

$$S_{i,j} = \frac{V(\varphi_{i,j}) + \text{Cov}(\varphi_{i,j}, \varphi_{i,k})}{V(Y)}, k = 2 \text{ if } j = 1 \quad S_{i,12} = \frac{V(\varphi_{i,12})}{V(Y)}. \quad (6.21)$$

The estimation procedure of these indices is quite similar to Procedure 1:

Procedure 2

1. Estimation of $(\eta_i)_{i=1, \dots, k}$: as reminded in Part 6.2.2 with Equations (6.2), $\eta_i = \mathbb{E}(Y|\mathbf{X}^{(i)}) - \mathbb{E}(Y)$. We use the non parametric estimation reminded in step 5 of Procedure 1 to get $\hat{\eta}_i$.
2. For $i = 1, \dots, k$, we apply step 2 to step 5 of Procedure 1, considering $\hat{\eta}_i$ as the new output.

If p is odd, the procedure is the same except that the influence of the independent variable is measured by a Sobol index, as it is independent from all the others. The next part is devoted to numerical examples.

6.5 Numerical examples

In this section, we study three examples with dependent input variables. For the first two illustrations, we consider IPDV models and a Gaussian mixture distribution on the input variables. The covariance matrices of the mixture satisfy conditions of Example 6.1.

We give estimations of our new indices, and compare these estimations to the true values, computed from expressions (6.6). We also compute dispersions of the estimators.

In [DVWG09], Da Veiga *et al.* proposed to estimate the classical Sobol indices as defined in (6.3) by nonparametric tools. Indeed, the local polynomial regression were used to estimate conditional moments $\mathbb{E}(Y|\mathbf{X}_{\mathbf{u}})$, $u \subseteq [1 : p]$. This method, used further, will be called Da Veiga procedure (DVP). Results given by DVP are compared with our method. We show that the usual sensitivity indices are not appropriate in the dependence frame, even if a relevant estimation method is used.

The last example is devoted to a more realistic case. The study case concerns the river flood inundation of an industrial site. In this example, input variables have different distributions, and some pairs are linearly correlated. We represent dispersions of the estimation of our new indices, and give some physical interpretations.

6.5.1 Two-dimensional IPDV model

Let consider the model

$$Y = \exp X_1 + X_1 + X_2.$$

Here, ν and $P_{\mathbf{X}}$ are of the form given by Example 6.1, with $m = \mu = 0$. Thus, the analytical decomposition of Y is

$$\eta_0 = \mathbb{E}(\exp X_1), \quad \eta_1 = \exp X_1 + X_1 - \mathbb{E}(\exp X_1), \quad \eta_2 = X_2.$$

For the application, we implement Procedure 1 in Matlab software. We proceed to $L = 50$ simulations and $n = 1000$ observations. Parameters were fixed at $\sigma_1 = \sigma_2 = 1$, $\varphi_1^2 = \varphi_2^2 = 0.5$, $\rho_{12} = 0.4$ and $\alpha = 0.2$.

In Table 6.1, we give the estimation of our indices and their standard deviation (indicated by $\pm$) on L simulations. In comparison, we give the analytical value of each index. We also give estimators of the classical Sobol indices with DVP.

		S_1	S_2	S_{12}	$\sum_u S_u$
New indices	Estimation	0.745 ± 0.016	0.219 ± 0.015	0.045 ± 0.012	1
	Analytical	0.7497	0.2503	0	1
DVP indices	Estimation	0.7774	0.1792	0.043	1.00

Table 6.1: Estimation of the new and DVP indices with $\rho_{12} = 0.4$

Notice that we obtain a quite good estimation of S_1 with our estimation procedure. $\hat{S}_2$ is lightly lower than expected, and, consequently, the estimation

error of the interaction term η_{12} is bigger than 0. In comparison, the DVP badly scales S_2 , even if for both methods, the inputs hierarchy is the same. In our method, it would be relevant to separate the variance part to the covariance one in the first order indices. Indeed, in this way, we would be able to get the part of variability explained by X_i alone in S_i , and its contribution hidden in the dependence with X_j . We note S_i^v the variance contribution alone, and S_i^c the covariance contribution, that is

$$S_i = \underbrace{\frac{V(\eta_i)}{V(Y)}}_{S_i^v} + \underbrace{\frac{\text{Cov}(\eta_i, \eta_j)}{V(Y)}}_{S_i^c}, \quad i = 1, 2, j \neq i.$$

The new indices estimations given in Table 6.1 are decomposed in Table 6.2.

		S_i^v	S_i^c	S_i
Estimated	X_1	0.6531	0.0918	0.7449
	X_2	0.1049	0.1138	0.2187
Analytical	X_1	0.6122	0.1375	0.7497
	X_2	0.1129	0.1375	0.2504

Table 6.2: Estimation of S_i^v and S_i^c with $\rho_{12} = 0.4$

For each index, the covariate X_1 explains 65% (in estimation, 61% in reality) of the part of the total variability. However, the contribution embedded in the correlation is not negligible as it represents between 9% and 11% (13.75% analytically) of the total variance. Considering the shape of the model, it is quite natural to get a higher contribution of X_1 . Also, as their dependence is quite important, with a covariance term equal to 0.4, we are not surprised by the relatively high value of S_1^c (resp. S_2^c).

6.5.2 Linear four-dimensional model

The test model is

$$Y = 5X_1 + 4X_2 + 3X_3 + 2X_4.$$

Let consider the case of two blocks $\mathbf{X}^{(1)} = (X_1, X_3)$ and $\mathbf{X}^{(2)} = (X_2, X_4)$ of correlated variables. $\mathbf{X}^{(i)}$, $i = 1, 2$ follows the Gaussian mixture. The analytical sensitivity indices are given by (6.21). For $L = 50$ simulations and $n = 1000$ observations, we took $\varphi_1^{2(1)} = \varphi_2^{2(1)} = 0.5$, $\varphi_1^{2(2)} = 0.7$, $\varphi_2^{2(2)} = 0.3$, $\rho_{13}^{(1)} = 0.4$, $\rho_{24}^{(2)} = 0.37$, $\alpha_1 = \alpha_2 = 0.2$ and $\Sigma^{(1)} = \Sigma^{(2)} = II_2$.

Figure 6.1 displays the dispersion of indices of first order for all variables and second order for grouped variables. The true values and the estimators of classical Sobol indices with DVP are also represented.

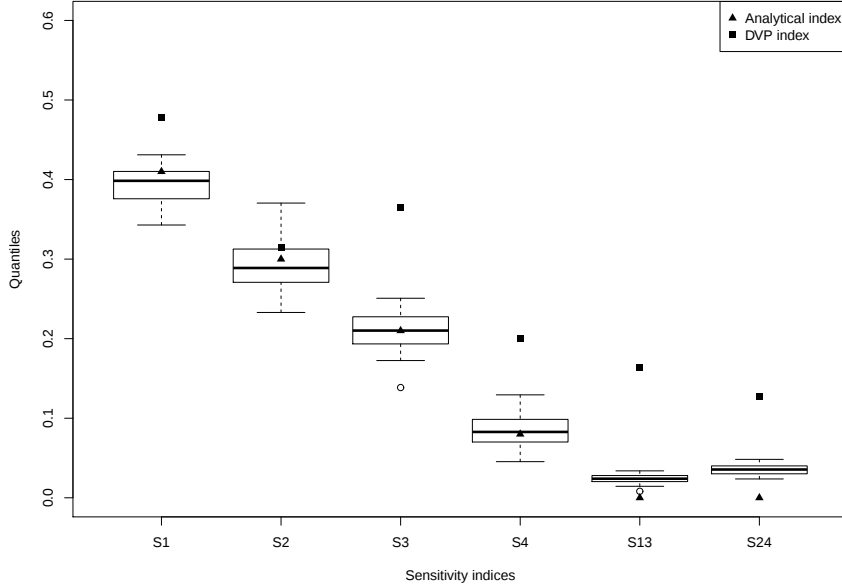


Figure 6.1: Boxplots representation of new indices-Comparison with analytical and DVP indices

On Figure 6.1, we see that X_1 has the biggest contribution, whereas the influence of X_4 is very low. It reflects well the model if we look at the coefficients of X_i , $i = 1, \dots, 4$. Notice that the interaction terms are well estimated, as they are close to 0. For each case, the dispersion on 50 simulations is very low. Furthermore, the DVP estimators are once again very high compared with the true indices values.

6.5.3 River flood inundation

Several SA methods have been studied on the simplified model of river flood inundation. A brief description of the model is given here, but more details can be found in [Ioo11, DR06]. The study case concerns an industrial site located near a river, and protected from it by a dyke. The goal is to study the water level with respect to the dyke height to prevent from inundation. The model has the following form

$$S = \underbrace{Z_v + h}_{Z_c} - H_d - C_b, \quad h = \left(\frac{Q}{BK_s \sqrt{\frac{Z_m - Z_v}{L}}} \right)^{0.6}.$$

This model is a crude simplification of the 1-D Saint Venant equations, when uniform and constant flow rate is assumed. The model output S is the maximal overflow that depends on eight random variables. H_d is a design

Variables	Meaning	Distribution
h	maximal annual water level	-
Q	maximal annual flow rate	Gumbel $G(1013, 558)$ tr. to $[500; 3000]$
K_s	Strickler friction coefficient	Normal $N(30, 8)$ tr. to $[15, +\infty[$
Z_v	river downstream level	Triangular $T(49, 50, 51)$
Z_m	river upstream level	Triangular $T(54, 55, 56)$
H_d	dyke height	Uniform $\mathcal{U}([7, 9])$
C_b	bank level	Triangular $T(55, 55.5, 56)$
L	length of the river stretch	Triangular $T(4990, 5000, 5010)$
B	river width	Triangular $T(295, 300, 305)$

Table 6.3: Description of inputs-output of the river flood model (tr. to=truncated to)

parameter. The randomness of other inputs is due to their spatio-temporal variability, or some inaccuracies of their estimation. Table 6.3 gives the meaning of each input variable, and how they are distributed. The river flow is represented in Figure 6.2.

Here, we suppose that (Q, K_s) is a correlated pair, with correlation coefficient $\rho = 0.5$. This correlation is admitted in real case, as we consider that the friction coefficient K_s increases with the flow rate Q . Also, (Z_v, Z_m) and (L, B) are assumed to be dependent with the same Pearson coefficient $\rho = 0.3$, because data are supposed to be simultaneously collected by the same measuring device. Correlated variables are simulated according to the algorithm given in [Sch09]. A theoretical background can be found in [Nel06]. We made $n = 1000$ model evaluations repeated $L = 50$ times. The dispersion of estimated indices is represented in Figure 6.3.

The most influential parameters are the flow rate Q , the downstream level Z_v , and the dyke height H_d . Q and H_d 's strong contribution makes sense here, as they represent the most important parameters to limit river flood.

6.6 Conclusions and perspectives

This paper gives a rigorous frame for general Hoeffding-Sobol decomposition in the case of dependent input variables. In the statistical procedure, we only consider the restricted case of IPDV models. Thus, for more general models, the mathematical properties of the statistical estimation of this decomposition remains a challenging open problem.

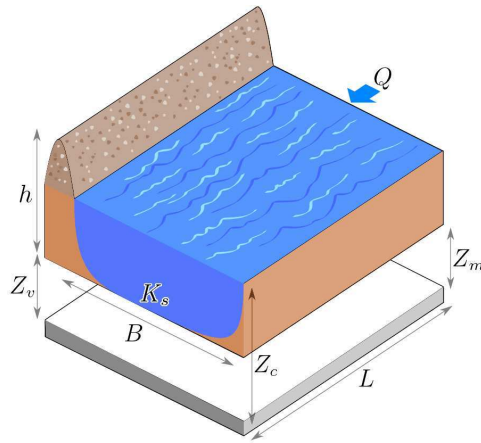


Figure 6.2: The river flood model

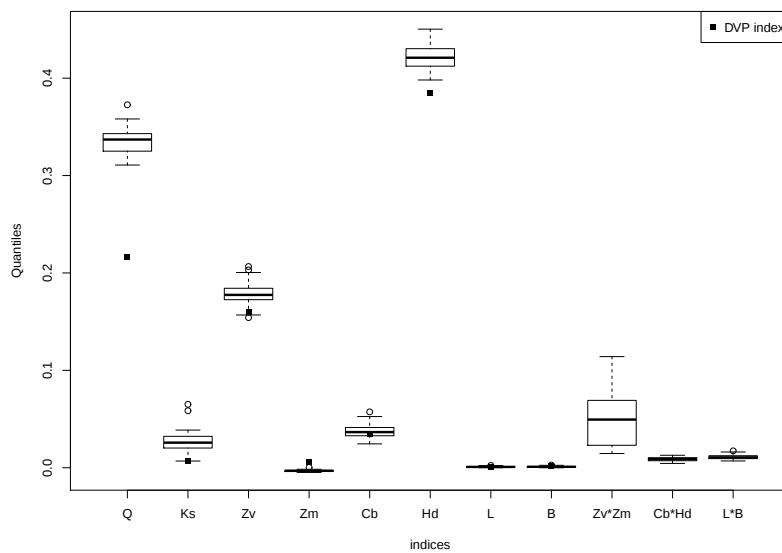


Figure 6.3: Boxplots of new indices for the river flood model

Introduction to Chapter 7

The next contribution is the submitted paper [CGP13]. The goal of this work is to complete the contribution done in the previous chapter by providing an efficient numerical method to estimate the sensitivity indices we previously defined.

For models with independent input variables, we observed that the components from the Hoeffding decomposition could be rewritten in terms of conditional expectations. Let remind the expression of the Hoeffding components,

$$\begin{aligned} \eta_{\emptyset} &= \mathbb{E}(Y) \\ \eta_i(X_i) &= \mathbb{E}(Y|X_i) - \mathbb{E}(Y), \\ \vdots &\quad \quad \quad \vdots \\ \eta_u(\mathbf{X}_u) &= \mathbb{E}(Y|\mathbf{X}_u) + \sum_{v \subset u} (-1)^{|u|-|v|} \mathbb{E}(Y|\mathbf{X}_v), \quad \forall u \in S \setminus \{\emptyset\}. \end{aligned}$$

Thus, as studied in Chapter 1 of Part I, there exist several efficient numerical techniques to provide a robust estimation of these functions. Most of these methods are based on a good approximation of the numerical integrations, and an appropriate choice of design of experiments.

For models with non independent inputs, we proposed in Chapter 6 to study a generalized functional decomposition of $Y = \eta(\mathbf{X})$, given by,

$$\eta(\mathbf{X}) = \sum_{u \in S} \eta_u(\mathbf{X}_u), \quad \mathbb{E}[\eta_u(\mathbf{X}_u)\eta_v(\mathbf{X}_v)] = 0, \quad \forall v \subset u, \quad \forall u \in S, \quad (6.22)$$

where the inputs $X_1, \dots, X_p$ admit a joint density $p_{\mathbf{X}}$ with respect to any product measure ν . The further work focuses on the estimation of the components $(\eta_u)_u$.

When ν is the Lebesgue measure, Hooker [Hoo07], followed by Li *et al.* [LR12], studies this decomposition under numerical aspects. On the one hand,

Hooker, shows that the orthogonality conditions imposed to the components of (6.22) admit an equivalent expression in terms of integrands. This equivalent expression can be extended to any product measure ν , as shown in Proposition 6.6.

Proposition 6.6. *Let u be a non empty set of S , and let $h_u \in L_{\mathbb{R}}^2$. The next two assertions are equivalent:*

1. *For all $v \subset u$, for all $h_v \in L_{\mathbb{R}}^2$,*

$$\mathbb{E}[h_u(\mathbf{X}_u)h_v(\mathbf{X}_v)] = \int_{\mathbb{R}^p} h_u(\mathbf{x}_u)h_v(\mathbf{x}_v)p_{\mathbf{X}}(\mathbf{x})d\nu(\mathbf{x}) = 0;$$

2. *For all $i \in u$,*

$$\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) = 0.$$

Proof of Proposition 6.6 .

Let $u \neq \emptyset$, and $h_u \in L_{\mathbb{R}}^2$ such that, for all $v \subset u$, $h_v \in L_{\mathbb{R}}^2$, we get

$$\int_{\mathbb{R}^p} h_u(\mathbf{x}_u)h_v(\mathbf{x}_v)p_{\mathbf{X}}(\mathbf{x})d\nu(\mathbf{x}) = 0.$$

Suppose now there exists $i \in u$ such that,

$$\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) \neq 0.$$

Set

$$h_v = \int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i), \quad v = u \setminus \{i\} \subset u.$$

Then, we get

$$\begin{aligned} \mathbb{E}[h_u(\mathbf{X}_u)h_v(\mathbf{X}_v)] &= \int_{\mathbb{R}^p} h_u(\mathbf{x}_u)h_v(\mathbf{x}_v)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_u(\mathbf{x}_u) \\ &= \int_{\mathbb{R}^p} h_u(\mathbf{x}_u) \left(\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) \right) p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_u(\mathbf{x}_u) \\ &= \int_{\mathbb{R}^{|u|-1}} \left(\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) \right)^2 d\nu_{u \setminus \{i\}}(\mathbf{x}_{u \setminus \{i\}}) \\ &> 0. \end{aligned}$$

There is contradiction, therefore $\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) = 0$.

Conversely, we assume that, for all $i \in u$,

$$\int_{\mathbb{R}} h_u(\mathbf{x}_u)p_{\mathbf{X}_u}(\mathbf{x}_u)d\nu_i(x_i) = 0.$$

Let $v \subset u$, and $i \in u \setminus v$. Then,

$$\begin{aligned}
\int_{\mathbb{R}^p} h_u(\mathbf{x}_u) h_v(\mathbf{x}_v) p_{\mathbf{X}} d\nu(\mathbf{x}) &= \int_{\mathbb{R}^{|u|}} h_u(\mathbf{x}_u) h_v(\mathbf{x}_v) p_{\mathbf{X}_u}(\mathbf{x}_u) d\nu_u(\mathbf{x}_u) \\
&= \int h_v(\mathbf{x}_v) \left(\int_{\mathbb{R}} h_u(\mathbf{x}_u) p_{\mathbf{X}_u}(\mathbf{x}_u) d\nu_i(x_i) \right) d\nu_{u \setminus \{i\}}(\mathbf{x}_{u \setminus \{i\}}) \\
&= 0
\end{aligned}$$

■

Later, Li *et al.* [LR12] deduce from Proposition 6.6 an analytic expression of each summand of the expansion (6.22). This formulation is given in Proposition 6.7.

Proposition 6.7. *The components of (6.22) can be expressed as follows,*

$$\begin{aligned}
\eta_{\emptyset} &= \int_{\mathbb{R}^p} \eta(\mathbf{x}) p_{\mathbf{X}}(\mathbf{x}) d\nu(\mathbf{x}) \\
\eta_u(\mathbf{x}_u) &= \int_{\mathbb{R}^{p-|u|}} \eta(\mathbf{x}) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u) - \sum_{v \subset u} \eta_v(\mathbf{x}_v) \\
&\quad - \sum_{u \cap v \neq \emptyset, v} \int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u), \quad u \in S \setminus \{\emptyset\}.
\end{aligned} \tag{6.23}$$

Proof of Proposition 6.7.

To show (6.23), we use the equivalence established in Proposition 6.6. First we have,

$$\begin{aligned}
\int_{\mathbb{R}^p} \eta(\mathbf{x}) p_{\mathbf{X}}(\mathbf{x}) d\nu(\mathbf{x}) &= \sum_{u \in S} \int_{\mathbb{R}^p} \eta(\mathbf{x}_u) p_{\mathbf{X}}(\mathbf{x}) d\nu(\mathbf{x}) \\
&= \eta_{\emptyset} \text{ because } \mathbb{E}[\eta_u(\mathbf{X}_u)] = 0 \ \forall \ u \in S \setminus \{\emptyset\}.
\end{aligned}$$

Let $u \neq \emptyset$. Then,

$$\begin{aligned}
\int_{\mathbb{R}^{p-|u|}} \eta(\mathbf{x}) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u) &= \sum_{v \in S} \int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u) \\
&= \sum_{u \cap v = \emptyset} \int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u) \\
&\quad + \sum_{u \cap v \neq \emptyset} \int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-u}(\mathbf{x}-u) d\nu_{-u}(\mathbf{x}-u).
\end{aligned}$$

If

- $u \cap v = \emptyset$, then
 - either $v = \emptyset$, so,

$$\int_{\mathbb{R}^{p-|u|}} \eta_{\emptyset} p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) = \eta_{\emptyset};$$

- or $v \neq \emptyset$ and $v \subset -u$. Therefore,

$$\begin{aligned} \int \eta_v(\mathbf{x}_v) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) &= \int \left(\int_{\mathbb{R}^{|v|}} \eta_v(\mathbf{x}_v) dP_{\mathbf{X}_v}(\mathbf{x}_v) \right) p_{\mathbf{x}-\mathbf{u} \setminus \mathbf{v}} d\nu_{-u \setminus \mathbf{v}}(\mathbf{x}-\mathbf{u} \setminus \mathbf{v}) \\ &= 0. \end{aligned}$$

- $u \cap v \neq \emptyset$. Set $w = u \cap v$. So,
 - either $w = v$, i.e. $v \subseteq u$, and

$$\int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) = \eta_v;$$

- or $w \neq v$, and

$$\int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) \text{ not necessarily null.}$$

As a consequence,

$$\begin{aligned} \int_{\mathbb{R}^{p-|u|}} \eta(\mathbf{x}) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) &= \sum_{v \subseteq u} \eta_v(\mathbf{x}_v) \\ &+ \sum_{u \cap v \neq \emptyset, v} \int \eta_v(\mathbf{x}_v) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}). \end{aligned}$$

Therefore, for all $u \in S \setminus \{\emptyset\}$,

$$\begin{aligned} \eta_u(\mathbf{x}_u) &= \int_{\mathbb{R}^{p-|u|}} \eta(\mathbf{x}) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}) - \sum_{v \subset u} \eta_v(\mathbf{x}_v) \\ &- \sum_{u \cap v \neq \emptyset, v} \int_{\mathbb{R}^{p-|u|}} \eta_v(\mathbf{x}_v) p_{\mathbf{x}-\mathbf{u}}(\mathbf{x}-\mathbf{u}) d\nu_{-u}(\mathbf{x}-\mathbf{u}). \end{aligned}$$

■

Proposition 6.7 shows that, unlike the case of independent distribution, the analytical expression of the components are impracticable under this form, because each component depends on the other terms that have not been identified yet. For this reason, the estimation of $(\eta_u)_{u \in S}$ by numerical integrations is intractable here. Nevertheless, we can see these summands as the projection of the model function η in a constrained space.

By using the equivalent expression of the hierarchical orthogonality given in Proposition 6.6, Hooker focuses on the use of a constrained minimization

problem to estimate each component of (6.22). Using a uniform grid as sample, Hooker concentrates the points on regions of high weight and adjusts the density function accordingly [Hoo07]. He also suggests to explore other type of design of experiments to properly estimate the density. Thus, the major challenge of this method is to be able to choose an appropriate design of experiments to estimate the inputs density. Also, the numerical resolution of a constrained minimization could be a difficult task, especially for high-dimensional models.

More recently, Li *et al.* [LR12] use a different approach. They express each component as a linear combination of truncated basis functions of $L^2_{\mathbb{R}}$. More precisely, interaction functions of order higher than one are substituted to a linear combination that consists of tensored basis functions, and combinations of lower-order interaction terms. The authors identify the induced coefficients by a constrained minimization, solved by a numerical technique, called D-MORPH. The D-MORPH technique is a sophisticated descent technique that consists in finding an optimal solution in a consistent problem that admits an infinity of solutions [LR10a, LR10b].

In Chapter 7, we propose another numerical procedure to estimate the components of the decomposition (6.22). Instead of using minimization problem-solving techniques, we build a functional system substituting properly each component. From truncated basis functions of $L^2_{\mathbb{R}}$, these systems aim at satisfying the hierarchical orthogonality constraints that summands must check. Thus, each component is represented by a linear combination of functions of the constructed system, where coefficients are deduced by least-squared estimation.

The objective is to build approximation spaces such that each of its functions satisfies the orthogonality conditions given in (6.22). The other objective is to be able to control the numerical cost of such estimation. Indeed, the expansion given in (6.22) consists of 2^p functions. If, for each summand η_u , the representation system admits L_u functions (as it will be discussed, the size L_u grows exponentially with the order $|u|$), the procedure requires to estimate $\sum_{u \in S} L_u$ coefficients, quantity that blows up with the model dimension p .

To remedy to this problem, variable selection strategy has already proved its efficiency in functional ANOVA through polynomial chaos [Bla09], and wavelets [Tou11]. Here, we also improve our original method by a variable selection technique, presented in Part II of the manuscript. The goal is to be able to reconstitute a maximal amount of model information in a minimal cost.

Chapter 7

Generalized Sobol sensitivity indices for dependent variables: numerical methods

Abstract: The hierarchically orthogonal functional decomposition of any measurable function η of a random vector $\mathbf{X} = (X_1, \dots, X_p)$ consists in decomposing $\eta(\mathbf{X})$ into a sum of increasing dimension functions depending only on a subvector of $\mathbf{X}$. Even when $X_1, \dots, X_p$ are assumed to be dependent, this decomposition is unique if components are hierarchically orthogonal. That is, two of the components are orthogonal whenever all the variables involved in one of the summands are a subset of the variables involved in the other. Setting $Y = \eta(\mathbf{X})$, this decomposition leads to the definition of generalized sensitivity indices able to quantify the uncertainty of Y with respect to the dependent inputs $\mathbf{X}$ [CGP12]. In this paper, a numerical method is developed to identify the component functions of the decomposition using the hierarchical orthogonality property. Further, the asymptotic properties of the components estimation is studied, as well as the numerical estimation of generalized sensitivity indices through a toy model. In addition, a model coming from real world illustrates the interest of the method.

Keywords: Dependent variables; extended basis; Greedy algorithm; LARS; sensitivity analysis; Sobol ANOVA decomposition.

7.1 Introduction

In a nonlinear regression model, input parameters can be subject to many sources of uncertainty. The objective of global sensitivity analysis is to identify and to rank the input variables that drive the uncertainty of the model output. The most popular methods are the variance-based ones [SCS00]. Among them, the Sobol indices are widely used [Sob93]. This last method

stands on the assumption that the incomes are independent. Under this assumption, Hoeffding [Hoe48] shows that the model output can be uniquely decomposed into a sum of increasing dimension functions, where the integrals of every summand over any of its own variables must be zero. A consequence of these conditions is that all summands of the decomposition are mutually orthogonal. Using this decomposition, Sobol shows that the global variance can also be decomposed as a sum of partial variances. Thus, the so-called Sobol sensitivity index for a group of inputs is the ratio between the partial variance associated to these inputs and the global variance [Sob93]. However, in models with dependent inputs, the use of Sobol indices may lead to a wrong interpretation because the sensitivity induced by the dependence between two factors is implicitly included in their Sobol indices. To handle this problem, a naive solution can consist in computing Sobol sensitivity indices for independent groups of dependent variables. First introduced by Sobol [Sob93], this idea is exploited in practice by Jacques *et al.* [JLD06]. Nevertheless, this technique implies to work with models having several independent groups of inputs. Furthermore, it does not allow to quantify the individual contribution of each input. A different way to deal with this issue has been initiated by Borgonovo *et al.* [Bor07, BCT11]. These authors define a new measure based on the joint distribution of $(Y, \mathbf{X})$. The new sensitivity indicator of an input X_i measures the shift between the output distribution and the same distribution conditionally to X_i . This moment free index has many properties and has been applied to some real applications [BT08, BCT12]. However, the dependence issue remains unsolved as we do not know how the conditional distribution is distorted by the dependence, and so how it impacts the sensitivity index. Another idea is to use the Gram-Schmidt orthogonalization procedure. In an early work, Bedford [Bed98] suggests to orthogonalize the conditional expectations and then to use the usual variance decomposition on this new orthogonal collection. He then uses the Monte Carlo simulation to compute the indices. In the same spirit, Mara *et al.* [MT12] use the same Gram-Schmidt tool to decorrelate the inputs, but perform polynomial regressions to approximate the model. In both papers, the decorrelation method depends on the ordering of the variables, making the procedure computationally expensive and difficult to interpret.

Following the construction of Sobol indices previously exposed, Xu *et al.* [XG08] propose to decompose the partial variance of an input into a correlated and an uncorrelated contribution in the context of linear models. This last work has been later extended by Li *et al.* with the concept of HDMR [LRY⁺10, LRR01]. In [LRY⁺10], the authors suggest to reconstruct the model function via classical basis (polynomials, splines,...), then to deduce the decomposition of the response variance as a sum of partial variances and covariances. Instead of classical basis, Caniou *et al.* [CS10] use a polynomial chaos expansion to approximate the initial model as far as the copula

theory to model the dependence structure [Nel06]. Thus, in all these papers, the authors choose a type of model reconstruction before proceeding to the splitting of the response variance.

In a previous paper [CGP12], we revisit the Hoeffding decomposition in a different way, bringing a new definition in the case of dependent inputs. Inspired by the pioneering work of Stone [Sto94] and Hooker [Hoo07], we show, under a weak assumption on the inputs distribution, that any model function can be decomposed into a sum of hierarchically orthogonal component functions. This means that two of these summands are orthogonal whenever all variables included in one of the components are also involved in the other. The decomposition leads to generalized Sobol sensitivity indices able to quantify the uncertainty brought by dependent inputs on the model.

The goal of this paper is to complete the work done in [CGP12] by providing an efficient numerical method for the estimation of the generalized Sobol sensitivity indices. In our previous paper [CGP12], we have proposed a statistical procedure based on projection operators to identify the components of the hierarchically orthogonal functional decomposition (HOFD). The method consists in projecting the model output onto constrained spaces to obtain a functional linear system. The numerical resolution of these systems relies on an iterative scheme that requires to estimate conditional expectations at each step. On one hand, this method is well tailored for independent pairs of dependent variables models. On the other hand, it is difficult to apply to more general models because of its computational cost. Hooker [Hoo07] has also worked on the estimation of the HOFD components. This author studies the component estimation via a minimization problem under constraints using a sample grid. In general, this procedure is also quite computational demanding. Moreover, it requires to get a prior on the inputs distribution at each evaluation point, or, at least, to be able to estimate them properly. In a recent article, Li *et al.* [LR12] come back on Hooker's work and also identify the HOFD components by a least-squares method. They propose to approximate these components using their expansions on suitable basis. They bypass some technical problem of degenerate design matrix by using a continuous descent technique [LR10a].

In this paper, we propose an alternative to directly construct a hierarchical orthogonal basis. Inspired by the usual Gram-Schmidt algorithm, the procedure consists in recursively constructing for each component a multi-dimensional basis that satisfies the hierarchical orthogonal conditions. This procedure will be referred as the Hierarchical Orthogonal Gram-Schmidt (HOGS) procedure. Then, each component of the decomposition can be properly estimated by a linear combination of this basis. The coefficients are then estimated by the usual least-squares method. Thanks to the HOGS procedure, we show that the design matrix has full rank, so the minimization problem admits a unique and explicit solution. Further, we study the asymptotic properties of the estimated components. Nevertheless, the prac-

tical estimation of the one-by-one component suffers from the curse of dimensionality when using the ordinary least-squares estimation. To handle this problem, we propose to estimate parameters of the model using variable selection methods. Two usual algorithms are briefly presented, and are adapted to our method. Further, the HOGS procedure coupled with these algorithms is experimented on numerical examples.

The paper is organized as follows. In Section 7.2, we give and discuss the general results on the HOFD. We remind Conditions (C.7) and (C.8) under which the HOFD is available. Further, we give the generalized Sobol sensitivity indices, and discuss their interpretation. Section 7.3 is devoted to the HOGS procedure. We introduce the appropriate notation, and give a detailed procedure. In Section 7.4, we adapt the least-squares estimation to our problem, and we point out the curse of dimensionality. Further, we discuss about variables selection via a penalized minimization. Section 7.5 brings asymptotic results on the component estimators, constructed according to the HOGS procedure of Section 7.3. In Section 7.6, we present numerical applications. The first example is a toy function, and its objective is to show the efficiency of the HOGS procedure coupled with variable selection methods in the sensitivity estimation. The last example concerns the pressure applied to a tank. In this example, we want to detect the most influent inputs in the model with our procedure.

7.2 Generalized Sobol sensitivity indices

Functional ANOVA models are specified by a sum of functions depending on an increasing number of variables. A functional ANOVA model is said to be additive if only main effects are included in the model. It is said to be saturated if all interaction terms are included in the model. However, the existence and the uniqueness of such decomposition is ensured by some identifiability constraints. When the inputs are independent, any regular model function is exactly a saturated ANOVA model with pairwise orthogonal components, as reminded in the introduction. It results that the contribution of any group of variables onto the model is measured by the Sobol index, bounded between 0 and 1. Moreover, the Sobol indices are summed to 1 [Sob93]. The use of such an index is not excluded in the dependence context, but the information due to the dependence is considered several times. This could lead to a wrong interpretation of the Sobol indices. In this section, we remind the main results established in Chastaing *et al.* [CGP12] when inputs can be non-independent. In this case, the saturated ANOVA model is established with weaker identifiability constraints than for the independent case. This leads to a generalization of the Sobol indices well suited to perform global sensitivity analysis when the inputs are not independent. First, we remind the general context and notation. The last part is dedicated

to the generalization of the Hoeffding-Sobol decomposition when inputs are potentially dependent. The definition of the generalized sensitivity indices follows.

7.2.1 First settings

We denote by $\subset$ the strict inclusion, that is $A \subset B \Rightarrow A \cap B \neq B$, whereas we use $\subseteq$ when equality is possible.

Consider a measurable function η of a random vector $\mathbf{X} = (X_1, \dots, X_p) \in \mathbb{R}^p$, $p \geq 1$, and let Y be the real-valued response variable defined as

$$\begin{aligned} Y : (\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}) &\rightarrow (\mathbb{R}, \mathcal{B}(\mathbb{R})) \\ \mathbf{X} &\mapsto \eta(\mathbf{X}) \end{aligned}$$

where the joint distribution of $\mathbf{X}$ is denoted by $P_{\mathbf{X}}$. For a σ -finite measure ν on $(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p))$, we assume that $P_{\mathbf{X}} \ll \nu$ and that $\mathbf{X}$ admits a density $p_{\mathbf{X}}$ with respect to ν , that is $p_{\mathbf{X}} = \frac{dP_{\mathbf{X}}}{d\nu}$.

Also, we assume that $\eta \in L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. As usual, we define the inner product $\langle \cdot, \cdot \rangle$ and the norm $\|\cdot\|$ of the Hilbert space $L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$ as

$$\begin{aligned} \langle h_1, h_2 \rangle &= \int h_1(\mathbf{x})h_2(\mathbf{x})p_{\mathbf{X}}d\nu(\mathbf{x}) = \mathbb{E}(h_1(\mathbf{X})h_2(\mathbf{X})), \quad h_1, h_2 \in L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}), \\ \|h\|^2 &= \langle h, h \rangle = \mathbb{E}(h(\mathbf{X})^2), \quad h \in L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}}). \end{aligned}$$

Here $\mathbb{E}(\cdot)$ denotes the expectation. Further, $V(\cdot) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))^2]$ denotes the variance, and $\text{Cov}(\cdot, *) = \mathbb{E}[(\cdot - \mathbb{E}(\cdot))(* - \mathbb{E}(*))]$ the covariance.

Let us denote $[1 : k] := \{1, 2, \dots, k\}$, $\forall k \in \mathbb{N}^*$, and let S be the collection of all subsets of $[1 : p]$. As misuse of notation, we will denote the sets $\{i\}$ by i , and $\{ij\}$ by ij . For $u \in S$ with $u = \{u_1, \dots, u_t\}$, we set the cardinality of u as $|u| = t$ and the random subvector $\mathbf{X}_{\mathbf{u}} := (X_{u_1}, \dots, X_{u_t})$. Conventionally, if $u = \emptyset$, $|u| = 0$, and $X_{\emptyset} = 1$. Also, we denote by $\mathbf{X}_{-\mathbf{u}}$ the complementary vector of $\mathbf{X}_{\mathbf{u}}$ (that is, $-u$ is the complementary set of u). The marginal density of $\mathbf{X}_{\mathbf{u}}$ (respectively $\mathbf{X}_{-\mathbf{u}}$) is denoted by $p_{\mathbf{X}_{\mathbf{u}}}$ (resp. $p_{\mathbf{X}_{-\mathbf{u}}}$).

Further, the mathematical structure of the functional ANOVA models is defined through subspaces $(H_u)_{u \in S}$ and $(H_u^0)_{u \in S}$ of $L^2_{\mathbb{R}}(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. $H_{\emptyset} \equiv H_{\emptyset}^0$ denotes the space of constant functions. For $u \in S \setminus \{\emptyset\}$, H_u is the space of square-integrable functions that depend only on $\mathbf{X}_{\mathbf{u}}$. The space H_u^0 is defined as:

$$H_u^0 = \{h_u \in H_u, \langle h_u, h_v \rangle = 0, \forall v \subset u, \forall h_v \in H_v^0\} = H_u \cap \left(\sum_{v \subset u} H_v^0 \right)^{\perp}.$$

Through the article, we will see that $\eta(\mathbf{X})$ can be written as a functional ANOVA model in terms of low-order components. We define $d := \max_u |u|$

the order of a functional ANOVA model. Thus, if the ANOVA model is additive, $d = 1$. If it is a saturated model, the order is maximal, and $d = p$.

7.2.2 Generalized Sobol sensitivity indices

Let us suppose that

$$\boxed{\begin{array}{c} P_{\mathbf{X}} \ll \nu \\ \text{where} \\ \nu(dx) = \nu_1(dx_1) \otimes \cdots \otimes \nu_p(dx_p) \end{array}} \quad (\text{C.7})$$

Our main assumption is :

$$\boxed{\exists 0 < M \leq 1, \forall u \subseteq [1 : p], \quad p_{\mathbf{X}} \geq M \cdot p_{\mathbf{X}_u} p_{\mathbf{X}_{-u}} \quad \nu\text{-a.e.}} \quad (\text{C.8})$$

Under these conditions, the following result states a general decomposition of η as a saturated functional ANOVA model, under the specific conditions of the spaces H_u^0 (defined in Section 7.2.1),

Theorem 7.1. *Let η be any function in $L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. Then, under (C.7) and (C.8), there exist functions $\eta_{\emptyset}, \eta_1, \dots, \eta_{\{1, \dots, p\}} \in H_{\emptyset} \times H_1^0 \times \cdots \times H_{\{1, \dots, p\}}^0$ such that the following equality holds :*

$$\begin{aligned} \eta(\mathbf{X}) &= \eta_{\emptyset} + \sum_{i=1}^p \eta_i(X_i) + \sum_{1 \leq i < j \leq p} \eta_{ij}(X_i, X_j) + \cdots + \eta_{\{1, \dots, p\}}(\mathbf{X}) \\ &= \sum_{u \in \mathcal{S}} \eta_u(\mathbf{X}_u). \end{aligned} \quad (7.1)$$

Moreover, this decomposition is unique.

To get a better understanding of Theorem 7.1, the reader could refer to its proof and further explanations in [CGP12]. Notice that, unlike the Sobol decomposition with independent inputs, the component functions of (7.1) are hierarchically orthogonal, and no more mutually orthogonal. Thus, further along the article, the obtained decomposition (7.1) will be abbreviated HOFD (for Hierarchically Orthogonal Functional Decomposition). Also, as mentioned in Chapter 6, the HOFD is said to be a generalized decomposition because it turns out to be the usual functional ANOVA decomposition when incomes are independent.

The general decomposition of the output $Y = \eta(\mathbf{X})$ given in Theorem 7.1 allows for decomposing the global variance as a simplified sum of covariance terms. Further below, we define the generalized sensitivity indices able to measure the contribution of any group of inputs in the model when inputs can be dependent :

Definition 7.1. *The sensitivity index S_u of order $|u|$ measuring the contribution of $\mathbf{X}_u$ into the model is given by :*

$$S_u = \frac{V(\eta_u(\mathbf{X}_u)) + \sum_{u \cap v \neq \emptyset, v} \text{Cov}(\eta_u(\mathbf{X}_u), \eta_v(\mathbf{X}_v))}{V(Y)} \quad (7.2)$$

More specifically, the first order sensitivity index S_i is given by :

$$S_i = \frac{V(\eta_i(X_i)) + \sum_{\substack{v \neq \emptyset \\ i \notin v}} \text{Cov}(\eta_i(X_i), \eta_v(\mathbf{X}_v))}{V(Y)} \quad (7.3)$$

These indices are called generalized Sobol sensitivity indices because if all inputs are independent, it can be shown that $\text{Cov}(\eta_u, \eta_v) = 0$, $\forall u \neq v$ [CGP12].

Proposition 7.1. *Under (C.7) and (C.8), the sensitivity indices S_u previously defined sums to 1, i.e. $\sum_{u \in S \setminus \{\emptyset\}} S_u = 1$.*

Interpretation of the sensitivity indices

As the covariance term $\text{Cov}(\eta_u, \sum_{u \cap v \neq \emptyset, v} \eta_v)$ in S_u can be negative, S_u is no

more bounded between 0 and 1 as in the independent case. Hence, the interpretation of our indices here is much less obvious. However, we could interpret a sensitivity index S_u given by (7.2) as follows. The form of a sensitivity index S_u allows for distinguishing two parts: the first part, $V(\eta_u)/V(Y)$ could be identified as the full contribution of $\mathbf{X}_u$, whereas the second part, $\text{Cov}(\eta_u, \sum_{u \cap v \neq \emptyset, v} \eta_v)/V(Y)$ could be interpreted as the contribution induced

by the dependence with the other terms of the decomposition. Thus, the covariance terms would play the role of compensation. If $S_u > 0$, the covariance contribution strengthens the variance term if it is positive, whereas it weakens the full contribution if not. In the case of a negative sensitivity index, the contribution induced by the dependence dominates in the index, that would show a small full contribution of the variable itself with respect to the total sum of covariances. As an illustration, we consider a model with $p = 2$ inputs. If the component η_1 is strongly negatively correlated with η_2 ,

the variations of η_1 is going to impact on the variations of η_2 . Thus, it may result in a negative index. This tells us about the importance of the dependent part here. Nevertheless, the covariance contribution is the same in S_1 and S_2 . Thus, if S_1 is greater than S_2 , the full contribution of S_1 will be bigger than the one in S_2 . In this case, we are able to rank the input parameters.

However, these results are not new, as they are developed in [CGP12]. As a continuity of this article, we propose here to estimate the functional ANOVA components in the second part. In Theorem 7.1, we show that each component η_u belongs to a subspace H_u^0 of $L^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. Thus, to estimate η_u , the most natural approach is to construct a good approximation space of H_u^0 . The next section aims at proposing a procedure to construct such a space.

7.3 The hierarchical orthogonal Gram-Schmidt procedure

In Section 7.2, we have seen that the generalized sensitivity indices are defined for any type of reference measures $(\nu_i)_{i \in [1:p]}$. From now and until the end, we will assume that $\nu_i, \forall i \in [1 : p]$, are diffuse measures. Indeed, the non diffuse measures raise additional issues in the results developed further that we will not address in this paper.

In a Hilbert space, it is usual to call in an orthonormal basis to express any of the space element as a linear combination of these basis. Further below, we will define the finite-dimensional spaces $H_u^L \subset H_u$ and $H_u^{0,L} \subset H_u^0, \forall u \in S$, as linear spans of some orthonormal systems that will be settled later. We take the notation $\text{Span}\{B\}$ to define the set of all finite linear combination of elements of B , also called the linear span of B .

Consider, for any $i \in [1 : p]$, a truncated orthonormal system $(\psi_{l_i}^i)_{l_i=1}^{L_i}$ of $L^2(\mathbb{R}, \mathcal{B}(\mathbb{R}), P_{X_i})$, with $L_i \geq 1$. Further, let us denote the vector of sizes as $L = {}^t(L_1, \dots, L_p)$. Also, we denote by $\mathbf{l}_u = (l_u^1, \dots, l_u^i, \dots)_{i \in u}$ the multi-index associated to the tensor-product of $(\otimes_{i \in u} \psi_{l_u^i}^i)$. Hence, $\mathbf{l}_u \in \times_{i \in u} [1 : L_i]$, where $\times_{i \in u} [1 : L_i]$ denotes the Cartesian product of the set $[1 : L_i]$, for $i \in u$.

To define properly the truncated spaces $H_u^L \subset H_u$, we need first to assume that

$$\boxed{\forall u \in S, \forall l_i \in [1 : L_i], \int (\prod_{i \in u} \psi_{l_i}^i(x_i))^2 p_{\mathbf{X}} d\nu(\mathbf{x}) < +\infty} \quad (\text{C.9})$$

Remark. A sufficient condition for (C.8) is to have $0 < M_1 \leq p_{\mathbf{X}} \leq M_2$ (see Section 6.3 of Chapter 6). In this particular case, it is enough to assume that $\int (\prod_{i \in [1:p]} \psi_{l_i}^i(x_i))^2 d\nu(\mathbf{x}) < +\infty$ to warrant (C.9).

Under (C.9), we define, $H_{\emptyset}^L = \text{Span}\{1\}$. Also, we set, $\forall i \neq j \in [1 : p]$,

$$H_i^L = \text{Span}\{1, \psi_1^i, \dots, \psi_{L_i}^i\},$$

$$H_{ij}^L = \text{Span}\{1, \psi_1^i, \dots, \psi_{L_i}^i, \psi_1^j, \dots, \psi_{L_j}^j, \psi_1^i \otimes \psi_1^j, \dots, \psi_{L_i}^i \otimes \psi_{L_j}^j\},$$

where $(\psi_{l_{\{ij\}}}^i \otimes \psi_{l_{\{ij\}}}^j)$, for $\mathbf{l}_{\{ij\}} = (l_{\{ij\}}^i, l_{\{ij\}}^j) \in [1 : L_i] \times [1 : L_j]$ is the tensor product orthonormal system of $H_i^L \otimes H_j^L$. More generally, for any u such that $|u| \geq 1$, we set

$$H_u^L = \text{Span}\left\{1, (\otimes_{i \in v} \psi_{l_v^i}^i)_{\mathbf{l}_v = (l_v^i)_{i \in v} \in \times_{i \in v} [1:L_i]}, \forall v \subseteq u, v \neq \emptyset\right\}.$$

Notice that $\dim(H_u^L) = 1 + \sum_{v \subseteq u, v \neq \emptyset} \prod_{i \in v} L_i$. Now, we define the corresponding

spaces up to the hierarchical orthogonality constraints. First, $H_{\emptyset}^{0,L} = H_{\emptyset}^L$. For $u \in S \setminus \{\emptyset\}$,

$$H_u^{0,L} = \{h_u \in H_u^L, \langle h_u, h_v \rangle = 0, \forall v \subset u, \forall h_v \in H_v^{0,L}\}.$$

From now, we denote by L_u the dimension of $H_u^{0,L}$, for $u \in S \setminus \{\emptyset\}$. By definition of $H_u^{0,L}$, we get $L_u = \dim(H_u^L) - [\sum_{v \subseteq u, v \neq \emptyset} \prod_{i \in v} L_i + 1] = \prod_{i \in u} L_i$.

Suppose that we observe an independent and identically distributed sample $(y^s, \mathbf{x}^s)_{s=1, \dots, n}$ of size n from the distribution of $(Y, \mathbf{X})$. We define the empirical inner product $\langle \cdot, \cdot \rangle_n$ and norm $\|\cdot\|_n$ as

$$\langle g_1, g_2 \rangle_n = \frac{1}{n} \sum_{s=1}^n g_1(\mathbf{x}^s) g_2(\mathbf{x}^s), \quad \|g\|_n^2 = \langle g, g \rangle_n.$$

We define the finite-dimensional linear subspaces $(G_{u,n}^{0,L})_{u \in S}$ as the approximating spaces of $(H_u^{0,L})_{u \in S}$, when the scalar product $\langle \cdot, \cdot \rangle$ is replaced by the empirical one. First, $G_{\emptyset,n}^{0,L} = H_{\emptyset}^L$, and, for $u \in S \setminus \{\emptyset\}$,

$$G_{u,n}^{0,L} = \{g_u \in H_u^L, \langle g_u, g_v \rangle_n = 0, \forall v \subset u, \forall g_v \in G_{v,n}^{0,L}\}.$$

Now we have determined $(H_u^{0,L})_{u \in S}$ and $(G_{u,n}^{0,L})_{u \in S}$, we want to build them. In the next procedure, we propose an iterative scheme to construct them, taking into account their specific properties of orthogonality.

Procedure 1 [HOGS]

1. *Initialization:* For any $i \in [1 : p]$, take a truncated orthonormal system $(\psi_{l_i}^i)_{l_i=0}^{L_i}$, $L_i \geq 1$, of $L^2(\mathbb{R}, \mathcal{B}(\mathbb{R}), P_{X_i})$ such that $\psi_0^i = 1$. Set

$$\phi_{l_i}^i = \psi_{l_i}^i, \quad \forall l_i \geq 1 \text{ and } H_i^{0,L} = \text{Span} \{ \phi_1^i, \dots, \phi_{L_i}^i \}.$$

As $(\phi_{l_i}^i)_{l_i=1}^{L_i}$ is an orthonormal system, if $h \in H_i^{0,L}$, $h(X_i) = \sum_{l_i=1}^{L_i} \beta_{l_i}^i \phi_{l_i}^i(X_i)$

satisfies $\mathbb{E}(h_i(X_i)) = 0$.

2. Let $u = \{ij\}$ with $i \neq j \in [1 : p]$. To build a basis $(\phi_{\mathbf{l}_u}^u)$ of $H_u^{0,L}$, we consider the spaces $H_i^{0,L}$ and $H_j^{0,L}$ constructed in Step 1, i.e.

$$H_i^{0,L} = \text{Span} \{ \phi_1^i, \dots, \phi_{L_i}^i \} \quad \text{and} \quad H_j^{0,L} = \text{Span} \{ \phi_1^j, \dots, \phi_{L_j}^j \}.$$

To simplify notation, we write $\{ij\} := ij$ and $\mathbf{l}_{\{ij\}} := \mathbf{l}_{ij} = (l_{ij}^i, l_{ij}^j)$, for $l_{ij}^i \in [1 : L_i]$ and $l_{ij}^j \in [1 : L_j]$. To construct $(\phi_{\mathbf{l}_{ij}}^{ij})$, we proceed as follows: for all $\mathbf{l}_{ij} = (l_{ij}^i, l_{ij}^j) \in [1 : L_i] \times [1 : L_j]$,

(a) set

$$\begin{aligned} \phi_{\mathbf{l}_{ij}}^{ij}(X_i, X_j) &= \phi_{l_{ij}^i}^i(X_i) \times \phi_{l_{ij}^j}^j(X_j) + \sum_{k=1}^{L_i} \lambda_{k, \mathbf{l}_{ij}}^i \phi_k^i(X_i) \\ &\quad + \sum_{k=1}^{L_j} \lambda_{k, \mathbf{l}_{ij}}^j \phi_k^j(X_j) + C_{\mathbf{l}_{ij}}^{ij} \end{aligned} \quad (7.4)$$

- (b) compute the $(L_i + L_j + 1)$ coefficients $(C_{\mathbf{l}_{ij}}^{ij}, (\lambda_{k, \mathbf{l}_{ij}}^i)_{k=1}^{L_i}, (\lambda_{k, \mathbf{l}_{ij}}^j)_{k=1}^{L_j})$ by solving

$$\langle \phi_{\mathbf{l}_{ij}}^{ij}, \phi_k^i \rangle = 0, \quad \forall k \in [1 : L_i] \quad (7.5)$$

$$\langle \phi_{\mathbf{l}_{ij}}^{ij}, \phi_k^j \rangle = 0, \quad \forall k \in [1 : L_j] \quad (7.6)$$

$$\langle \phi_{\mathbf{l}_{ij}}^{ij}, 1 \rangle = 0. \quad (7.7)$$

- (c) When removing the constant term $C_{\mathbf{l}_{ij}}^{ij}$, the linear system that consists in (7.5)-(7.6), combined with (7.4), is equivalent to a matrix system of the form,

$$A_\phi^{ij} \lambda^{ij} = D_{ij}^{\mathbf{l}_{ij}}, \quad (7.8)$$

where

$$A_{\phi}^{ij} = \begin{pmatrix} \mathbb{E}(\Phi_i^t \Phi_i) & \mathbb{E}(\Phi_i^t \Phi_j) \\ \mathbb{E}(\Phi_j^t \Phi_i) & \mathbb{E}(\Phi_j^t \Phi_j) \end{pmatrix}, \quad \begin{aligned} (\Phi_i)_k &= \phi_k^i, & k \in [1 : L_i], \\ (\Phi_j)_k &= \phi_k^j, & k \in [1 : L_j], \end{aligned}$$

and

$$\lambda^{ij} = \begin{pmatrix} \lambda_{1, \mathbf{l}_{ij}}^i \\ \vdots \\ \lambda_{L_i, \mathbf{l}_{ij}}^i \\ \lambda_{1, \mathbf{l}_{ij}}^j \\ \vdots \\ \lambda_{L_j, \mathbf{l}_{ij}}^j \end{pmatrix}, \quad D_{ij}^{L_{ij}} = - \begin{pmatrix} \langle \phi_{l_{ij}}^i \times \phi_{l_{ij}}^j, \phi_1^i \rangle \\ \vdots \\ \langle \phi_{l_{ij}}^i \times \phi_{l_{ij}}^j, \phi_{L_i}^i \rangle \\ \vdots \\ \langle \phi_{l_{ij}}^i \times \phi_{l_{ij}}^j, \phi_1^j \rangle \\ \vdots \\ \langle \phi_{l_{ij}}^i \times \phi_{l_{ij}}^j, \phi_{L_j}^j \rangle \end{pmatrix}.$$

Notice that A_{ϕ}^{ij} is a Gram matrix, and, as shown in Lemma B.1 of Appendix B, A_{ϕ}^{ij} is a definite positive matrix. Therefore, the system (7.8) admits a unique solution $\tilde{\lambda}^{ij} = ((\tilde{\lambda}_{k, \mathbf{l}_{ij}}^i)_{k=1}^{L_i}, (\tilde{\lambda}_{k, \mathbf{l}_{ij}}^j)_{k=1}^{L_j})$. The constant $C_{\mathbf{l}_{ij}}^{ij}$ can then be deduced as

$$C_{\mathbf{l}_{ij}}^{ij} = -\mathbb{E} \left[\phi_{l_{ij}}^i \otimes \phi_{l_{ij}}^j (X_i, X_j) + \sum_{k=1}^{L_i} \tilde{\lambda}_{k, \mathbf{l}_{ij}}^i \phi_k^i (X_i) + \sum_{k=1}^{L_j} \tilde{\lambda}_{k, \mathbf{l}_{ij}}^j \phi_k^j (X_j) \right].$$

Hence, we can deduce that (7.4) has a unique expression.

At the end, set

$$H_{ij}^{0,L} = \text{Span} \left\{ \phi_1^{ij}, \dots, \phi_{L_{ij}}^{ij} \right\}, \quad \text{with } L_{ij} = L_i \times L_j.$$

3. Iterate the same procedure with $u = \{u_1, \dots, u_k\} \in S$, that is $|u| = k$, $k \geq 1$. Suppose that, for any $v \in S$ such that $1 \leq |v| \leq k-1$, we get

$$H_v^{0,L} = \text{Span} \left\{ \phi_{\mathbf{l}_v}^v, \mathbf{l}_v = (l_v^i)_{i \in v} \in \times_{i \in v} [1 : L_i] \right\}, \quad L_v := \dim(H_v^{0,L}) = \prod_{i \in v} L_i.$$

To construct a basis $(\phi_{\mathbf{l}_u}^u)_{\mathbf{l}_u \in \times_{i \in u} [1 : L_i]}$ of $H_u^{0,L}$, we proceed as follows: for all $\mathbf{l}_u = (l_u^{u_1}, \dots, l_u^{u_k}) \in [1 : L_{u_1}] \times \dots \times [1 : L_{u_k}]$,

(a) set

$$\phi_{\mathbf{l}_u}^u(\mathbf{X}_u) = (\phi_{l_1}^{u_1} \otimes \cdots \otimes \phi_{l_k}^{u_k})(\mathbf{X}_u) + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\mathbf{l}_v \in \times_{i \in v} [1:L_i]} \lambda_{\mathbf{l}_v, \mathbf{l}_u}^v \phi_{\mathbf{l}_v}^v(\mathbf{X}_v) + C_{\mathbf{l}_u}^u \quad (7.9)$$

(b) compute the $(1 + \sum_{\substack{v \subset u \\ v \neq \emptyset}} L_v)$ coefficients $(C_{\mathbf{l}_u}^u, (\lambda_{\mathbf{l}_v, \mathbf{l}_u}^v)_{\mathbf{l}_v \in \times_{i \in v} [1:L_i], v \subset u})$ by solving

$$\begin{cases} \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle = 0, & \forall v \subset u, \forall \mathbf{l}_v \in \times_{i \in v} [1:L_i] \\ \langle \phi_{\mathbf{l}_u}^u, 1 \rangle = 0. \end{cases} \quad (7.10)$$

(c) Exactly like in Step 2, the linear system (7.10) is equivalent to a matrix system of the form $A_\phi^u \lambda^u = D_u^{\mathbf{l}_u}$, when $C_{\mathbf{l}_u}^u$ has been removed. The matrix A_ϕ^u is a Gram matrix involving block terms $\mathbb{E}(\Phi_v(\mathbf{X}_v)^t \Phi_w(\mathbf{X}_w))_{v,w \subset u}$, where $\Phi_v(\mathbf{X}_v)$ is a vector of length L_v such that $(\Phi_v(\mathbf{X}_v))_{\mathbf{l}_v} = \phi_{\mathbf{l}_v}^v(\mathbf{X}_v)$, for $\mathbf{l}_v \in [1:L_v]$. λ^u is the unknown vector of coefficients, whereas $D_u^{\mathbf{l}_u}$ involves block terms $-\langle \otimes_{i=1}^k \phi_{l_i}^{u_i}, \Phi_v \rangle_{v \subset u}$. Lemma B.1 ensures the uniqueness of the system, therefore (7.9) admits a unique expression.

Set

$$H_u^{0,L} = \text{Span} \left\{ \phi_{\mathbf{l}_u}^u, \mathbf{l}_u \in \times_{i \in [1:k]} [1:L_{u_i}] \right\}, \text{ with } L_u := \dim(H_u^{0,L}) = \prod_{i=1}^k L_{u_i}.$$

The construction of $(G_{u,n}^{0,L})_{u \in S}$ is very similar to the $(H_u^{0,L})_{u \in S}$ one. However, as the spaces $(G_{u,n}^{0,L})_{u \in S}$ depend on the observed sample, their construction requires to assume that the sample size n is larger than the sizes $L_i, i \in [1:p]$.

To build $G_{i,n}^{0,L}, \forall i \in [1:p]$, we use the usual Gram-Schmidt procedure on $(\psi_{l_i}^i)_{l_i=0}^{L_i}$ to get an orthonormal system $(\varphi_{l_i}^i)_{l_i=1}^{L_i}$ with respect to the empirical inner product $\langle \cdot, \cdot \rangle_n$.

To build $(G_{u,n}^{0,L})_{u \in S}$, for u such that $|u| \geq 2$, it is enough to replace $\langle \cdot, \cdot \rangle$ by $\langle \cdot, \cdot \rangle_n$ and to use the HOGS procedure. At the end, we write

$$G_{u,n}^{0,L} = \text{Span} \left\{ \varphi_{\mathbf{l}_u}^u, \mathbf{l}_u \in \times_{i \in u} [1:L_i] \right\}, \forall u \in S \setminus \{\emptyset\}.$$

In practice, polynomials or splines basis functions [DS06] will be considered. Also, for practical reasons and motivation exposed in Section 7.5, we will approximate the model output by an ANOVA model of order at most $d = 3$. In the next section, we discuss the practical estimation of generalized Sobol

sensitivity indices using least-squares minimization. Further, we also discuss the curse of dimensionality, and propose some variable selection methods to handle it.

7.4 Estimation of the generalized sensitivity indices

7.4.1 Least-Squares estimation

The effects $(\eta_u)_{u \in S}$ in the HOFD (7.1) satisfy

$$(\eta_u)_{u \in S} = \underset{\substack{(\tilde{\eta}_u)_{u \in S} \\ \tilde{\eta}_u \in H_u^0}}{\text{Argmin}} \mathbb{E}[(Y - \sum_{u \in S} \tilde{\eta}_u(\mathbf{X}_u))^2] \quad (7.11)$$

Notice that $\eta_\emptyset$, the expected value of Y , has not interest for the sensitivity indices estimation. Thus, $\tilde{Y} := Y - \mathbb{E}(Y)$ replaces Y in (7.11). Also, the residual term $\eta_{\{1, \dots, p\}}$ is removed from (7.11) and it is estimated afterwards.

In Section 7.3, we defined the approximating spaces $G_{u,n}^{0,L}$ of H_u^0 , for $u \in S \setminus \{\emptyset\}$. Thus, the minimization problem (7.11) may be replaced by its empirical version,

$$\min_{(\beta_{\mathbf{l}_u})_{\mathbf{l}_u, u}} \frac{1}{n} \sum_{s=1}^n \left[\tilde{y}^s - \sum_{\substack{u \subset [1:p] \\ u \neq \emptyset}} \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \beta_{\mathbf{l}_u}^u \varphi_{\mathbf{l}_u, n}^u(\mathbf{x}_u^s) \right]^2, \quad (7.12)$$

where $\tilde{y}^s := y^s - \bar{y}$, $\bar{y} := \frac{1}{n} \sum_{s=1}^n y^s$, and where every component η_u is replaced by a linear combination of functions $(\varphi_{\mathbf{l}_u, n}^u)_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \in G_{u,n}^{0,L}$ constructed according to the HOGS Procedure of Section 7.3. The equivalent matrix form of (7.12) is

$$\min_{\boldsymbol{\beta}} \|\mathbb{Y} - \mathbb{X}_{\boldsymbol{\varphi}} \boldsymbol{\beta}\|_n^2, \quad (7.13)$$

where $\mathbb{Y}_s = y^s - \bar{y}$, $\mathbb{X}_{\boldsymbol{\varphi}} = \begin{pmatrix} \boldsymbol{\varphi}_1 & \dots & \boldsymbol{\varphi}_u & \dots \end{pmatrix} \in \times_{u \in S} \mathcal{M}_{n, L_u}(\mathbb{R})$, where $\times_{u \in S} \mathcal{M}_{n, L_u}(\mathbb{R})$ denotes the cartesian product of real entries matrices with n rows and L_u columns.

For $u \in S$, $(\boldsymbol{\varphi}_u)_{s, \mathbf{l}_u} = \varphi_{\mathbf{l}_u, n}^u(\mathbf{x}_u^s)$, $\forall s \in [1 : n]$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$. $(\boldsymbol{\beta})_{\mathbf{l}_u, u} = \beta_{\mathbf{l}_u}^u$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$, $\forall u \subset [1 : p]$, $u \neq \emptyset$.

Remind that the dimension of spaces $H_u^{0,L}$ is denoted by $L_u = \prod_{i \in u} L_i$, $\forall u \in S \setminus \{\emptyset\}$. Thus, the number of parameters to be estimated in (7.13) is equal to

$m := \sum_{\substack{u \subset [1:p] \\ u \neq \emptyset}} L_u$. Let us remark that it would be numerically very expensive to

consider the estimation of all these coefficients. Even for small ANOVA order d , the number of terms blows up with the dimensionality of the problem, and so does the number of model evaluations when using an ordinary least-squares regression scheme. As an illustration, take $d = 3$, $p = 8$ and $L_i = L = 5$, $\forall i \in [1 : p]$. In this case, $m = 7740$ parameters are to be estimated, which could be a difficult task in practice. To handle this problem, many variable selection methods have been considered in the field of statistics. The next section aims at briefly exposing the variable selection methods via a penalized regression. We particularly focus on the ℓ_0 penalty [Tem11] and on the Lasso regression [Tib96].

7.4.2 The variable selection methods

For simplicity, we denote by m the number of parameters in (7.13). The variable selection methods usually deal with the penalized regression

$$\min_{\boldsymbol{\beta}} \|\mathbb{Y} - \mathbb{X}_{\varphi} \boldsymbol{\beta}\|_{\mathbf{n}}^2 + \lambda J(\boldsymbol{\beta}), \quad (7.14)$$

where $J(\cdot)$ is positive valued for $\boldsymbol{\beta} \neq 0$, and where $\lambda \geq 0$ is a tuning parameter. The most intuitive approach is to consider the ℓ_0 -penalty $J(\boldsymbol{\beta}) = \|\boldsymbol{\beta}\|_0$,

where $\|\boldsymbol{\beta}\|_0 = \sum_{j=1}^m \mathbb{I}(\beta_j \neq 0)$. Indeed, the ℓ_0 regularization aims at selecting non-zero coefficients, thus at removing the useless parameters from the model.

The greedy approximation [Tem11] offers a series of strategies to deal with the ℓ_0 -penalty. Nevertheless, the ℓ_0 regularization is a non convex function, and suffers from the statistical instability, as mentioned in [Bre95, Tib96]. A convex relaxation of the optimization problem can be viewed with the Lasso regression [Tib96]. Indeed, the Lasso regression corresponds to the ℓ_1 -penalty, i.e. (7.14) with $J(\boldsymbol{\beta}) = \|\boldsymbol{\beta}\|_1$, and $\|\boldsymbol{\beta}\|_1 = \sum_{j=1}^m |\beta_j|$.

The Lasso offers a good compromise between a rough selection of non-zero elements, and a ridge regression ($J(\boldsymbol{\beta}) = \sum_{j=1}^m \beta_j^2$) that only shrinks coefficients, but is known to be stable [BvdG11, HTF09].

To offer a good panel to the reader, we will adapt our method to the ℓ_0 and to the ℓ_1 regularization. The adaptive forward-backward greedy (*FoBa*) algorithm proposed in Zhang [Zha11] is exploited here to deal with the ℓ_0 penalization. From a dictionary $\mathcal{D}$ that can be large and/or redundant, the *FoBa* algorithm is an iterative scheme that sequentially selects and deletes the element of $\mathcal{D}$ that has the least impact on the fit. The aim of the algorithm is to efficiently select a

limited number of predictors. The advantage of such approach is that it is very intuitive, and easy to implement. In our problem, the *FoBa* algorithm is applied on the whole set of basis functions. It can then happen that none basis function is retained for the estimation of a HOFD component. In this case, as we want to estimate each component of the HOFD, the coefficient corresponding to this component is set to be zero.

Initiated by Efron *et al.* [EHJT04], the modified LARS algorithm is further adapted to our problem to deal with the Lasso regression. The LARS is a general iterative technique that builds up the regression function by successive steps. The adaptation of LARS to Lasso (the modified LARS) is inspired by the homotopy method proposed by Osborne *et al.* [Os00]. The main advantage of the modified LARS algorithm is that it builds up the whole regularized solution path $\{\hat{\beta}(\lambda), \lambda \in \mathbb{R}\}$, exploiting the property of piecewise linearity of the solutions with respect to λ [BvdG11, OPT00].

In the next part, both *FoBa* algorithm and the modified LARS algorithm are adapted to our problem and they are compared via numerical examples.

7.4.3 Summary of the estimation procedure

Provided an initial good choice of orthonormal systems $(\psi_{l_i}^i)_{l_i=0}^{L_i}$, for $i \in [1 : p]$, we first construct the approximating spaces $G_{u,n}^{0,L}$ of H_u^0 for $|u| \leq d$, and $d < p$, thanks to the HOGS Procedure of Section 7.3. A HOFD component η_u is then a projection onto $G_{u,n}^{0,L}$, whose coefficients are defined by least-squares estimation. To bypass the curse of dimensionality, the *FoBa* algorithm or the modified LARS algorithm is used. Once the HOFD components are estimated, we deduce the empirical estimation of the generalized Sobol sensitivity indices given in Definition 7.1.

7.5 Asymptotic results

We suppose that

$$\eta(\mathbf{X}) \approx \eta^R(\mathbf{X}) = \sum_{u \in S} \eta_u^R(\mathbf{X}), \quad \text{with} \quad \eta_u^R = \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u) \in H_u^{0,L}.$$

In the following, we give the convergence properties of the estimator $\hat{\eta}^R$ to η^R , with

$$\hat{\eta}^R(\mathbf{X}) := \sum_{u \in S} \hat{\eta}_u^R(\mathbf{X}_u), \quad \text{with} \quad \hat{\eta}_u^R(\mathbf{X}_u) = \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \hat{\beta}_{\mathbf{l}_u}^u \varphi_{\mathbf{l}_u,n}^u(\mathbf{X}_u) \in G_{u,n}^{0,L}.$$

where $(\hat{\beta}_{\mathbf{l}_u}^u)_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]}$, $u \in S$ are estimated by the usual least-squared problem (7.13). Thus, we are interested in the convergence results when the order of truncation $L_u = \prod_{i \in u} L_i$ ($u \in S$), is fixed.

Proposition 7.2 gives the convergence result.

Proposition 7.2. *Assume that*

$$Y = \eta^R(\mathbf{X}) + \varepsilon, \quad \text{where } \eta^R(\mathbf{X}) = \sum_{u \in S} \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u) \in H_u^{0,L},$$

with $\mathbb{E}(\varepsilon) = 0$, $\mathbb{E}(\varepsilon^2) = \sigma_*^2$, $\mathbb{E}(\varepsilon \cdot \phi_{\mathbf{l}_u}^u(\mathbf{X}_u)) = 0$, $\forall \mathbf{l}_u \in \times_{i \in u} [1:L_i]$, $\forall u \in S$.
 $(\beta_0 = (\beta_{\mathbf{l}_u}^{u,0})_{\mathbf{l}_u, u} \text{ is the true parameter})$.

Further, let us consider the least-squares estimation $\hat{\eta}^R$ of η^R using the sample $(y^s, \mathbf{x}^s)_{s \in [1:n]}$ from the distribution of $(Y, \mathbf{X})$, and the functions $(\varphi_{\mathbf{l}_u}^u)_{\mathbf{l}_u, u}$, that is

$$\hat{\eta}^R(\mathbf{X}) = \sum_{u \in S} \hat{\eta}_u^R(\mathbf{X}_u), \quad \text{where } \hat{\eta}_u^R(\mathbf{X}_u) = \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \hat{\beta}_{\mathbf{l}_u}^u \varphi_{\mathbf{l}_u, n}^u(\mathbf{X}_u) \in G_{u, n}^{0,L},$$

where $\hat{\beta} = \underset{\beta \in \Theta}{\text{Argmin}} \|\mathbb{Y} - \mathbb{X}_\varphi \beta\|_n^2$.

If we assume that

(H.1) The distribution $P_{\mathbf{X}}$ is equivalent to $\otimes_{i=1}^p P_{X_i}$;

(H.2) For any $u \in S$, any $\mathbf{l}_u \in \times_{i \in u} [1:L_i]$, $\|\phi_{\mathbf{l}_u}^u\|_n = 1$ and $\|\varphi_{\mathbf{l}_u, n}^u\|_n = 1$

(H.3) For any $i \in [1:p]$, any $l_i \in [1:L_i]$, the fourth moment of $\varphi_{l_i, n}^i$ is finite.

Then,

$$\|\hat{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0 \quad \text{when } n \rightarrow +\infty. \quad (7.15)$$

For further interest, the detailed proof of Proposition 7.2 is postponed to Appendix B.

Our aim here is to study how the approximating spaces $G_{u, n}^{0,L}$, constructed with the previous procedure, behave when $n \rightarrow +\infty$. However, we assume that the order of truncation L_u ($u \in S$) is fixed. By extending the work of Stone [Sto94], Huang [Hua98] explores the convergence properties of functional ANOVA models when and L_u is not fixed anymore. Nevertheless, the results are obtained for a general model space H_u , and its approximating space G_u , $\forall u \in S$. In [Hua98], the author states that if the basis functions are m -smooth and bounded, $\|\hat{\eta} - \eta\|$ converges in probability. For polynomials, Fourier transforms or splines, he specifically shows that $\|\hat{\eta} - \eta\| = O_p(n^{-\frac{2m}{2m+d}})$ (See [Hua98] p. 257), where d is the ANOVA order. Thus, to get a good estimation in practice, it is in our interest to have a small order d in a model. In the next numerical applications, we use this theoretical result to substitute the initial model to a functional ANOVA of order at most $d = 3$.

7.6 Application

In this section, we consider two numerical applications. The first model is a toy function studied in Li *et al.* [LR12]. It is used to study the numerical properties of the practical method summarized in Section 7.4.3. The second application is the study of a shell subject to an internal pressure. From a finite elements code, we estimate the generalized sensitivity indices of input parameters implied into the model. The goal is to quantify the sensitivity of each input variable in a context of dependence.

7.6.1 A test case

Remind that the HOGS procedure is used to construct basis functions adapted to the hierarchical orthogonality constraints. Hence, the same method improved by the adaptive greedy algorithm will be abbreviated GHOGS method (G for Greedy). The method improved by the modified LARS algorithm will be called LHOGS (L for LARS).

The estimation procedure suggested in Chapter 6, and reminded in the introduction, is compared to the G/LHOGS methods. It will be denoted POM (for Projection Operators Method), relatively to the projection operators it uses.

Let $\mathbf{X} \sim N(0, \Sigma)$ and the model function,

$$Y = g_1(X_1, X_2) + g_2(X_2) + g_3(X_3),$$

with

$$\begin{aligned} g_1(X_1, X_2) &= [a_1 X_1 + a_0][b_1 X_2 + b_0] \\ g_2(X_2) &= c_2 X_2^2 + c_1 X_2 + c_0 \\ g_3(X_3) &= d_3 X_3^3 + d_2 X_3^2 + d_1 X_3 + d_0, \end{aligned}$$

and

$$\Sigma = \begin{pmatrix} \sigma_1^2 & \gamma\sigma_1\sigma_2 & 0 \\ \gamma\sigma_1\sigma_2 & \sigma_2^2 & 0 \\ 0 & 0 & \sigma_3^2 \end{pmatrix}.$$

Condition (C.8) does not allow to use the normal distribution, but rather the mixture Gaussian one [CGP12]. However, the Gaussian distribution allows for computing a HOFD decomposition, as done in [LR12]. Moreover, if the research of solutions is restricted to the polynomial spaces, the uniqueness of the HOFD components given in [LR12] is ensured, whatever the type of distribution. Thus, the analytical form of the generalized Sobol indices can be deduced in this case.

	S_1	S_2	S_3	S_{12}	S_{13}	S_{23}	S_{123}	$\ \hat{\beta}\ _0$
Analytical	0.4429	0.4718	0.0763	0.0091	0	0	0	-
POM	0.4402 (0.02)	0.4718 (0.04)	0.0810 (0.001)	-0.0014 (0.001)	- -	- -	- -	-
GHOGS	0.4499 (0.027)	0.4647 (0.033)	0.0754 (0.025)	0.0030 (0.003)	0 (0)	0 (0)	0.0070 (0.002)	5 to 7
LHOGS	0.4534 (0.027)	0.4688 (0.031)	0.0793 (0.026)	0.0060 (0.002)	0.0013 (0.001)	0.0010 (0.001)	-0.0098 (0.0002)	22 to 207

Table 7.1: Sensitivity indices estimation with the POM and the G/LHOGS methods

We take $a_0 = c_1 = d_0 = 1$, $a_1 = b_0 = c_2 = d_1 = d_2 = 2$ and $b_1 = c_0 = d_3 = 3$. The variations are fixed at $\sigma_1 = \sigma_2 = 0.2$, $\sigma_3 = 0.18$ and $\gamma = 0.6$. To show the interest of the greedy and Lasso application, we proceed to $n = 200$ model evaluations repeated 50 times. For each component, we choose a Hermite basis of degree 10. Thus, the number of parameters $m = 330 > n = 200$. In view of analytical results given in [LR12], it is clear that only 8 among 330 basis functions are necessary to recover entirely the model. Table 7.1 helps for comparing the POM with the GHOGS/LHOGS procedures on the sensitivity indices estimation and their standard deviation (indicated into brackets). Table 7.1 also provides the number of non-zero estimated coefficients for the *FoBa* and the modified LARS algorithm. The estimated components $\hat{\eta}_1$, $\hat{\eta}_2$ and $\hat{\eta}_3$ are represented in Figure 7.1.

The advantage of the GHOGS and of LHOGS methods is to be able to estimate all interaction indices, whereas the POM only estimates interaction indices involved in dependent pairs [CGP12]. Even if many of the non-zero coefficients in LARS are close to zero, this method tends to estimate a large number of non-zero parameters in comparison with the GHOGS procedure. Through this example, the greedy recovers properly the information with a relevant selection of coefficients.

7.6.2 The tank pressure model

The case study concerns a shell closed by a cap and subject to an internal pressure. Figure 7.2 illustrates a simulation of tank distortion. We are interested in the von Mises stress [vM13] on the point y labelled in Figure 7.2. The von Mises stress allows for predicting material yielding which occurs when it reaches the material yield strength. The selected point y corresponds to the point for which the von Mises stress is maximal in the tank. Therefore, we want to prevent the tank from material damage induced by plastic deforma-

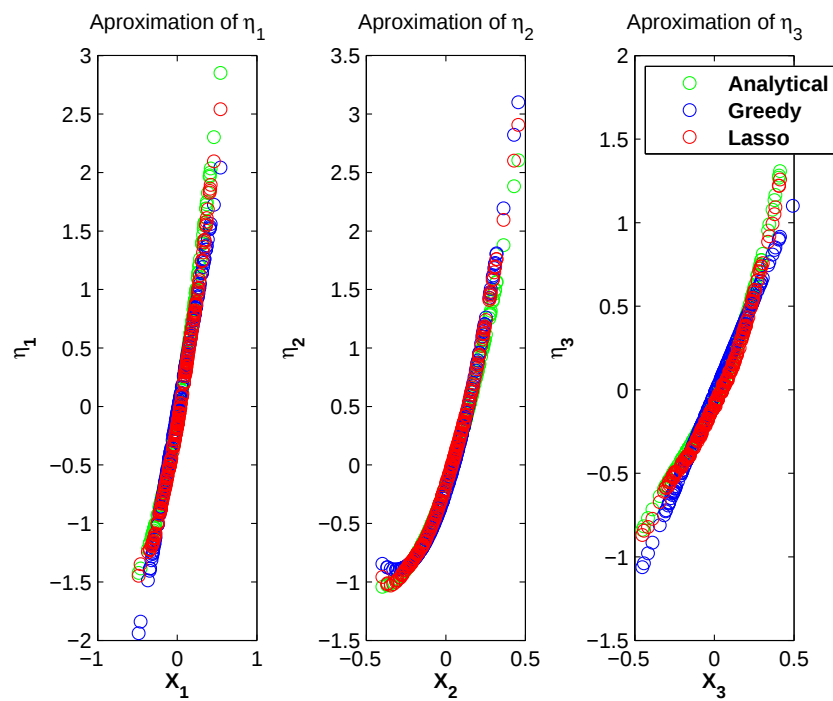
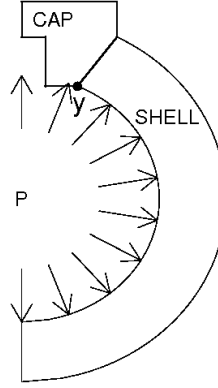


Figure 7.1: Estimation of the HOFD components by the G/LHOGS methods

Figure 7.2: Tank distortion at point y

tions. To offer a large panel of tanks able to resist to the internal pressure, a manufacturer wants to know the most contributive parameters to the von Mises criterion variability. In the model we propose, the von Mises criterion depends on three geometrical parameters: the shell internal radius (R_{int}), the shell thickness (T_{shell}), and the cap thickness (T_{cap}). It also depends on five physical parameters concerning the Young's modulus (E_{shell} and E_{cap}) and the yield strength ($\sigma_{y,shell}$ and $\sigma_{y,cap}$) of the shell and the cap. The last parameter is the internal pressure (P_{int}) applied to the shell. The system is modeled by a 2D finite elements code ASTER.

In table 7.2, we give the input distributions.

The geometrical parameters are uniformly distributed because of the large choice left for the tank building. The correlation γ between the geometrical parameters is induced by the constraints of manufacturing processes. The physical inputs are normally distributed and their uncertainty are due to the manufacturing process and the properties of the elementary constituents variabilities. The large variability of P_{int} in the model corresponds to the different internal pressure values which could be applied to the shell by the user.

To measure the contribution of the correlated inputs to the output variability, we estimate the generalized sensitivity indices by the practical method exposed in Section 7.4. We proceed to $n = 1000$ simulations over 50 runs. We use the 5-spline functions for the geometrical parameters and the Hermite basis functions of degree 7 for the physical parameters. The first order indices dispersions are displayed in Figure 7.3 for both Greedy and LARS algorithm.

We observe first that the HOGS procedure applied with the greedy and the LARS techniques give very similar results. The first four physical param-

Inputs	Distribution
R_{int}	$\mathcal{U}([1800; 2200]), \gamma(R_{int}, T_{shell}) = 0.85$
T_{shell}	$\mathcal{U}([360; 440]), \gamma(T_{shell}, T_{cap}) = 0.3$
T_{cap}	$\mathcal{U}([180; 220]), \gamma(T_{cap}, R_{int}) = 0.3$
E_{cap}	$\alpha N(\mu, \Sigma) + (1 - \alpha)N(\mu, \Omega)$
$\sigma_{y, cap}$	$\alpha = 0.02, \mu = \begin{pmatrix} 210 \\ 500 \end{pmatrix}, \Sigma = \begin{pmatrix} 350 & 0 \\ 0 & 29 \end{pmatrix}, \Omega = \begin{pmatrix} 175 & 81 \\ 81 & 417 \end{pmatrix}$
E_{shell}	$\alpha N(\mu, \Sigma) + (1 - \alpha)N(\mu, \Omega)$
$\sigma_{y, shell}$	$\alpha = 0.02, \mu = \begin{pmatrix} 70 \\ 300 \end{pmatrix}, \Sigma = \begin{pmatrix} 117 & 0 \\ 0 & 500 \end{pmatrix}, \Omega = \begin{pmatrix} 58 & 37 \\ 37 & 250 \end{pmatrix}$
P_{int}	$N(80, 10)$

Table 7.2: Description of inputs of the shell model

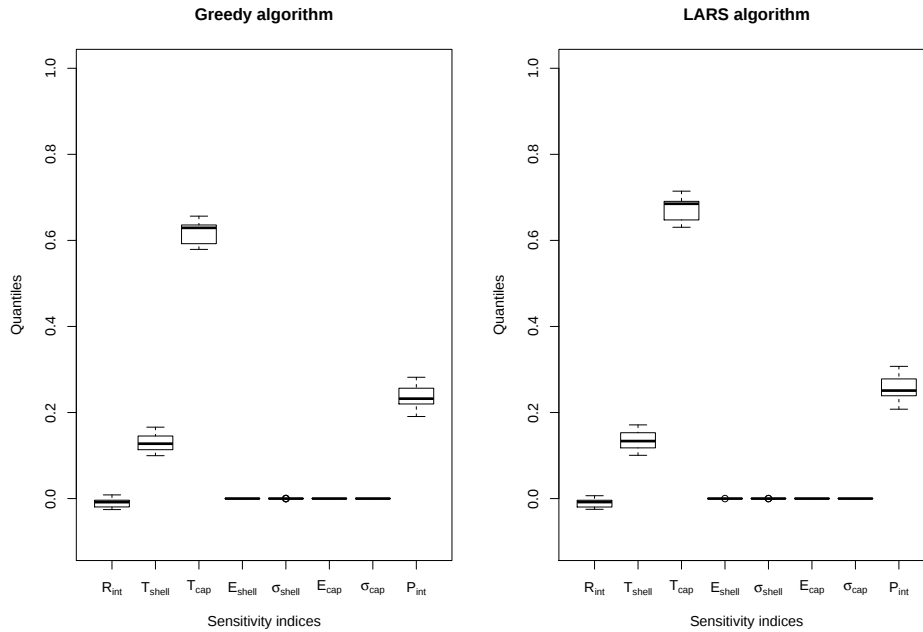


Figure 7.3: Boxplot representations of the first order sensitivity indices

ters are independent from the other inputs, and their effects are null, so we can deduce that they do not have any influence in the model. Also, even if the internal pressure plays an important role, the strongest contribution comes from the correlated set of geometrical inputs $(R_{int}, T_{shell}, T_{cap})$. The sensitivity index related to the shell radius (R_{int}) is negative, so the covariance induced by the dependence dominates in the index, showing that either there is a strong negative covariance part or the full contribution of the variable is small. In the first case, it shows that R_{int} is influent through its correlation. In the second one, the input R_{int} is not an influent variable in the model. The sensitivity indices of the shell thickness (T_{shell}) and the cap thickness (T_{cap}) reveal that these two variables have a strong influence in the model. Thus, to scale down the variability in the model, we should reduce the cap thickness variability first. Because of the strong correlation between the shell radius and its thickness, one should reduce the variability of both parameters.

Appendices

Appendix A

Generalized decomposition for dependent variables

This appendix gathers the proofs and complementary notes of the statements given in Chapter 6.

A.1 Generalized Hoeffding decomposition

A.1.1 Generalized decomposition for dependent inputs

The upcoming proof follows the guideline of the proof of Lemma 3.1 in Stone [Sto94].

Proof of Lemma 6.1.

By induction on the cardinal of T , let show that

$$\mathcal{H}(n) : \forall T | \#(T) = n, \quad \mathbb{E}[(\sum_{u \in T} h_u(\mathbf{X}))^2] \geq \delta^{\#(T)-1} \sum_{u \in T} \mathbb{E}[h_u^2(\mathbf{X})].$$

- $\mathcal{H}(1)$ is obviously true, as T is reduced to a singleton
- Let $n \in \mathbb{N}^*$. Suppose that $\mathcal{H}(n')$ is true for all $1 \leq n' \leq n$. Let T such that $\#(T) = n + 1$. We want to prove $\mathcal{H}(n + 1)$.
Choose a maximal set r of T , i.e. r is not a proper subset of any set u in T . We show first that

$$\mathbb{E}[(\sum_{u \in T} h_u(\mathbf{X}))^2] \geq M \cdot \mathbb{E}(h_r^2(\mathbf{X})). \quad (\text{A.1})$$

- If $\#(r) = p$, by definition of H_r^0 , we get $\mathbb{E}[(\sum_{u \in T} h_u(\mathbf{X}))^2] \geq \mathbb{E}(h_r^2(\mathbf{X})) \geq M \mathbb{E}(h_r^2(\mathbf{X}))$ as $M \leq 1$.

- If $1 \leq \#(r) \leq p-1$, set $\mathbf{X} = (X_1, X_2)$, where $X_1 = (X_l)_{l \notin r}$ and $X_2 = (X_l)_{l \in r}$. By Condition (C.2), it follows that

$$p_{\mathbf{X}} \geq M \cdot p_{X_1} p_{X_2}.$$

As a consequence,

$$\begin{aligned} \mathbb{E}[(\sum_{u \in T} h_u(\mathbf{X}))^2] &= \int_{\mathcal{X}_1} \int_{\mathcal{X}_2} [h_r(x_2) + \sum_{u \neq r} h_u(x_1, x_2)]^2 p_{\mathbf{X}} \nu(dx_1, dx_2) \\ &\geq M \int_{\mathcal{X}_1} \int_{\mathcal{X}_2} [h_r(x_2) \\ &\quad + \sum_{u \neq r} h_u(x_1, x_2)]^2 p_{X_1} p_{X_2} \nu_1(dx_1) \nu_2(dx_2) \\ &\geq M \int_{\mathcal{X}_1} \mathbb{E}[(h_r(X_2) + \sum_{u \neq r} h_u(x_1, X_2))^2] p_{X_1} \nu_1(dx_1) \end{aligned}$$

when $\mathcal{X}_i$ denotes the support of X_i , $i = 1, 2$. By maximality of r and by definition of H_r^0 ,

- If $u \subset r$, h_u only depends on X_2 and by orthogonality,

$$\mathbb{E}(h_u(X_2)h_r(X_2)) = 0.$$

- If $u \not\subset r$, h_u depends on X_1 fixed at x_1 , and $X_2^u = (X_l)_{l \in r \cap u}$, so $h_u(x_1, \cdot) \in H_{r \cap u}^0$, with $r \cap u \subset r$, it comes then

$$\mathbb{E}(h_u(x_1, X_2)h_r(X_2)) = 0.$$

Thus,

$$\begin{aligned} \mathbb{E}[(\sum_{u \in T} h_u(\mathbf{X}))^2] &\geq M \int_{\mathcal{X}_1} \mathbb{E}(h_r^2(X_2)) p_{X_1} \nu_1(dx_1) \\ &= M \cdot \mathbb{E}(h_r^2(\mathbf{X})). \end{aligned}$$

So (A.1) holds for any size of any maximal sets of T .

By using (A.1) with $\tilde{h}_r = h_r$ and $\tilde{h}_u = -\beta h_u$, $\forall u \neq r$, we get

$$\mathbb{E}[(h_r(\mathbf{X}) - \beta \sum_{u \neq r} h_u(\mathbf{X}))^2] \geq M \mathbb{E}(h_r^2(\mathbf{X})). \quad (\text{A.2})$$

Now, for a better understanding, we set $x = h_r(\mathbf{X})$ and $y = \sum_{u \neq r} h_u(\mathbf{X})$.

Thus, (A.2) may be rewritten as

$$\mathbb{E}[(x - \beta y)^2] \geq M \mathbb{E}(x^2). \quad (\text{A.3})$$

Taking $\beta = \frac{\mathbb{E}(xy)}{\mathbb{E}(y^2)}$ in (A.3), it follows that:

$$\mathbb{E}(x^2) - \frac{\mathbb{E}(xy)^2}{\mathbb{E}(y^2)} \geq M\mathbb{E}(x^2).$$

Hence,

$$\mathbb{E}(xy)^2 \leq (1 - M) \cdot \mathbb{E}(x^2)\mathbb{E}(y^2) \quad (\text{A.4})$$

Further, as $\mathbb{E}[(x + y)^2] \geq 0$, it comes that,

$$\begin{aligned} \mathbb{E}(x^2) + \mathbb{E}(y^2) &\geq -2\mathbb{E}(xy) \\ &\geq -\frac{2}{\sqrt{1-M}}\mathbb{E}(xy), \quad \text{as } \sqrt{1-M} \leq 1 \\ &\geq 2\sqrt{\mathbb{E}(x^2)}\sqrt{\mathbb{E}(y^2)} \quad \text{by (A.4)} \end{aligned} \quad (\text{A.5})$$

Inequalities (A.5) and (A.4) imply that

$$\begin{aligned} \mathbb{E}[(x + y)^2] &\geq \mathbb{E}(x^2) + \mathbb{E}(y^2) - 2\sqrt{1-M}\sqrt{\mathbb{E}(x^2)}\sqrt{\mathbb{E}(y^2)} \\ &\geq (1 - \sqrt{1-M})[\mathbb{E}(x^2) + \mathbb{E}(y^2)] \\ &\geq \delta[\mathbb{E}(x^2) + \mathbb{E}(y^2)] \end{aligned} \quad (\text{A.6})$$

Now, we replace x and y by their true value. As $\mathcal{H}(n)$ is supposed to be true and (A.6) holds, it follows that:

$$\begin{aligned} \mathbb{E}[(\sum_u h_u(\mathbf{X}))^2] &\geq \delta \left[\mathbb{E}(h_r^2(\mathbf{X})) + \delta^{n-1} \sum_{u \neq r} \mathbb{E}(h_u^2(\mathbf{X})) \right] \\ &\geq \delta^n \sum_u \mathbb{E}(h_u^2(\mathbf{X})) \quad \text{as } \delta \in]0, 1] \\ &= \delta^{\#(T)-1} \sum_u \mathbb{E}(h_u^2(\mathbf{X})). \end{aligned}$$

Hence, $\mathcal{H}(n+1)$ holds.

We can deduce that $\mathcal{H}(n)$ is true for any collection T of $[1 : p]$. ■

Proof of Theorem 6.1.

Let define the vector space $K^0 = \{ \sum_{u \in S^-} h_u(\mathbf{X}_u), h_u \in H_u^0, \forall u \in S^- \}$.

In the first step, we will prove that K^0 is a complete space to prove the existence and uniqueness of the projection of η in K^0 , thanks to the projection theorem [Lue97].

Secondly, we will show that η is exactly equal to the decomposition into H^0 , and finally, we will see that each term of the summand is unique.

- We show that H_u^0 is closed into H_u (as H_u is a Hilbert space).

Let $(h_{n,u})_n$ be a convergent sequence of H_u^0 with $h_{n,u} \rightarrow h_u$. As $(h_{n,u})_n \in H_u^0 \subset H_u$ complete, $h_u \in H_u$. Let $v \subset u$, and $h_v \in H_v^0$:

$$\begin{array}{ccc} \langle h_u - h_{n,u}, h_v \rangle & = & \langle h_u, h_v \rangle - \langle h_{n,u}, h_v \rangle \\ \downarrow & & \parallel \\ 0 & & 0 \quad \text{as } H_u^0 \perp H_v^0. \end{array}$$

Thus, $\langle h_u, h_v \rangle = 0$, so that $h_u \in H_u^0$. H_u^0 is then a complete space.

Let $(h_n)_n$ be a Cauchy sequence in K^0 and we show that each component is of Cauchy and that $h_n \rightarrow h \in K^0$.

As $h_n \in K^0$, $h_n = \sum_{u \in S^-} h_{n,u}$, $h_{n,u} \in H_u^0$. It follows that:

$$\begin{aligned} \|h_n - h_m\|^2 &= \left\| \sum_u (h_{n,u} - h_{m,u}) \right\|^2 \\ &\geq \delta^{\#(S^-)-1} \sum_{u \in S^-} \|h_{n,u} - h_{m,u}\|^2 \quad \text{by Inequality (6.4).} \end{aligned}$$

As $(h_n)_n$ is a Cauchy sequence, by the above inequality, $(h_{n,u})_n$ is also Cauchy. As $h_{n,u} \rightarrow h_u \in H_u^0$, we deduce that $h_n \xrightarrow[n \rightarrow \infty]{} \sum_{u \in S^-} h_u = h \in K^0$.

Thus, K^0 is complete. By the projection theorem, we can deduce there exists a unique element into K^0 such that:

$$\left\| \eta - \sum_{u \in S^-} \eta_u \right\|^2 \leq \|\eta - h\|^2 \quad \forall h \in K^0.$$

- Decomposition of η : following Hooker [Hoo07], we introduce the residual term as

$$\eta_{1,\dots,p}(X_1, \dots, X_p) = \eta(X_1, \dots, X_p) - \sum_{u \in S^-} \eta_u(\mathbf{X}_u).$$

By projection, $\langle \eta - \sum_{v \in S^-} \eta_v, h_u \rangle = 0 \quad \forall u \in S^-, \forall h_u \in H_u^0$. Hence,

$\eta(\mathbf{X}) = \sum_{u \in S} \eta_u(\mathbf{X}_u)$, $\eta_u \in H_u^0$, $\forall u \in S$, and this decomposition is well defined.

- Terms of the summand are unique: assume that $\eta = \sum_{u \in S} \eta_u = \sum_{u \in S} \tilde{\eta}_u$, $\tilde{\eta}_u \in H_u^0$.

By Lemma 6.1, it follows that

$$\left. \begin{aligned} \sum_{u \in S} (\eta_u - \tilde{\eta}_u) &= 0 \\ \left\| \sum_{u \in S} (\eta_u - \tilde{\eta}_u) \right\|^2 &\geq \delta^{\#(S)-1} \sum_{u \in S} \|\eta_u - \tilde{\eta}_u\|^2 \end{aligned} \right\} \Rightarrow \|\eta_u - \tilde{\eta}_u\|^2 = 0 \quad \forall u \in S$$

■

A.1.2 Generalized sensitivity indices

Proof of Proposition 6.1.

Under (C.1) and (C.2), Theorem 6.1 states the existence and the uniqueness of the decomposition of η as $\eta(\mathbf{X}) = \sum_{u \in S} \eta_u(\mathbf{X}_u)$, with $\eta_u \in H_u^0 \perp H_v^0$,

$\forall v \subset u, u, v \in S$.

Therefore, $\mathbb{E}(\eta(\mathbf{X})) = \mathbb{E} \left[\sum_{u \in S} \eta_u(\mathbf{X}_u) \right] = \eta_\emptyset$, and

$$\begin{aligned} V(Y) = V(\eta(\mathbf{X})) &= \mathbb{E}(\eta^2(\mathbf{X})) - \eta_\emptyset^2 \\ &= \sum_{u \neq \emptyset} \mathbb{E}(\eta_u^2(\mathbf{X}_u)) + \sum_{u \neq v} \mathbb{E}(\eta_u(\mathbf{X}_u) \eta_v(\mathbf{X}_v)) \\ &= \sum_{u \neq \emptyset} V(\eta_u(\mathbf{X}_u)) + \sum_{\substack{u \neq \emptyset \\ u \neq [1:p]}} \sum_{\substack{v \neq \emptyset \\ u \not\subset v, v \not\subset u}} \mathbb{E}(\eta_u(\mathbf{X}_u) \eta_v(\mathbf{X}_v)) \\ &= \sum_{u \neq \emptyset} \left[V(\eta_u(\mathbf{X}_u)) + \sum_{\substack{v \\ u \cap v \neq u, v}} \text{Cov}(\eta_u(\mathbf{X}_u), \eta_v(\mathbf{X}_v)) \right]. \end{aligned}$$

Thus, the sensitivity of $\mathbf{X}_u$ in the model may be measured by the comparison between $V(\eta_u(\mathbf{X}_u)) + \sum_{\substack{v \\ u \cap v \neq u, v}} \text{Cov}[\eta_u(\mathbf{X}_u), \eta_v(\mathbf{X}_v)]$ and the global variance

$V(Y)$. Thus, we define (6.6) and (6.7) as sensitivity measures. The equality (6.8) follows from the decomposition of the global variance.

■

A.2 Examples of distribution function

A.2.1 Boundedness of the inputs density function

Proof of Example 6.1.

– ν is a product of measure as $\frac{d\nu}{d\nu_L} = \prod_{i=1}^p \nu_i(x_i)$, with $\nu_i \sim N(m_i, \sigma_i^2)$.

– $p_{\mathbf{X}}$ is given by

$$\begin{aligned} p_{\mathbf{X}}(\mathbf{x}) &= \frac{dP_{\mathbf{X}}}{d\nu}(\mathbf{x}) = \frac{dP_{\mathbf{X}}}{d\nu_L} \times \frac{d\nu_L}{d\nu}(\mathbf{x}) \\ &= \alpha + (1 - \alpha) \left| \frac{\Sigma}{\Omega} \right|^{1/2} \exp -\frac{1}{2} t(\mathbf{x} - m)(\Omega^{-1} - \Sigma^{-1})(\mathbf{x} - m) \quad (\text{A.7}) \end{aligned}$$

First, we have $p_{\mathbf{X}}(\mathbf{x}) \geq \alpha > 0$. Further, the sufficient and necessary condition to have $p_{\mathbf{X}} \leq M_2 < \infty$ is to get $(\Omega^{-1} - \Sigma^{-1})$ positive definite. Indeed, if $(\Omega^{-1} - \Sigma^{-1})$ admits a negative eigenvalue, $p_{\mathbf{X}}$ can not be bounded. Thus, $0 < \alpha \leq p_{\mathbf{X}} \leq M_2$ iff $(\Omega^{-1} - \Sigma^{-1})$ is positive definite. ■

A.2.2 Examples of distribution of two inputs

Proof of Proposition 6.3.

Condition (C.5) is immediate with Equation (6.9). Let prove that (C.5) is equivalent to (C.6).

If (C.6) holds, then $c(u, v) \geq M$. Conversely, we assume that $0 < M < 1$, and

$$\tilde{C}(u, v) = \frac{C(u, v) - Muv}{1 - M}.$$

It is enough to show it is a copula: Obviously, $\tilde{C}(0, u) = C(u, 0) = 0$, $\tilde{C}(1, u) = \tilde{C}(u, 1) = u \quad \forall u \in [0, 1]$. By second order derivation, it comes that $\tilde{c}(u, v) = \frac{c(u, v) - M}{1 - M}$, so $\tilde{c}(u, v) \geq 0$ by hypothesis (C.5). ■

Let consider the class of Archimedian copulas,

$$C(u, v) = \varphi^{-1}[\varphi(u) + \varphi(v)], \quad u, v \in [0, 1], \quad (\text{A.8})$$

where the generator φ is a non negative two times differentiable function defined on $[0, 1]$ with $\varphi(1) = 0$, $\varphi'(u) < 0$ and $\varphi''(u) > 0$, $\forall u \in [0, 1]$.

A sufficient condition for (C.5) is given in Proposition A.1:

Proposition A.1. *If there exist $M_1, M_2 > 0$ such that:*

$$-\varphi'(u) \geq M_1 \quad \forall u \in [0, 1] \quad (\text{A.9})$$

$$\frac{d}{du} \left(\frac{1}{2} \frac{1}{\varphi'(u)^2} \right) \geq M_2, \quad \forall u \in [0, 1]. \quad (\text{A.10})$$

Then, Condition (C.5) holds.

The proof is straightforward. Now, we will see two illustrative Archimedian copulas satisfying Proposition A.1.

Example A.1. Let $\alpha < 0$, $\theta > 0$ and β with $\beta < -\alpha e^{-\theta}$. Set

$$\varphi_2(x) = -\frac{\alpha}{\theta}e^{-\theta x} + \beta x + \left(\frac{\alpha}{\theta}e^{-\theta} - \beta\right), \quad x \in [0, 1]. \quad (\text{A.11})$$

Example A.2. Let $C < 0$ and set

$$\varphi_3(x) = x \ln x + (C - 1)x + (1 - C), \quad x \in [0, 1]. \quad (\text{A.12})$$

A.3 Estimation

A.3.1 Model of $p = 2$ input variables

Proof of Proposition 6.5.

- We first show first that $(\mathcal{S})$ admits an unique solution. Under (C.1) and (C.2), by Theorem 6.1, the decomposition of $\eta(\mathbf{X})$ is unique and

$$\eta(X_1, X_2) = \eta_\emptyset + \eta_1(X_1) + \eta_2(X_2) + \eta_{12}(X_1, X_2).$$

$$\text{with } \begin{cases} \eta_\emptyset \in H_\emptyset \\ \eta_i \in H_i^0 \perp H_\emptyset, \quad i = 1, 2 \\ \eta_{12} \in H_{12}^0 \perp H_i^0, \quad i = 1, 2, H_{12}^0 \perp H_\emptyset. \end{cases}$$

Thus,

$$\begin{pmatrix} \text{I} & 0 & 0 & 0 \\ 0 & \text{I} & P_{H_1^0} & 0 \\ 0 & P_{H_2^0} & \text{I} & 0 \\ 0 & 0 & 0 & \text{I} \end{pmatrix} \begin{pmatrix} \eta_\emptyset \\ \eta_1 \\ \eta_2 \\ \eta_{12} \end{pmatrix} = \begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix}. \quad (\text{A.13})$$

So $(\eta_\emptyset, \eta_1, \eta_2, \eta_{12})$ is solution of $(\mathcal{S})$. Now, assume there exists an another solution of the system, say $(\tilde{\eta}_\emptyset, \dots, \tilde{\eta}_{[1:p]}) \in H_\emptyset \times \dots \times H_{[1:p]}^0$, then

$$\begin{aligned} \begin{cases} \eta_\emptyset - \tilde{\eta}_\emptyset = 0 \\ \eta_1 - \tilde{\eta}_1 + P_{H_1^0}(\eta_2 - \tilde{\eta}_2) = 0 \\ P_{H_2^0}(\eta_1 - \tilde{\eta}_1) + \eta_2 - \tilde{\eta}_2 = 0 \\ \eta_{12} - \tilde{\eta}_{12} = 0 \end{cases} &\Rightarrow \begin{cases} \eta_\emptyset = \tilde{\eta}_\emptyset \\ P_{H_1^0}(\eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2) = 0 \\ P_{H_2^0}(\eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2) = 0 \\ \eta_{12} = \tilde{\eta}_{12} \end{cases} \\ &\Rightarrow \begin{cases} \eta_\emptyset = \tilde{\eta}_\emptyset \\ \eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2 \in H_1^{0\perp} \cap H_2^{0\perp} \\ \eta_{12} = \tilde{\eta}_{12}. \end{cases} \end{aligned}$$

As $\eta_1 - \tilde{\eta}_1 \in H_1^0$ and $\eta_2 - \tilde{\eta}_2 \in H_2^0$, it follows that

$$\begin{aligned} & \begin{cases} \langle \eta_1 - \tilde{\eta}_1, \eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2 \rangle = 0 \\ \langle \eta_2 - \tilde{\eta}_2, \eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2 \rangle = 0 \end{cases} \\ & \Rightarrow \begin{cases} \|\eta_1 - \tilde{\eta}_1\|^2 + \langle \eta_1 - \tilde{\eta}_1, \eta_2 - \tilde{\eta}_2 \rangle = 0 \\ \|\eta_2 - \tilde{\eta}_2\|^2 + \langle \eta_1 - \tilde{\eta}_1, \eta_2 - \tilde{\eta}_2 \rangle = 0 \end{cases} \\ & \Rightarrow \|\eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2\|^2 = 0 \\ & \Rightarrow \eta_1 - \tilde{\eta}_1 + \eta_2 - \tilde{\eta}_2 = 0. \end{aligned}$$

As 0 can be uniquely decomposed into H^0 as a sum of zero, then, $\eta_1 - \tilde{\eta}_1 = \eta_2 - \tilde{\eta}_2 = 0$.

– Let now compute
$$\begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix}.$$

First of all, it is obvious that the constant term satisfies $\eta_\emptyset = \mathbb{E}(\eta)$ and that η_{12} is obtained by subtracting from η all other terms on the right of the decomposition.

Now, let us use the projector's property of embedded spaces. Indeed, as $H_i^0 \subset H_i$, $\forall i = 1, 2$, it comes

$$P_{H_i^0}(\eta) = P_{H_i^0}(P_{H_i}(\eta)) = P_{H_i^0}[\underbrace{\mathbb{E}(\eta|X_i)}_{\varphi(X_i)}].$$

φ is a function of X_i , so it can be decomposed into the following expression:

$$\varphi(X_i) = h_\emptyset + \varphi_i(X_i), \quad h_\emptyset \in H_\emptyset, \quad \varphi_i \in H_i^0,$$

with $h_\emptyset = \mathbb{E}(\varphi) = \mathbb{E}(\eta)$. Hence, one can easily deduce $P_{H_i^0}(\eta)$, $i = 1, 2$, as the term $\varphi_i = \mathbb{E}(\eta|X_i) - \mathbb{E}(\eta)$

We obtain

$$\begin{pmatrix} P_{H_\emptyset}(\eta) \\ P_{H_1^0}(\eta) \\ P_{H_2^0}(\eta) \\ P_{H_{12}^0}(\eta) \end{pmatrix} = \begin{pmatrix} \mathbb{E}(\eta) \\ \mathbb{E}(\eta|X_1) - \mathbb{E}(\eta) \\ \mathbb{E}(\eta|X_2) - \mathbb{E}(\eta) \\ \eta - \mathbb{E}(\eta|X_1) - \mathbb{E}(\eta|X_2) + \mathbb{E}(\eta) \end{pmatrix}. \quad (\text{A.14})$$

■

Appendix B

Generalized sensitivity indices and numerical methods

In this part, we restate and prove Proposition 7.2 of Section 7.5 of Chapter 7. For sake of clarity, we first define and recall some notation that will be used further.

B.1 Notation

B.1.1 Reminder

First, as mentioned in Section 7.4.1 of Chapter 7, we assume that Y is centered. Recall that, $\forall i \in [1 : p]$, $L_i := L_{\{i\}}$ is the dimension of the spaces $H_i^{0,L}$ and $G_{i,n}^{0,L}$. More generally, $\phi_{\{i\}} := \phi_i$. Also, $L_u := \prod_{i \in u} L_i$ is the dimension of the spaces $H_u^{0,L}$ and $G_{u,n}^{0,L}$.

For $u \in S$, $\mathbf{l}_u = (l_u^i)_{i \in u}$ is a multi-index of $\times_{i \in u} [1 : L_i]$, where $\times_{i \in u} [1 : L_i]$ is the Cartesian product of the sets $[1 : L_i]$, for $i \in u$.

We refer $(\phi_{\mathbf{l}_u}^u)_{\mathbf{l}_u \in \times_{i \in u} [1 : L_i]}$ as the basis of $H_u^{0,L}$ and $(\varphi_{\mathbf{l}_{u,n}}^u)_{\mathbf{l}_{u,n} \in \times_{i \in u} [1 : L_i]}$ as the basis of $G_{u,n}^{0,L}$ constructed according to HOGS Procedure of Section 7.3 of the main document. Thus, these functions all lie in $L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$. For simplicity, we will also assume that

$$\|\phi_{\mathbf{l}_u}^u\| = 1, \quad \text{and} \quad \|\varphi_{\mathbf{l}_{u,n}}^u\|_n = 1.$$

$\langle \cdot, \cdot \rangle$ and $\|\cdot\|$ are used as the inner product and norm on $L_{\mathbb{R}}^2(\mathbb{R}^p, \mathcal{B}(\mathbb{R}^p), P_{\mathbf{X}})$,

$$\langle h_1, h_2 \rangle = \int h_1(\mathbf{x}) h_2(\mathbf{x}) p_{\mathbf{X}} d\nu(\mathbf{x}), \quad \|h\|^2 = \langle h, h \rangle,$$

while $\langle \cdot, \cdot \rangle_n$ and $\|\cdot\|_n$ denote the empirical inner product and norm, that is

$$\langle g_1, g_2 \rangle_n = \frac{1}{n} \sum_{s=1}^n g_1(\mathbf{x}^s) g_2(\mathbf{x}^s), \quad \|g\|_n^2 = \langle g, g \rangle_n,$$

when $(y^s, \mathbf{x}^s)_{s=1, \dots, n}$ is the n -sample of observations from the distribution of $(Y, \mathbf{X})$.

B.1.2 New settings

We set $m := \sum_{u \in S} L_u$ the number of parameters in the regression model. De-

note, for all $u \in S$, $\Phi_u(\mathbf{X}_u) \in (L^2(\mathbb{R}, \mathcal{B}(\mathbb{R}), P_{\mathbf{X}}))^{L_u}$, with $(\Phi_u(\mathbf{X}_u))_{\mathbf{l}_u} = \phi_{\mathbf{l}_u}^u(\mathbf{X}_u)$, and by β any vector of $\Theta \subset \mathbb{R}^m$, where $(\beta)_{\mathbf{l}_u, u} = \beta_{\mathbf{l}_u}^u$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$.

Recall that, for $a, b \in \mathbb{N}^*$, $\mathcal{M}_{a,b}(\mathbb{R})$ denotes the set of all real matrices with a rows and b columns. Set $\mathbb{X}_\varphi = \begin{pmatrix} \varphi_1 & \dots & \varphi_u & \dots \end{pmatrix} \in \times_{u \in S} \mathcal{M}_{n, L_u}(\mathbb{R})$, where, for $u \in S$, $(\varphi_u)_{s, \mathbf{l}_u} = \varphi_{\mathbf{l}_u, n}^u(\mathbf{x}_u^s)$, $\forall s \in [1 : n]$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$.

Also, we set $\mathbb{X}_\phi = \begin{pmatrix} \phi_1 & \phi_2 & \dots \end{pmatrix} \in \times_{u \in S} \mathcal{M}_{n, L_u}(\mathbb{R})$, where, for $u \in S$, $(\phi_u)_{s, \mathbf{l}_u} = \phi_{\mathbf{l}_u}^u(\mathbf{x}_u^s)$, $\forall s \in [1 : n]$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$.

Denote by A_ϕ be the $m \times m$ Gram matrix whose block entries are $(\mathbb{E}(\Phi_u(\mathbf{X}_u)^t \Phi_v(\mathbf{X}_v)))_{u, v \in S}$.

At last, we define the functions

$$M_n(\beta) = \|\mathbb{Y} - \mathbb{X}_\varphi \beta\|_n^2, \quad (\text{B.1})$$

where $\mathbb{Y}_s = y^s$, $\forall s \in [1 : n]$. The Euclidean norm will be denoted $\|\cdot\|_2$.

Proposition . Assume that

$$Y = \eta^R(\mathbf{X}) + \varepsilon, \quad \text{where } \eta^R(\mathbf{X}) = \sum_{u \in S} \sum_{\mathbf{l}_u \in \times_{i \in u} [1 : L_i]} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u) \in H_u^{0,L},$$

with $\mathbb{E}(\varepsilon) = 0$, $\mathbb{E}(\varepsilon^2) = \sigma_*^2$, $\mathbb{E}(\varepsilon \cdot \phi_{\mathbf{l}_u}^u(\mathbf{X}_u)) = 0$, $\forall \mathbf{l}_u \in \times_{i \in u} [1 : L_i]$, $\forall u \in S$.

$(\beta_0 = (\beta_{\mathbf{l}_u}^{u,0})_{\mathbf{l}_u, u})$ is the true parameter).

Further, let us consider the least-squares estimation $\hat{\eta}^R$ of η^R using the sample $(y^s, \mathbf{x}^s)_{s \in [1:n]}$ and the functions $(\varphi_{\mathbf{l}_u}^u)_{\mathbf{l}_u, u}$, that is

$$\hat{\eta}^R(\mathbf{X}) = \sum_{u \in S} \hat{\eta}_u^R(\mathbf{X}_u), \quad \text{where } \hat{\eta}_u^R(\mathbf{X}_u) = \sum_{\mathbf{l}_u \in \times_{i \in u} [1 : L_i]} \hat{\beta}_{\mathbf{l}_u}^u \varphi_{\mathbf{l}_u, n}^u(\mathbf{X}_u) \in G_{u,n}^{0,L},$$

where $\hat{\beta} = \underset{\beta \in \Theta}{\text{Argmin}} M_n(\beta)$.

If we assume that

(H.1) The distribution $P_{\mathbf{X}}$ is equivalent to $\otimes_{i=1}^p P_{X_i}$;

(H.2) For any $u \in S$, any $\mathbf{l}_u \in \times_{i \in u} [1 : L_i]$, $\|\phi_{\mathbf{l}_u}^u\| = 1$ and $\|\varphi_{\mathbf{l}_u, n}^u\|_n = 1$

(H.3) For any $i \in [1 : p]$, any $l_i \in [1 : L_i]$, the fourth moment of $\varphi_{l_i, n}^i$ is finite.

Then,

$$\|\hat{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0 \text{ when } n \rightarrow +\infty. \quad (\text{B.2})$$

The proof of Proposition is broken up into Lemmas B.1-B.4. To prove (B.2), we introduce $\bar{\eta}^R$ as the following approximation of η^R ,

$$\bar{\eta}^R = \sum_{u \in S} \bar{\eta}_u^R(\mathbf{X}_u) = \sum_{u \in S} \sum_{\mathbf{l}_u \in \times_{i \in u} [1 : L_i]} \beta_{\mathbf{l}_u}^{u,0} \varphi_{\mathbf{l}_u, n}^u(\mathbf{X}_u) = \mathbb{X}_{\varphi} \beta_0,$$

and we write the triangular inequality,

$$\|\hat{\eta}^R - \eta^R\| = \|\hat{\eta}^R - \bar{\eta}^R + \bar{\eta}^R - \eta^R\| \leq \|\hat{\eta}^R - \bar{\eta}^R\| + \|\bar{\eta}^R - \eta^R\|. \quad (\text{B.3})$$

Thus, it is enough to prove that $\|\hat{\eta}^R - \bar{\eta}^R\| \xrightarrow{a.s.} 0$, and that $\|\bar{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0$. Lemmas B.3 and B.4 deal with convergence results on $\|\bar{\eta}^R - \eta^R\|$ and on $\|\hat{\eta}^R - \bar{\eta}^R\|$, respectively. Lemmas B.1, B.2 are preliminary results to prove Lemmas B.3 and B.4.

B.2 Preliminary results

Lemma B.1. *If (H.1) holds, then A_{ϕ} is a non singular matrix.*

Proof of Lemma B.1.

First of all notice that when we consider a Gram matrix, by a classical argument on the associated quadratic form, the full rank of this matrix holds if and only if the associated functional vector has full rank in L^2 [Hol88].

To begin with, set, for all $i \in [1 : p]$, $\Psi_i = \begin{pmatrix} 1 \\ \Phi_i \end{pmatrix}$ and $G_i := \mathbb{E}(\Psi_i^t \Psi_i)$.

As $(\phi_{l_i}^i)_{l_i=1}^{L_i}$ is an orthonormal system, we obviously get $G_i = \mathbf{I}_{(L_i+1) \times (L_i+1)}$. Now we may rewrite the tensor product $\otimes_{i=1}^p G_i$ as

$$\otimes_{i=1}^p G_i = \otimes_{i=1}^p \mathbb{E}(\Psi_i^t \Psi_i) = \int \otimes_{i=1}^p [\Psi_i(x_i)^t \Psi_i(x_i)] dP_{X_1}(x_1) \cdots dP_{X_p}(x_p). \quad (\text{B.4})$$

We obviously have $\otimes_{i=1}^p G_i = \mathbf{I}$. So that, using the remark of the beginning of the proof, the system $\otimes_{i=1}^p [\Psi_i^t \Psi_i] = \left(1 \quad (\otimes_{i \in u} \Phi_i)_{\substack{u \subseteq [1:p] \\ u \neq \emptyset}} \right)$ is linearly independent $(\otimes_{i=1}^p P_{X_i})$ -a.e.

As we assumed that $\otimes_{i=1}^p P_{X_i}$ and $P_{\mathbf{X}}$ are equivalent by (H.1), we get that $\left(1 \quad (\otimes_{i \in u} \Phi_i)_{\substack{u \subseteq [1:p] \\ u \neq \emptyset}} \right)$ is linearly independent $P_{\mathbf{X}}$ -a.e..

Now, we may conclude as in the classical Gram-Schmidt construction. Indeed, the construction of the system $(\Phi_u)_{u \in S}$ involves an invertible triangular matrix. ■

Lemma B.2. *Let $u, v \in S$ and $\mathbf{l}_u \in \times_{i \in u} [1 : L_i]$, $\mathbf{l}_v \in \times_{i \in v} [1 : L_i]$. Assume that (H.2) holds. Further, assume that $\|\varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u\| \xrightarrow{a.s.} 0$, $\|\varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u\|_n \xrightarrow{a.s.} 0$, $\|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\| \xrightarrow{a.s.} 0$ and $\|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n \xrightarrow{a.s.} 0$. Then, the following results hold:*

- (i) $\|\varphi_{\mathbf{l}_u, n}^u\| \xrightarrow{a.s.} 1$ and $\|\varphi_{\mathbf{l}_u, n}^u\|_n \xrightarrow{a.s.} 1$;
- (ii) $\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v \rangle \xrightarrow{a.s.} \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle$ and $\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v \rangle_n \xrightarrow{a.s.} \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle$;
- (iii) $\langle \varphi_{\mathbf{l}_u, n}^u, \varphi_{\mathbf{l}_v, n}^v \rangle \xrightarrow{a.s.} \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle$ and $\langle \varphi_{\mathbf{l}_u, n}^u, \varphi_{\mathbf{l}_v, n}^v \rangle_n \xrightarrow{a.s.} \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle$;
- (iv) For $u \in S$, with $|u| \geq 1$, and $\mathbf{l}_u \in \times_{i \in u} [1 : L_i]$,

$$\left\| \prod_{i \in u} \varphi_{\mathbf{l}_i, n}^i - \prod_{i \in u} \phi_{\mathbf{l}_i}^i \right\|_n \xrightarrow{a.s.} 0 \implies \langle \prod_{i \in u} \varphi_{\mathbf{l}_i, n}^i, \varphi_{\mathbf{l}_v, n}^v \rangle_n \xrightarrow{a.s.} \langle \prod_{i \in u} \phi_{\mathbf{l}_i}^i, \phi_{\mathbf{l}_v}^v \rangle.$$

Proof of Lemma B.2.

The first point (i) is trivial. Now, we have, by (H.2),

$$\begin{aligned} |\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v \rangle - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| &= |\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v \rangle| \\ &\leq \underbrace{\|\phi_{\mathbf{l}_u}^u\|}_{=1} \|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\| \xrightarrow{a.s.} 0. \end{aligned}$$

Further,

$$\begin{aligned} |\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| &\leq |\langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v \rangle_n| + |\langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| \\ &= \|\phi_{\mathbf{l}_u}^u\|_n \|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n + |\langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle|. \end{aligned}$$

By the usual strong law of large numbers, $|\langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| \xrightarrow{a.s.} 0$. By hypothesis, $\|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n \xrightarrow{a.s.} 0$, and $\|\phi_{\mathbf{l}_u}^u\|_n \xrightarrow{a.s.} 1$ by (H.2). Hence, (ii) holds.

The point (iii) follows from

$$\begin{aligned} |\langle \varphi_{\mathbf{l}_u, n}^u, \varphi_{\mathbf{l}_v, n}^v \rangle - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| &= |\langle \varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v \rangle + \langle \phi_{\mathbf{l}_u}^u, \varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v \rangle| \\ &\leq \|\varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u\| \|\varphi_{\mathbf{l}_v, n}^v\| + \|\phi_{\mathbf{l}_u}^u\| \|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|. \end{aligned}$$

By assumptions, $\|\varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u\| \xrightarrow{a.s.} 0$ and $\|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\| \xrightarrow{a.s.} 0$. Thus, the first point of (iii) is satisfied, as $\|\varphi_{\mathbf{l}_v, n}^v\| \xrightarrow{a.s.} \|\phi_{\mathbf{l}_v}^v\| = 1$ (by (i)). Further,

$$\begin{aligned} |\langle \varphi_{\mathbf{l}_u, n}^u, \varphi_{\mathbf{l}_v, n}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| &\leq |\langle \varphi_{\mathbf{l}_u, n}^u, \varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v \rangle_n| + |\langle \varphi_{\mathbf{l}_u, n}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| \\ &= \|\varphi_{\mathbf{l}_u, n}^u\|_n \|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n + |\langle \varphi_{\mathbf{l}_u, n}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle|. \end{aligned}$$

First, $\|\varphi_{\mathbf{l}_u, n}^u\|_n = 1$. By hypothesis, $\|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n \xrightarrow{a.s.} 0$.

By (ii), $|\langle \varphi_{\mathbf{l}_u, n}^u, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \phi_{\mathbf{l}_u}^u, \phi_{\mathbf{l}_v}^v \rangle| \xrightarrow{a.s.} 0$, so we can conclude.

Let show (iv). We have,

$$\begin{aligned} \left| \langle \prod_{i \in u} \varphi_{\mathbf{l}_u, n}^i, \varphi_{\mathbf{l}_v, n}^v \rangle_n - \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle \right| &\leq \left| \langle \prod_{i \in u} \varphi_{\mathbf{l}_u, n}^i - \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \varphi_{\mathbf{l}_v, n}^v \rangle_n \right| \\ &\quad + \left| \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \varphi_{\mathbf{l}_v, n}^v \rangle_n - \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle \right| \\ &\leq \left\| \prod_{i \in u} \varphi_{\mathbf{l}_u, n}^i - \prod_{i \in u} \phi_{\mathbf{l}_u}^i \right\|_n + \left\| \prod_{i \in u} \phi_{\mathbf{l}_u}^i \right\|_n \|\varphi_{\mathbf{l}_v, n}^v - \phi_{\mathbf{l}_v}^v\|_n \\ &\quad + \left| \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle \right|. \end{aligned}$$

By the strong law of large numbers, $\left| \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle_n - \langle \prod_{i \in u} \phi_{\mathbf{l}_u}^i, \phi_{\mathbf{l}_v}^v \rangle \right| \xrightarrow{a.s.} 0$, and we can conclude with the previous arguments. ■

B.3 Main convergence results

Lemma B.3. *Remind that the true regression function is*

$$\eta^R(\mathbf{X}) = \sum_{u \in S} \eta_u^R(\mathbf{X}_u), \text{ where } \eta_u^R(\mathbf{X}_u) = \sum_{\substack{\mathbf{l}_u \in \times_{i \in u} [1:L_i]}} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u) \in H_u^{0,L}.$$

Further, let $\bar{\eta}^R$ be the approximation of η^R ,

$$\bar{\eta}^R(\mathbf{X}) = \sum_{u \in S} \bar{\eta}_u^R(\mathbf{X}_u), \text{ where } \bar{\eta}_u^R(\mathbf{X}_u) = \sum_{\substack{\mathbf{l}_u \in \times_{i \in u} [1:L_i]}} \beta_{\mathbf{l}_u}^{u,0} \varphi_{\mathbf{l}_u,n}^u(\mathbf{X}_u) \in G_{u,n}^{0,L}.$$

Then, under (H.2)-(H.3), we have

$$\|\bar{\eta}_u^R - \eta_u^R\| \xrightarrow{a.s.} 0 \quad \forall u \in S, \quad \text{and} \quad \|\bar{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0.$$

Proof of Lemma B.3.

For any $u \in S$,

$$\|\bar{\eta}_u^R - \eta_u^R\| = \left\| \sum_{\mathbf{l}_u=1}^{L_u} \beta_{\mathbf{l}_u}^{u,0} \varphi_{\mathbf{l}_u,n}^u - \sum_{\mathbf{l}_u=1}^{L_u} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u \right\| \leq \sum_{\mathbf{l}_u=1}^{L_u} |\beta_{\mathbf{l}_u}^{u,0}| \cdot \|\varphi_{\mathbf{l}_u,n}^u - \phi_{\mathbf{l}_u}^u\|.$$

Let us show that $\|\varphi_{\mathbf{l}_u,n}^u - \phi_{\mathbf{l}_u}^u\| \xrightarrow{a.s.} 0$. Actually, the proof of this convergence requires the use of Lemma B.2, so we also have to show that $\|\varphi_{\mathbf{l}_u,n}^u - \phi_{\mathbf{l}_u}^u\|_n \xrightarrow{a.s.} 0$. These two results are going to be proved by a double induction on $|u| \geq 1$ and on $\mathbf{l}_u \in \times_{i \in u} [1:L_i]$. We set

$$(\mathcal{H}_k) \quad \forall u, |u| = k, \quad \begin{cases} \|\varphi_{\mathbf{l}_u,n}^u - \phi_{\mathbf{l}_u}^u\| \xrightarrow{a.s.} 0 \\ \|\varphi_{\mathbf{l}_u,n}^u - \phi_{\mathbf{l}_u}^u\|_n \xrightarrow{a.s.} 0 \quad \forall \mathbf{l}_u = (l_i^i)_{i \in u} \in \times_{i \in u} [1:L_i]. \end{cases}$$

Let us show that $(\mathcal{H}_k)$ is true for any $k \leq d$:

- Let $u = \{i\}$, so $k = 1$. We used the Gram-Schmidt procedure on $(\phi_{l_i}^i)_{l_i=1}^{L_i}$ to construct $(\varphi_{l_i,n}^i)_{l_i=1}^{L_i}$. Let us show by induction on l_i that $\|\varphi_{l_i,n}^i - \phi_{l_i}^i\| \xrightarrow{a.s.} 0$, $\forall s_i \in [1:L_i]$. Set

$$(\mathcal{H}'_{l_i}) \quad \begin{cases} \|\varphi_{l_i,n}^i - \phi_{l_i}^i\| \xrightarrow{a.s.} 0 \\ \|\varphi_{l_i,n}^i - \phi_{l_i}^i\|_n \xrightarrow{a.s.} 0. \end{cases}$$

- For $l_i = 1$, $\varphi_{1,n}^i = \frac{\phi_1^i - \langle \phi_1^i, \varphi_{0,n}^i \rangle_n \varphi_{0,n}^i}{\underbrace{\|\phi_1^i - \langle \phi_1^i, \varphi_{0,n}^i \rangle_n \varphi_{0,n}^i\|_n}_{D_n^{i,1}}}$, with $\varphi_{0,n}^i = \phi_0^i = 1$. So,

$$\|\varphi_{1,n}^i - \phi_1^i\| \leq \left\| \frac{1 - D_n^{i,1}}{D_n^{i,1}} \phi_1^i \right\| + \left\| \frac{\langle \phi_1^i, 1 \rangle_n}{D_n^{i,1}} \right\|.$$

As $(\phi_1^i)_{l_i=1}^{L_i}$ is an orthonormal system, we get $|\langle \phi_1^i, 1 \rangle_n| \xrightarrow{a.s.} \mathbb{E}(\phi_1^i) = 0$ and $D_n^{i,1} \xrightarrow{a.s.} \|\phi_1^i\| = 1$.

Therefore, $\|\varphi_{1,n}^i - \phi_1^i\| \xrightarrow{a.s.} 0$. Also,

$$\|\varphi_{1,n}^i - \phi_1^i\|_n \leq \frac{1 - D_n^{i,1}}{D_n^{i,1}} \|\phi_1^i\|_n + \frac{|\langle \phi_1^i, 1 \rangle_n|}{D_n^{i,1}}.$$

Exactly with the same previous argument, we conclude that $\|\varphi_{1,n}^i - \phi_1^i\|_n \xrightarrow{a.s.} 0$, then $(\mathcal{H}'_1)$ is true.

- Let $l_i \in [1 : L_i]$. Suppose that $(\mathcal{H}'_k)$ is true for any $k \leq l_i$. Let us show $(\mathcal{H}'_{l_i+1})$ holds. By construction, we get,

$$\varphi_{l_i+1}^i = \frac{\phi_{l_i+1}^i - \sum_{k=0}^{l_i} \langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n \cdot \varphi_{k,n}^i}{\underbrace{\|\phi_{l_i+1}^i - \sum_{k=0}^{l_i} \langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n \cdot \varphi_{k,n}^i\|_n}_{D_n^{i,l_i+1}}}.$$

So,

$$\|\varphi_{l_i+1,n}^i - \phi_{l_i+1}^i\| \leq \left\| \frac{1 - D_n^{i,l_i+1}}{D_n^{i,l_i+1}} \phi_{l_i+1}^i \right\| + \sum_{k=0}^{l_i} \left\| \frac{1}{D_n^{i,l_i+1}} \langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n \cdot \varphi_{k,n}^i \right\|.$$

For all $k \leq l_i$,

$$\begin{aligned} |\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n| &= |\langle \phi_{l_i+1}^i, \phi_k^i \rangle_n + \langle \phi_{l_i+1}^i, \varphi_{k,n}^i - \phi_k^i \rangle_n| \\ &\leq |\langle \phi_{l_i+1}^i, \phi_k^i \rangle_n| + |\langle \phi_{l_i+1}^i, \varphi_{k,n}^i - \phi_k^i \rangle_n| \\ &\leq |\langle \phi_{l_i+1}^i, \phi_k^i \rangle_n| + \|\phi_{l_i+1}^i\|_n \cdot \|\varphi_{k,n}^i - \phi_k^i\|_n. \end{aligned}$$

By induction, $\|\varphi_{k,n}^i - \phi_k^i\| \xrightarrow{a.s.} 0$. By the usual law of large numbers, $|\langle \phi_{l_i+1}^i, \phi_k^i \rangle_n| \xrightarrow{a.s.} |\langle \phi_{l_i+1}^i, \phi_k^i \rangle| = 0$ as the system $(\phi_l^i)_{l=1}^{L_i}$ is orthonormal. As $\|\phi_{l_i+1}^i\|_n \xrightarrow{a.s.} \|\phi_{l_i+1}^i\| = 1$, we deduce that $|\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n| \xrightarrow{a.s.} 0$. And,

$$\|\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n \varphi_{k,n}^i\| \leq \|\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n\|^2 \underbrace{\|(\varphi_{k,n}^i)^2\|^{1/2}}_{<+\infty} \quad \text{by (H.3)}.$$

Also,

$$D_n^{i,l_i+1} \xrightarrow{a.s.} \|\phi_{l_i+1}^i\| = 1 \Rightarrow \|\varphi_{l_i+1,n}^i - \phi_{l_i+1}^i\| \xrightarrow{a.s.} 0.$$

Now,

$$\|\varphi_{l_i+1,n}^i - \phi_{l_i+1}^i\|_n \leq \frac{1 - D_n^{i,l_i+1}}{D_n^{i,l_i+1}} \|\phi_{l_i+1}^i\|_n + \sum_{k=0}^{l_i} \frac{1}{D_n^{i,l_i+1}} |\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n|.$$

With the previous arguments, $\langle \phi_{l_i+1}^i, \varphi_{k,n}^i \rangle_n \xrightarrow{a.s.} 0$, and $D_n^{i,l_i+1} \xrightarrow{a.s.} 1$. Then, we conclude that $(\mathcal{H}'_{l_i+1})$ is true.

Therefore, $(\mathcal{H}_1)$ is satisfied.

- Let $k \in [1 : p]$. Suppose now that $(\mathcal{H}_{|\tilde{u}|})$ is true for any $1 \leq |\tilde{u}| \leq k-1$, and any $l_{\tilde{u}} \in \times_{i \in \tilde{u}} [1 : L_i]$. Show that $(\mathcal{H}_k)$ is satisfied. Let u be such that $|u| = k$.

First, as $(\mathcal{H}_{|\tilde{u}|})$ is true for any $1 \leq |\tilde{u}| \leq k-1$, results (i)-(ii)-(iii) of Lemma B.2 can be applied to any couple $(\phi_{l_u}^u, \varphi_{l_{u,n}}^u)$ such that $|u| \leq k-1$.

Further, we have seen that, for any $l_u = (l_u^i) \in \times_{i \in u} [1 : L_i]$,

$$\varphi_{l_{u,n}}^u = \prod_{i \in u} \varphi_{l_{u,n}^i}^i + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\substack{l_v \in \times_{j \in v} [1 : L_j]}} \lambda_{l_v, l_u}^{v,n} \varphi_{l_v,n}^v + C_{l_u}^{n,u},$$

where $(C_{l_u}^{n,u}, (\lambda_{l_v, l_u}^{v,n})_{l_v, v \subset u})$ are computed by the resolution of the following equations

$$\begin{cases} \langle \varphi_{l_{u,n}}^u, \varphi_{l_v,n}^v \rangle_n = 0, & \forall v \subset u, \forall l_v = (l_v^j) \in \times_{j \in v} [1 : L_j] \\ \langle \varphi_{l_{u,n}}^u, 1 \rangle_n = 0. \end{cases} \quad (\text{B.5})$$

The resolution of (B.5) leads to the resolution of a linear system, when removing $C_{l_u}^{n,u}$, of the type

$$A^{u,n} \Lambda^{u,n} = D^{l_u,n},$$

where $\Lambda^{u,n}$ is the vector of unknown parameters $(\lambda_{l_v, l_u}^{v,n})_{l_v \in \times_{j \in v} [1 : L_j], v \subset u}$, $A^{u,n}$ is the matrix whose block entries are $(\langle \varphi_{v_1}, \varphi_{v_2} \rangle_n)_{v_1, v_2 \subset u}$, and $D^{l_u,n}$ involves block entries $(-\langle \otimes_{i \in u} \varphi_{l_u^i}^i, \varphi_{v_i} \rangle_n)_{v_i \subset u}$.

Also, the theoretical construction of the functions $(\phi_{l_u}^u)_{l_u, u}$ consists in setting

$$\phi_{l_u}^u = \prod_{i \in u} \phi_{l_u^i}^i + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\substack{l_v \in \times_{j \in v} [1 : L_j]}} \lambda_{l_v, l_u}^v \phi_{l_v}^v + C_{l_u}^u,$$

where $(C_{l_u}^u, (\lambda_{l_v, l_u}^v)_{l_v, v \subset u})$ are computed by the resolution of the following equations

$$\begin{cases} \langle \phi_{l_u}^u, \phi_{l_v}^v \rangle = 0, & \forall v \subset u, \forall l_v = (l_v^j) \in \times_{j \in v} [1 : L_j] \\ \langle \phi_{l_u}^u, 1 \rangle = 0. \end{cases} \quad (\text{B.6})$$

The resolution of (B.6) leads to the resolution of a linear system, when removing $C_{l_u}^u$, of the type

$$A_{\phi}^u \Lambda^u = D^{l_u},$$

where Λ^u is the vector of unknown parameters $(\lambda_{\mathbf{l}_v, \mathbf{l}_u}^v)_{\mathbf{l}_v \in \times_{j \in v} [1:L_j], v \subset u}$, A_ϕ^u is the matrix whose block entries are $(\langle \phi_{v_1}, \phi_{v_2} \rangle)_{v_1, v_2 \subset u}$, and $D^{\mathbf{l}_u}$ involves block entries $(-\langle \otimes_{i \in u} \phi_{l_u^i}^i, \phi_{v_i} \rangle)_{v_i \subset u}$.

We have,

$$\begin{aligned}
 \|\varphi_{\mathbf{l}_u, \mathbf{n}}^u - \phi_{\mathbf{l}_u}^u\| &= \left\| \prod_{i \in u} \varphi_{l_u^i, \mathbf{n}}^i - \prod_{i \in u} \phi_{l_u^i}^i + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\mathbf{l}_v \in \times_{j \in v} [1:L_j]} (\lambda_{\mathbf{l}_v, \mathbf{l}_u}^{v, \mathbf{n}} \varphi_{\mathbf{l}_v, \mathbf{n}}^v - \lambda_{\mathbf{l}_v, \mathbf{l}_u}^v \phi_{\mathbf{l}_v}^v) \right. \\
 &\quad \left. + C_{\mathbf{l}_u}^{n, u} - C_{\mathbf{l}_u}^u \right\| \\
 &\leq \left\| \prod_{i \in u} \varphi_{l_u^i, \mathbf{n}}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\| + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\mathbf{l}_v \in \times_{j \in v} [1:L_j]} \left\| \lambda_{\mathbf{l}_v, \mathbf{l}_u}^{v, \mathbf{n}} \varphi_{\mathbf{l}_v, \mathbf{n}}^v - \lambda_{\mathbf{l}_v, \mathbf{l}_u}^v \phi_{\mathbf{l}_v}^v \right\| \\
 &\quad + |C_{\mathbf{l}_u}^{n, u} - C_{\mathbf{l}_u}^u|,
 \end{aligned} \tag{B.7}$$

and,

$$\begin{aligned}
 \|\varphi_{\mathbf{l}_u, \mathbf{n}}^u - \phi_{\mathbf{l}_u}^u\|_{\mathbf{n}} &\leq \left\| \prod_{i \in u} \varphi_{l_u^i, \mathbf{n}}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\|_{\mathbf{n}} + \sum_{\substack{v \subset u \\ v \neq \emptyset}} \sum_{\mathbf{l}_v \in \times_{j \in v} [1:L_j]} \left\| \lambda_{\mathbf{l}_v, \mathbf{l}_u}^{v, \mathbf{n}} \varphi_{\mathbf{l}_v, \mathbf{n}}^v - \lambda_{\mathbf{l}_v, \mathbf{l}_u}^v \phi_{\mathbf{l}_v}^v \right\|_{\mathbf{n}} \\
 &\quad + |C_{\mathbf{l}_u}^{n, u} - C_{\mathbf{l}_u}^u|.
 \end{aligned} \tag{B.8}$$

First, we show that

$$\left\{ \begin{array}{l} \left\| \prod_{i \in u} \varphi_{l_u^i, \mathbf{n}}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\| \xrightarrow{a.s.} 0, \quad \forall l_u^i \in [1 : L_i] \\ \left\| \prod_{i \in u} \varphi_{l_u^i, \mathbf{n}}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\| \xrightarrow{a.s.} 0, \quad \forall l_u^i \in [1 : L_i]. \end{array} \right. \tag{B.9}$$

Each univariate function $\varphi_{l_u^i, \mathbf{n}}^i$ is constructed from $(\phi_k^i)_{k \leq l_u^i}$ by the Gram-Schmidt procedure. Thus,

$$\begin{aligned}
\left\| \prod_{i \in u} \varphi_{l_u^i, n}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\| &= \left\| \prod_{i \in u} \left[\frac{\phi_{l_u^i}^i - \sum_{k=0}^{l_u^i-1} \langle \phi_{l_u^i}^i, \varphi_{k,n}^i \rangle_n \cdot \varphi_{k,n}^i}{D_n^{i, l_u^i}} \right] - \prod_{i \in u} \phi_{l_u^i}^i \right\| \\
&\leq \left\| \prod_{i \in u} \phi_{l_u^i}^i \left(\prod_{i \in u} \frac{1}{D_n^{i, l_u^i}} - 1 \right) \right\| \\
&\quad + \sum_{\substack{s+t=k \\ (s,t) \in \mathbb{N} \times \mathbb{N}^* \\ 1 \leq i_1 < \dots < i_s \leq k \\ 1 \leq j_1 < \dots < j_t \leq k}} \sum_{k=0}^{l_u^1-1} \dots \sum_{k=0}^{l_u^t-1} \|a_{i_1} \dots a_{i_s} \cdot b_{j_1} \dots b_{j_t}\|
\end{aligned}$$

where $a_i = \phi_{l_u^i}^i / D_n^{i, l_u^i}$, and $b_j = \langle \phi_{l_u^j}^j, \varphi_{k,n}^j \rangle_n \cdot \varphi_{k,n}^j / D_n^{j, l_u^j}$. As already proved, for all $i \in [1 : p]$, $l_u^i \in [1 : L_i]$,

$$D_n^{i, l_u^i} \xrightarrow{a.s.} 1.$$

Also, we previously showed that

$$\left\| \langle \phi_{l_u^j}^j, \varphi_{k,n}^j \rangle_n \varphi_{k,n}^j \right\| \xrightarrow{a.s.} 0, \quad \left\| \langle \phi_{l_u^j}^j, \varphi_{k,n}^j \rangle_n \varphi_{k,n}^j \right\|_n \xrightarrow{a.s.} 0, \quad \forall j, \forall l_u^j.$$

Thus, we conclude that (B.9) is satisfied.

Secondly, as $\left\| \prod_{i \in u} \varphi_{l_u^i, n}^i - \prod_{i \in u} \phi_{l_u^i}^i \right\|_n \xrightarrow{a.s.} 0$, Assertion (iv) of Lemma B.2 can be applied. Assertion (iii) claims that $A^{u,n}$ tends to the theoretical matrix A_ϕ^u . Also, by (iv) of Lemma B.2, $D^{l_u, n} \xrightarrow{a.s.} D^{l_u}$. Hence, $\Lambda^{u,n} \xrightarrow{a.s.} \lambda^u$. We also deduce that $C_{l_u}^{n,u} \xrightarrow{a.s.} C_{l_u}^u$.

Consequently, by induction, we deduce that every piece of the right-hand side of (B.7) (respectively (B.8)) tends to 0, so is $\|\varphi_{l_u, n}^u - \phi_{l_u}^u\|$ (resp. $\|\varphi_{l_u, n}^u - \phi_{l_u}^u\|_n$). Hence, $(\mathcal{H}_k)$ is satisfied.

As a conclusion, $\|\varphi_{l_u, n}^u - \phi_{l_u}^u\| \xrightarrow{a.s.} 0, \forall l_u \in \times_{i \in u} [1 : L_i], \forall u \in S$. Hence, we deduce that

$$\|\bar{\eta}_u^R - \eta_u^R\| \xrightarrow{a.s.} 0, \quad \forall u \in S,$$

and

$$\|\bar{\eta}^R - \eta^R\| \leq \sum_{u \in S} \|\bar{\eta}_u^R - \eta_u^R\| \implies \|\bar{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0.$$

■

Lemma B.4. Recall that $\hat{\beta} = \underset{\beta \in \Theta}{\text{Argmin}} M_n(\beta)$. If (H.1)–(H.3) hold, then

$$\left\| \hat{\beta} - \beta_0 \right\|_2 \xrightarrow{a.s.} 0 \quad (\text{B.10})$$

Moreover,

$$\left\| \hat{\eta}^R - \bar{\eta}^R \right\| \xrightarrow{a.s.} 0. \quad (\text{B.11})$$

Proof of Lemma B.4.

First, we remind the true regression model,

$$Y = \eta^R(\mathbf{X}) + \varepsilon, \quad \text{where} \quad \eta^R(\mathbf{X}) = \sum_{u \in S} \sum_{\mathbf{l}_u \in \times_{i \in u} [1:L_i]} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u), \quad (\text{B.12})$$

with $\mathbb{E}(\varepsilon) = 0$, $\mathbb{E}(\varepsilon^2) = \sigma_*^2$, $\mathbb{E}(\varepsilon \cdot \phi_{\mathbf{l}_u}^u(\mathbf{X}_u)) = 0$, $\forall \mathbf{l}_u \in \times_{i \in u} [1:L_i]$, $\forall u \in S$,
 and $\beta_0 = (\beta_{\mathbf{l}_u}^{u,0})_{\mathbf{l}_u, u \in S}$ the true parameter.

Let $\tilde{\beta}$ be the solution of the least-squared problem

$$\min_{\beta \in \Theta} \left\| \mathbb{Y} - \mathbb{X}_\phi \beta \right\|_n^2. \quad (\text{B.13})$$

Due to Lemma B.1, $({}^t X_\phi \mathbb{X}_\phi)^{-1}$ is well defined. Thus,

$$\tilde{\beta} - \beta_0 = \left(\frac{{}^t \mathbb{X}_\phi \mathbb{X}_\phi}{n} \right)^{-1} \cdot \frac{{}^t \mathbb{X}_\phi \cdot \varepsilon}{n}.$$

By the law of large numbers, $\left(\frac{{}^t \mathbb{X}_\phi \cdot \varepsilon}{n} \right)_u \xrightarrow{a.s.} \mathbb{E}(\varepsilon \cdot \phi_{\mathbf{l}_u}^u(\mathbf{X}_u)) = 0$, $\forall u \in S$.

Moreover, $\frac{{}^t \mathbb{X}_\phi \mathbb{X}_\phi}{n} \xrightarrow{a.s.} A_\phi$, where A_ϕ is defined in the new settings. Thus,
 $\left\| \tilde{\beta} - \beta_0 \right\|_2 \xrightarrow{a.s.} 0$.

Under (H.2)–(H.3), we have, by Proof of Lemma B.3, that

$$\left\| \varphi_{\mathbf{l}_u, n}^u - \phi_{\mathbf{l}_u}^u \right\| \xrightarrow{a.s.} 0. \quad (\text{B.14})$$

We are going to use (B.14) to show that $\left\| \hat{\beta} - \tilde{\beta} \right\|_2 \xrightarrow{a.s.} 0$.

As $\hat{\beta}$ is the solution of the ordinary least squares problem, we get $\hat{\beta} = ({}^t \mathbb{X}_\varphi \mathbb{X}_\varphi)^{-1} {}^t \mathbb{X}_\varphi \mathbb{Y}$ because, as seen later ${}^t \mathbb{X}_\varphi \mathbb{X}_\varphi \xrightarrow{a.s.} A_\phi$, that is invertible.

We define the usual matrix norm $\|\cdot\|_2$ as $\|A\|_2 := \sup_{\|x\|_2=1} \|Ax\|_2$. The Frobe-

nus matrix norm is defined as $\|A\|_F := \sqrt{\text{Trace}(A^t A)}$, and $\|A\|_2 \leq \|A\|_F$ [GvL96].

We use this inequality to get

$$\begin{aligned}
\|\hat{\beta} - \tilde{\beta}\|_2 &= \|(({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\varphi - ({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi)\mathbb{Y}\|_2 \\
&\leq \|({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\varphi - ({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi\|_2 \cdot \|\mathbb{Y}\|_2 \\
&= \sqrt{n} \|({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\varphi - ({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi\|_F \cdot \left\| \frac{\mathbb{Y}}{\sqrt{n}} \right\|_2.
\end{aligned}$$

First, by (B.12),

$$\left\| \frac{\mathbb{Y}}{\sqrt{n}} \right\|_2 = \left\| \frac{\mathbb{X}_\phi\beta_0 + \varepsilon}{\sqrt{n}} \right\|_2 \leq \left\| \frac{\mathbb{X}_\phi\beta_0}{\sqrt{n}} \right\|_2 + \left\| \frac{\varepsilon}{\sqrt{n}} \right\|_2$$

$$\begin{aligned}
\text{and } \left\| \frac{\mathbb{X}_\phi\beta_0}{\sqrt{n}} \right\|_2 &= \sqrt{\frac{1}{n} \sum_{s=1}^n \left[\sum_{u \in S} \sum_{\mathbf{l}_u} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{x}_{\mathbf{u}^s}) \right]^2} \xrightarrow{a.s.} ({}^t\beta_0 A_\phi \beta_0)^{1/2}. \text{ Also,} \\
\left\| \frac{\varepsilon}{\sqrt{n}} \right\|_2 &\xrightarrow{a.s.} \sqrt{\mathbb{E}(\varepsilon^2)} = \sigma_*. \text{ Hence, } \left\| \frac{\mathbb{Y}}{\sqrt{n}} \right\|_2 \leq ({}^t\beta_0 A_\phi \beta_0)^{1/2} + \sigma_* < \infty.
\end{aligned}$$

Now, let us consider $\sqrt{n} \|({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\varphi - ({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi\|_F$. After computation, we get

$$\begin{aligned}
n \|({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\varphi - ({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi\|_F^2 &= \text{Trace}\left[\left(\frac{{}^t\mathbb{X}_\phi\mathbb{X}_\phi}{n}\right)^{-1}\right] + \text{Trace}\left[\left(\frac{{}^t\mathbb{X}_\varphi\mathbb{X}_\varphi}{n}\right)^{-1}\right] \\
&\quad - 2 \text{Trace}\left[\left(\frac{{}^t\mathbb{X}_\phi\mathbb{X}_\phi}{n}\right)^{-1} \cdot \frac{{}^t\mathbb{X}_\phi\mathbb{X}_\varphi}{n} \cdot \left(\frac{{}^t\mathbb{X}_\varphi\mathbb{X}_\varphi}{n}\right)^{-1}\right]. \\
&- \frac{{}^t\mathbb{X}_\phi\mathbb{X}_\phi}{n} \xrightarrow{a.s.} A_\phi \\
&- \frac{{}^t\mathbb{X}_\varphi\mathbb{X}_\varphi}{n} \xrightarrow{a.s.} A_\phi, \text{ by (iii) of Lemma B.2} \\
&- \frac{{}^t\mathbb{X}_\phi\mathbb{X}_\varphi}{n} \xrightarrow{a.s.} A_\phi \text{ by (ii) of Lemma B.2.}
\end{aligned}$$

Under (H.1), using the result of Lemma B.1, A_ϕ is invertible. Then,

$$n \|({}^t\mathbb{X}_\phi\mathbb{X}_\phi)^{-1}{}^t\mathbb{X}_\phi - ({}^t\mathbb{X}_\varphi\mathbb{X}_\varphi)^{-1}{}^t\mathbb{X}_\phi\|_F^2 \xrightarrow{a.s.} \text{Trace}(A_\phi^{-1}) + \text{Trace}(A_\phi^{-1}) - 2\text{Trace}(A_\phi^{-1}) = 0.$$

Thus, $\|\hat{\beta} - \tilde{\beta}\|_2 \xrightarrow{a.s.} 0$. We conclude that

$$\|\hat{\beta} - \beta_0\|_2 \leq \|\hat{\beta} - \tilde{\beta}\|_2 + \|\tilde{\beta} - \beta_0\|_2 \xrightarrow{a.s.} 0.$$

At last,

$$\|\hat{\eta}^R - \bar{\eta}^R\| \leq \|\mathbb{X}_\varphi\|_2 \|\hat{\beta} - \beta_0\|_2 \xrightarrow{a.s.} 0.$$

■

Finally, as a consequence of Lemmas [B.3-B.4](#),

$$\|\hat{\eta}^R - \eta^R\| \leq \|\hat{\eta}^R - \bar{\eta}^R\| + \|\bar{\eta}^R - \eta^R\| \xrightarrow{a.s.} 0.$$

Notes: When adding a penalization $\frac{\lambda_n}{n}J(\beta)$, $\lambda_n \geq 0$, to the function $M_n(\beta)$ defined in New settings, the convergence of $\hat{\eta}^R$ to η^R is not straightforward. If $J(\beta) = \sum_{u \in S} \sum_{l_u} (\beta_{l_u}^u)^2$, the result is immediate. However, the ℓ_0 or the ℓ_1 penalization is of our interest. In this case, some strong assumptions of uniform convergence on the penalized function $M_n(\beta)$ would be necessary to get the consistency of $\hat{\beta}$ (to mimic the convergence on the M-estimators [\[VDV98\]](#)).

Part IV

Conclusions and Perspectives

Conclusions and method improvements

The ANOVA decomposition is a very attractive mathematical object for sensitivity analysis. For models with independent inputs, the Hoeffding decomposition is unanimously admitted in the sensitivity analysis community because it leads to the definition of a meaningful index to quantify the variability of a set of inputs on the output. However, for models with dependent incomes, there exist several decompositions motivated by different purposes. In this thesis, we focused on a specific ANOVA decomposition, inspired by the earlier work of Stone [Sto94] and Hooker [Hoo07]. Here, we were interested by a rigorous and relevant decomposition of the theoretical model function that depends on non-independent input variables. Moreover, the aim was to provide a clear and unambiguous generalized ANOVA decomposition that recovers the Hoeffding one when the inputs are independent. In particular, we have tried to shed some insight on the application of such decomposition for sensitivity analysis. We proposed a variance-based sensitivity index that generalizes the classical Sobol index to the case of a general form of the inputs distribution. Provided that the generalized decomposition is unique under the assumption that its components are hierarchically orthogonal, we deduced numerical methods able to identify these components. For both methods, our generalized sensitivity indices are deduced by empirical estimation of the covariance terms.

The first method consists in projecting the model output onto constrained spaces to obtain a functional linear system. The numerical resolution of these systems requires the estimation of conditional expectations in an iterative scheme. However, due to its computational cost, this method is only well-tailored for low-dimensional or pairwise dependent variables models. The second method is inspired by the meta-modelling approach. To identify each summand of the decomposition, we directly construct a meta-model adapted to the hierarchical orthogonality conditions. From an initial truncated orthonormal system (polynomials, splines, wavelets,...), we recursively construct systems subject to the orthogonal conditions. Thus, each component of the decomposition is well approximated by a linear combination of hierarchically orthogonal systems. The corresponding coefficients are then estimated by least-squares estimation.

However, this procedure generates a curse of dimensionality for two reasons. The first one is the number of ANOVA summands, which exponentially grows with the model dimension. Indeed, if p is the model dimension, the

ANOVA decomposition consists in 2^p components of increasing dimensions. The second reason is the degree of the initial orthonormal systems. It also grows very quickly when the number of components is getting large. As done in the second part of the contribution, the penalized regression is the choice we made to get a sparse selection of functions able to reconstitute the major part of the components' information. Using both the greedy algorithm to numerically solve the ℓ_0 penalty and the LARS method for the ℓ_1 penalty, we only provide convergence results for the usual least-squares estimation, i.e. without a penalization. The perspective would be to testify that a method used to solve the penalized regression provides a consistent estimator. The boosting algorithm [Sch90, BY10] is a greedy strategy that offers a triple advantage. The first one is that it is a very intuitive process resulting an easy implementation in practice. The boosting does variable selection and shrinkage, similar to the Lasso. Indeed, at each step, the selection is made in such a way that the correlation between the current residual and the element to be selected is maximized. At last, the boosting algorithm provides theoretical results on the consistency of its estimator [Buh06, CCAGV13]. To sum up, the boosting offers a good compromise between a greedy and a Lasso strategy in practice, and also offers consistency of the estimator. To explain the boosting strategy, we adopt the notation of Chapter 7, Part III. Remember that we want to solve the following regression problem,

$$Y = \eta^R(\mathbf{X}) + \varepsilon, \quad \eta^R(\mathbf{X}) = \sum_{u \in S} \sum_{\substack{\mathbf{l}_u \in \times_{i \in u} [1:L_i]}} \beta_{\mathbf{l}_u}^{u,0} \phi_{\mathbf{l}_u}^u(\mathbf{X}_u) \in H_u^{0,L}.$$

The aim is to recover the unknown coefficients $(\beta_{\mathbf{l}_u}^{u,0})_{\mathbf{l}_u, u}$ when the n -sample of observations $(y^s, \mathbf{x}^s)_{s=1, \dots, n}$ is observed in high-dimensional paradigm. We define a finite dictionary $\mathcal{D} = \{\phi_1^1, \dots, \phi_{l_1}^1, \dots, \phi_1^u, \dots, \phi_{l_u}^u, \dots\}$. Boosting algorithm generates successive approximations of η^R , using linear combination of elements of $\mathcal{D}$. The quantity $\hat{G}_k(\eta^R)$ denotes the approximation of η^R at step k of the algorithm. The boosting algorithm is described in Algorithm 7.

Following the work of Buhlmann [Buh06], clarified and improved by Champion *et al.* [CCAGV13], it can be shown that, under some smooth conditions of the elements of $\mathcal{D}$, and under the condition that the number of elements of $\mathcal{D}$ is of order $\exp n^\alpha$, $\alpha < 1$, that,

$$\mathbb{E} \left\| \eta^R - \hat{G}_{k_n}(\eta^R) \right\|_n \xrightarrow{\mathbb{P}} 0,$$

where $k_n = C \log(n)$ is the theoretical number of iterations, with C a constant term.

Algorithm 7: The $\mathbb{L}_2$ -boosting**Input:** Observations $(y^s, \mathbf{x}^s)_{s=1, \dots, n}$, $\gamma \in]0, 1]$ and $k_{up} \in \mathbb{N}^*$ **Output:** $\hat{G}_{k_{up}}(\eta^R)$ *Initialization:*

$$\hat{G}_0(\eta^R) = 0;$$

for $k = 1$ **to** k_{up} **do**1. Select $\phi_{\mathbf{l}_{u_k}}^{u_k} \in \mathcal{D}$ such that

$$|\langle Y - \hat{G}_{k-1}(\eta^R), \phi_{\mathbf{l}_{u_k}}^{u_k} \rangle_n| = \max_{\phi_{\mathbf{l}_u}^u \in \mathcal{D}} |\langle Y - \hat{G}_{k-1}(\eta^R), \phi_{\mathbf{l}_u}^u \rangle_n|$$

2. Compute the new approximation of η^R as

$$\hat{G}_k(\eta^R) = \hat{G}_{k-1}(\eta^R) + \gamma \langle Y - \hat{G}_{k-1}(\eta^R), \phi_{\mathbf{l}_{u_k}}^{u_k} \rangle_n \cdot \phi_{\mathbf{l}_{u_k}}^{u_k}$$

end

However, for our consideration, the functions $(\phi_{\mathbf{l}_u}^u)_{\mathbf{l}_u, u}$ are replaced by their empirical version $(\varphi_{\mathbf{l}_{u,n}}^u)_{\mathbf{l}_u, u}$. We actually want to estimate the coefficients $(\beta_{\mathbf{l}_u}^{u,0})_{u, \mathbf{l}_u}$ using the following approximation $\bar{\eta}^R$ of η^R ,

$$\bar{\eta}^R = \sum_{u \in S} \bar{\eta}_u^R(\mathbf{X}_u) = \sum_{u \in S} \sum_{\substack{\mathbf{l}_u \in \times \\ i \in u}} \beta_{\mathbf{l}_u}^{u,0} \varphi_{\mathbf{l}_{u,n}}^u(\mathbf{X}_u) \in G_{u,n} 0, L.$$

Thus, the new dictionary $\mathcal{D}_{\text{new}}$ would be

$$\mathcal{D}_{\text{new}} = \{\varphi_{\mathbf{l}_{1,n}}^1, \dots, \varphi_{\mathbf{l}_{1,n}}^1, \dots, \varphi_{\mathbf{l}_{1,n}}^u, \dots, \varphi_{\mathbf{l}_{u,n}}^u, \dots\}.$$

The modified version of Algorithm 7 consists in replacing the dictionary $\mathcal{D}$ by $\mathcal{D}_{\text{new}}$, i.e. the theoretical functions by their empirical version. The approximation $\hat{G}_{k_n}(\eta^R)$ is replaced by $\hat{G}_{k_n}(\bar{\eta}^R)$. It follows that, to show that the consistency of $\hat{G}_{k_n}(\bar{\eta}^R)$, we should understand the construction of the functions $(\varphi_{\mathbf{l}_{u,n}}^u)_{\mathbf{l}_u, u}$, and study the rate of convergence of $\|\varphi_{\mathbf{l}_{u,n}}^u - \phi_{\mathbf{l}_u}^u\|$ with respect to the size of the dictionary $\mathcal{D}$ (or, indifferently, the size of $\mathcal{D}_{\text{new}}$). This work is in progress.

Also, the perspective would be apply this method to real-world problems, to work on the physical interpretation of our generalized sensitivity indices. Indeed, unlike the Sobol index, our index can be negative. Anyone might wonder about the physical meaning of such result, and so we do. However, it would be difficult to consider a variance-based measure without any co-variance terms when dependence in the model is admitted.

Another outlook would be to work on the condition under which we construct our generalized sensitivity indices. As seen in Chapter 6 of Part III, Condition (C.2) can be translated with copula functions for bidimensional models. The perspective would be to find a weaker condition than (C.2), or, at least to get a large panel of sufficient condition of (C.2) in terms of copulas. Thus, our indices would cover a wider class of problems for which the dependence structure can be approximated by copulas.

These improvements would give an additional argument for using the measure we proposed in our contribution. However, it is hard to imagine an estimation method without having a large sample of observations. Indeed, the estimation technique is based on a regression problem, that requires a large sample of observations in practice to gain robustness. For many complex computer codes, that sometimes depend on a high number of input parameters, the use of regression procedure requires a large number of model evaluations, which is intractable for time expensive codes. Also, as we mentioned above, the number of ANOVA components implies a curse of dimensionality, that constitutes an additional issue for expensive codes. In practice, we assume that only the main effects and low-order interactions contain the major part of the model behaviour, but very few theoretical arguments confirm this assumption. Thus, the order truncation leads to an error of approximation that can be hardly controlled.

Toward new decompositions

To overcome the problem of expensive code in sensitivity analysis, a common practice, already mentioned in this manuscript, consists in approximating the complex code by a meta-model (also called surrogate model). As described in Part I, Li *et al.* [LRY⁺10] considers a ANOVA representation of the output (also called ANCOVA for ANalysis of COvariance and VAriance [Can12]), in which each component is approximated by cubic splines. Recently, Caniou *et al.* [Can12] come back on the ANCOVA, but focus on the use of polynomial chaos to surrogate the model function. Instead of considering an arbitrary choice of basis functions, Caniou *et al.* [Can12] expand the surrogate model as a sum of lower-order polynomial chaos with respect to the inputs marginal distributions. They then apply the ANCOVA on a sample of observations from the joint distribution of $\mathbf{X}$ to estimate the sensitivity indices. In one hand, this technique is a way to tackle the problem of expensive code by the use of a meta-model. On the other hand, it is a relevant way to consider the ANCOVA decomposition, because the decomposition recovers the polynomial chaos expansion in the particular case of independent inputs [Sud08]. However, we are still face to the need to estimate a maximal number of effects from the decomposition. Moreover, as exposed in Part I, the choice of

the polynomial chaos is inherent to the variables distribution, so considering the marginal distributions may mislead the final results in case of strong dependence.

To remedy to the various issues we just exposed, the perspective is to carry on working on the expansion of a surrogate model. More specifically, we are interested by the meta-modelling class of Gaussian processes (GPs) [Aro09, RW06]. That is, we consider that the prior knowledge about the model function $\eta(\mathbf{x})$ can be modelled by a GP $Z(\mathbf{x})$, where, by definition, $Z(\mathbf{x})$, $\mathbf{x} \in \mathbb{R}^p$ is a collection of gaussian random variables defined on the probability space $(\Omega_Z, \mathcal{A}_Z, P_Z)$ indexed by $\mathbb{R}^p$, with values in $\mathbb{R}$. $Z(\mathbf{x})$ is completely specified by its mean function $m(\mathbf{x}) = \mathbb{E}_Z[Z(\mathbf{x})]$, and its covariance kernel $k(\mathbf{x}, \tilde{\mathbf{x}}) = \text{Cov}_Z(Z(\mathbf{x}), Z(\tilde{\mathbf{x}}))$. Thus,

$$\eta(\mathbf{x}) \simeq Z(\mathbf{x}), \quad Z(\mathbf{x}) \sim \text{GP}(m(\mathbf{x}), k(\mathbf{x}, \tilde{\mathbf{x}})).$$

When $(y^l, \mathbf{x}^l)_{l=1, \dots, n}$ is a n -sample from the distribution of $(Y, \mathbf{X})$, the aim is to predict $Z(\mathbf{x})$ conditionally to the observations $\mathbf{y} := (y^1, \dots, y^n)$, i.e. we consider the conditional distribution

$$[Z(\mathbf{x})|Z_n = \mathbf{y}] = \text{GP}(\mu(\mathbf{x}), s^2(\mathbf{x}, \tilde{\mathbf{x}})), \quad (7.15)$$

where $Z_n = (Z(\mathbf{x}^1), \dots, Z(\mathbf{x}^n))$ is a Gaussian vector. The functions $\mu(\mathbf{x})$ and $s^2(\mathbf{x}, \tilde{\mathbf{x}})$ depend on the $m(\mathbf{x})$, and on $k(\mathbf{x}, \tilde{\mathbf{x}})$ computing at the observation points $(\mathbf{x}^1, \dots, \mathbf{x}^n)$. For more details on the exact shape of μ and s^2 , one can refer to [RW06, SWN03]. The function $\mu(\mathbf{x})$ is used as a predictor, whereas the conditional variance $s^2(\mathbf{x}, \tilde{\mathbf{x}})$ is the mean squared error of this predictor.

Thus, the conditional distribution $[Z(\mathbf{x})|Z_n = \mathbf{y}]$ offers an attractive expression that can be exploited in sensitivity analysis. Among an amount of contribution for independent input variables, one could cite Oakley & O'Hagan who consider the predictive mean of $[Z(\mathbf{x})|Z_n = \mathbf{y}]$ to construct variance-based sensitivity indices [OO04]. This work has been later clarified by Marrel *et al.* [MLR09].

Durrande *et al.* [DGRC13] were interested by a new class of ANOVA kernels, that consists in writing univariate kernel as follows

$$k(x_i, \tilde{x}_i) = 1 + k_0^i(x_i, \tilde{x}_i), \quad i \in [1 : p],$$

where the kernel k_0^i is chosen in such a way that

$$[Z(x_i)|Z_0 = 0] = \text{GP}(0, k_0^i(x_i, \tilde{x}_i)),$$

where $Z(x_i) \sim \text{GP}(0, k^i(x_i, \tilde{x}_i))$ can be any GP indexed by $\mathbb{R}$, and

$$Z_0 = \int_{\mathbb{R}} Z(x) p_{X_i} dx \sim N\left(0, \iint_{\mathbb{R}^2} k^i(x, \tilde{x}) p_{X_i}(x) p_{X_i}(\tilde{x}) dx d\tilde{x}\right)$$

By extension, the multivariate covariance kernel is a product of univariate kernel, i.e. $k(\mathbf{x}, \tilde{\mathbf{x}}) = \prod_{i=1}^p [1 + k_0^i(x_i, \tilde{x}_i)]$. Now, let $Z(\mathbf{x})$ be a GP with

$$Z(\mathbf{x}) \sim \text{GP}(0, k(\mathbf{x}, \tilde{\mathbf{x}})), \quad k(\mathbf{x}, \tilde{\mathbf{x}}) = \prod_{i=1}^p [1 + k_0^i(x_i, \tilde{x}_i)].$$

In this way, $Z(\mathbf{x})$ admits a covariance kernel that can be written as a sum of increasing dimension kernels centered with respect to the marginal distribution P_{X_i} . Indeed, by GP properties and by integration, it can be shown that $\int_{\mathbb{R}} k_0^i(x, \tilde{x}) dP_{X_i} = 0$, and

$$k(\mathbf{x}, \tilde{\mathbf{x}}) = \prod_{i=1}^p [1 + k_0^i(x_i, \tilde{x}_i)] = 1 + \sum_{u \in S} \prod_{i \in u} k_0^i(x_i, \tilde{x}_i).$$

From this new class of kernels, Durrande *et al.* are concerned by the mean prediction $\mu(\mathbf{x})$ of $Z(\mathbf{x})$ conditionally to the observations $\mathbf{y}$ given by (7.15), where, thanks to the specific form of the kernel $k(\mathbf{x}, \tilde{\mathbf{x}})$, $\mu(\mathbf{x})$ can be decomposed as follows,

$$\mu(\mathbf{x}) = \sum_{u \in S} \mu_u(\mathbf{x}_u).$$

Further, Durrande *et al.* defines sensitivity indices as

$$S_u^\mu = \frac{V[\mu_u(\mathbf{X}_u)]}{V[\mu(\mathbf{X})]}.$$

In our perspective, we exploit both the class of kernels given by Durrande *et al.* [DGRC13], and the decomposition of the polynomial chaos with respect to the marginal given by [Can12] to propose a new decomposition. Instead of decomposing the predictive mean $\mu(\mathbf{x})$ of a GP, the idea is to decompose the GP itself, resulting a GPs expansion as a sum of GPs indexed by increasing dimension vector spaces, that is

$$Z(\mathbf{x}) = Z_0 + \sum_{u \in S \setminus \{\emptyset\}} Z_u(\mathbf{x}_u).$$

The advantage in comparison with the predictive mean is that the GPs expansion allows for getting an explicit prediction error, characterized by the covariance kernel. Thus, we can have a control on the meta-modelling

error. Also, the use of GP does not require a large number of observations to be well approximated [SWN03, SWMW89]. We consider there is quite a lot of advantages in considering a GPs decomposition. Nevertheless, many research work should be conducted for both theoretical and practical aspects in order to get a meaningful measure of sensitivity well-fitted for models with dependent input variables.

Bibliography

Bibliography

- [ADM97] M. Avellaneda, G. Davis, and S. Mallat. Adaptive greedy approximations. *Constructive approximation*, 13:57–98, 1997.
- [AMH78] M.M. Ali, N.N. Mikhail, and M.S. Haq. A class of bivariate distributions including the bivariate logistic. *Journal of Multivariate Analysis*, 8(3):405–412, 1978.
- [Aro09] N. Aronszajn. Theory of reproducing kernels. *Transactions of the American Mathematical Society*, 68(3):337–404, 2009.
- [AS65] M. Abramowitz and I.A. Stegun. *Handbook of mathematical functions: with formulas, graphs, and mathematical tables*, volume 55. Dover publications, 1965.
- [Bar12] D. Barber. *Bayesian reasoning and machine learning*. Cambridge University Press, New York, 2012.
- [BCDD08] A. Barron, A. Cohen, W. Dahmen, and R. DeVore. Approximation and learning by greedy algorithms. *The Annals of Statistics*, 36(1):64–94, 2008.
- [BCT11] E. Borgonovo, W. Castaings, and S. Tarantola. Moment independent importance measures: New results and analytical test cases. *Risk Analysis*, 31(3):404–428, 2011.
- [BCT12] E. Borgonovo, W. Castaings, and S. Tarantola. Model emulation and moment-independent sensitivity analysis: An application to environmental modelling. *Environmental Modelling & Software*, 34:105–115, 2012.
- [Bed98] T. Bedford. Sensitivity indices for (tree-)dependent variables. In K. Chan, S. Tarantola, and F. Campolongo, editors, *SAMO98*, pages 17–20. Second International Symposium on Sensitivity Analysis of Model Output, 1998.
- [Ber09] D. Bertsimas. Optimization methods. Technical report, MIT, 2009. Disponible à <http://ocw.mit.edu/courses>.
- [Bla09] G. Blatman. *Adaptive sparse polynomial chaos expansions for uncertainty propagation and sensitivity analysis*. PhD thesis, Université BLAISE PASCAL - Clermont II, 2009.

- [Bor07] E. Borgonovo. A new uncertainty importance measure. *Reliability Engineering & System Safety*, 92:771–784, 2007.
- [Bou04] J. Boutahar. *Méthodes de réduction et de propagation d'incertitudes: application à un modèle de chimie-transport pour la modélisation et la simulation des impacts*. PhD thesis, Ecole des Ponts ParisTech, 2004.
- [Bre95] L. Breiman. Better subset regression using the nonnegative garrote. *Technometrics*, 37(4):373–384, 1995.
- [BT08] E. Borgonovo and S. Tarantola. Moment-independent and variance-based sensitivity analysis with correlations: an application to the stability of a chemical reactor. *Wiley interscience*, 40:687–698, 2008.
- [Buh06] P. Buhlmann. Boosting for high-dimensional linear models. *The Annals of Statistics*, 34(2):559–583, 2006.
- [BvdG11] P. Buhlmann and S. van de Geer. *Statistics for high-dimensional data*. Springer, Berlin, 2011.
- [BY10] P. Bühlmann and B. Yu. Boosting. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(1):69–74, 2010.
- [Can12] Y. Caniou. *Analyse de sensibilité globale pour les modèles imbriqués et multiéchelles*. PhD thesis, Université Blaise Pascal - Clermont II, 2012.
- [CB00] C. Couvreur and Y. Bresler. On the optimality of the backward greedy algorithm for the subset selection problem. *SIAM Journal on Matrix Analysis and Applications*, 21(3):797–808, 2000.
- [CCAGV13] M. Champion, C. Cierco-Ayrolles, S. Gadat, and M. Vignes. Sparse regression and support recovery with l2-boosting algorithm. 2013.
- [CGP12] G. Chastaing, F. Gamboa, and C. Prieur. Generalized hoeffding-sobol decomposition for dependent variables - Application to sensitivity analysis. *Electronic Journal of Statistics*, 6:2420–2448, 2012.
- [CGP13] G. Chastaing, F. Gamboa, and C. Prieur. Generalized sobol sensitivity indices for dependent variables: Numerical methods. Disponible à <http://arxiv.org/abs/1303.4372>, 2013.
- [CHT00] M. H. Chun, S. J. Han, and N. I. Tak. An uncertainty importance measure using a distance metric for the change in a cumulative distribution function. *Reliability Engineering & System Safety*, 70(3):313–321, 2000.
- [Cia98] P.G. Ciarlet. *Introduction à l'analyse numérique matricielle et à l'optimisation*. Dunod, 1998.

- [CIBN05] D.G. Cacuci, M. Ionescu-Bujor, and I.M. Navon. *Sensitivity and Uncertainty Analysis, Volume II: Applications to Large-Scale Systems*, volume 2. Chapman & Hall/CRC, 2005.
- [CLMM09] T. Crestaux, O. Le Maître, and J.M. Martinez. Polynomial chaos expansion for sensitivity analysis. *Reliability Engineering & System Safety*, 94(7):1161–1172, 2009.
- [CLS78] R.I. Cukier, H.B. Levine, and K.E. Shuler. Nonlinear sensitivity analysis of multiparameter model systems. *The Journal of Computational Physics*, 26:1–42, 1978.
- [CM47] R.H. Cameron and W.T. Martin. The orthogonal development of non-linear functionals in series of fourier-hermite functionals. *The Annals of Mathematics*, 48(2):385–392, 1947.
- [Coo97] R.M. Cooke. Markov and entropy properties of tree-and vine-dependent variables. In *Proceedings of the ASA Section of Bayesian Statistical Science*, volume 27, 1997.
- [CS75] J.H. Cukier, R.I. and Schaibly and K.E. Shuler. Study of the sensitivity of coupled reaction systems to uncertainties in rate coefficients. III. analysis of the approximations. *Journal of Chemical physics*, 63(3):431–440, 1975.
- [CS07] A. Charpentier and J. Segers. Lower tail dependence for archimedean copulas: characterizations and pitfalls. *Insurance: Mathematics and Economics*, 40(3):525–532, 2007.
- [CS09] J. Castle and N. Shephard. *The Methodology and Practice of Econometrics: A Festschrift in Honour of David F. Hendry*. OUP Oxford, 2009.
- [CS10] Y. Caniou and B. Sudret. Distribution-based global sensitivity analysis using polynomial chaos expansion. In *Procedia social and behavioral sciences 2*, pages 7625–7626. Sixth International Conference on Sensitivity Analysis of Model Output, 2010.
- [Cyb09] A.B. Cybakov. *Introduction to nonparametric estimation*. Springer, 2009.
- [DDFG01] A. Doucet, N. De Freitas, and N. Gordon. *Sequential Monte Carlo methods in practice*, volume 1. Springer, New York, 2001.
- [Dev86] L. Devroye. *Non-uniform random variate generation*, volume 4. Springer-Verlag, New York, 1986.
- [DeV98] R.A. DeVore. Nonlinear approximation. *Acta Numerica*, pages 51–150, 1998.
- [DGRC13] N. Durrande, D. Ginsbourger, O. Roustant, and L. Carraro. Reproducing kernels for spaces of zero mean functions. Application to sensitivity analysis. *Journal of Multivariate Analysis*, 115:57–67, 2013.

- [DP73] J. De Pillis. Gauss-seidel convergence for operators on hilbert space. *SIAM Journal of Numerical Analysis*, 10(1):112–122, 1973.
- [DP10] J. Dick and F. Pillichshammer. *Digital nets and sequences: discrepancy theory and quasi-Monte Carlo integration*. Cambridge University Press, Cambridge, 2010.
- [DR06] E. De Rocquigny. La maîtrise des incertitudes dans un contexte industriel-1ere partie: une approche méthodologique globale basée sue des exemples. *Journal de la Société Française de Statistique*, 147(2):33–71, 2006.
- [DS06] J.J. Dreesbeke and G. Saporta. *Approches non paramétriques en régression*. Technip, Paris, 2006.
- [DV07] S. Da Veiga. *Analyse d’incertitudes et de sensibilité- Application aux modèles de cinétique chimique*. PhD thesis, Université Toulouse III, 2007.
- [DVWG09] S. Da Veiga, F. Wahl, and F. Gamboa. Local polynomial estimation for sensitivity analysis on models with correlated inputs. *Technometrics*, 51(4):452–463, 2009.
- [EHJT04] B. Efron, T. Hastie, I. Johnstone, and R. Tibshirani. Least angle regression. *The Annals of Statistics*, 32(2):407–451, 2004.
- [EMS02] P. Embrechts, A. McNeil, and D. Straumann. Correlation and dependence in risk management: properties and pitfalls. *Risk management: value at risk and beyond*, pages 176–223, 2002.
- [ES81] B. Efron and C. Stein. The jackknife estimate of variance. *The Annals of Statistics*, 9(3):586–596, 1981.
- [FF93] I.E. Frank and J.H. Friedman. A statistical view of some chemometrics regression tools(with discussion). *Technometrics*, 35:109–148, 1993.
- [FFK02] H.-B. Fang, K.-T. Fang, and S. Kotz. The meta-elliptical distributions with given marginals. *Journal of Multivariate Analysis*, 82:1–16, 2002.
- [FG96] J. Fan and I. Gijbels. *Local polynomial modelling and its application*. Chapman & Hall, London, 1996.
- [Fis25] R.A. Fisher. *Statistical methods for Research Workers*. Oliver and Boyd, Edinburgh, 1925.
- [FKN90] K.-T. Fang, S. Kotz, and K.-W. Ng. *Symmetric Multivariate and Related Distributions -Monographs on Statistics and Applied Probability*. Chapman and Hall, London, 1990.
- [FS01] N.I. Fisher and P. Switzer. Graphical assessment of dependence. *The American Statistician*, 55(3):233–239, 2001.

- [Gal88] F. Galton. Co-relations and their measurement, chiefly from anthropometric data. In *Proceedings of the Royal Society of London*, volume 45, pages 135–145. The Royal Society, 1888.
- [GB03] C. Genest and J.C. Boies. Detecting dependence with kendall plots. *The American Statistician*, 57(4):275–284, 2003.
- [GF07] C. Genest and .C. Favre. Everything you always wanted to know about copula modeling but were afraid to ask. *Journal of Hydrologic Engineering*, 12(4):347–368, 2007.
- [GJK⁺13] F. Gamboa, A. Janon, T. Klein, A. Lagnoux-Renaudie, and C. Prieur. Statistical inference for sobol pick freeze monte carlo method. Disponible à <http://hal.inria.fr/hal-00804668>, 2013.
- [GR04] C. Genest and B. Rémillard. Test of independence and randomness based on the empirical copula process. *Test*, 13(2):335–369, 2004.
- [GS03] R.G. Ghanem and P.D. Spanos. *Stochastic finite elements: a spectral approach*. Dover publications, 2003.
- [Gu02] C. Gu. *Smoothing Spline ANOVA Models*. Springer, Berlin, 2002.
- [GvL96] G.H. Golub and C.F. van Loan. *Matrix computations*. The Johns Hopkins University Press, Baltimore, 1996.
- [HI86] S. Hora and R. Iman. A comparison of maximum/bounding and bayesian/monte carlo for fault tree uncertainty analysis. Technical Report SAND85-2839, Sandia Laboratories, 1986.
- [Hoe48] W. Hoeffding. A class of statistics with asymptotically normal distribution. *The annals of Mathematical Statistics*, 19(3):293–325, 1948.
- [Hol88] R.B. Holmes. On random correlation matrices. Technical Report 803, Massachusetts Institute of Technology, 1988.
- [Hoo07] G. Hooker. Generalized functional anova diagnostics for high-dimensional functions of dependent variables. *Journal of Computational and Graphical Statistics*, 16(3):709–732, 2007.
- [HS96] T. Homma and A. Saltelli. Importance measures in global sensitivity analysis of nonlinear models. *Reliability Engineering & System Safety*, 52(1):1–17, 1996.
- [HT84] T. Hastie and R. Tibshirani. *Generalized additive models*. Defense Technical Information Center, 1984.
- [HTF09] T. Hastie, R. Tibshirani, and J. Friedman. *The elements of statistical learning*. Springer, New York, 2009.

- [Hua98] J.Z. Huang. Projection estimation in multiple regression with application to functional anova models. *The annals of statistics*, 26(1):242–272, 1998.
- [IC82] R.L. Iman and W.J. Conover. A distribution-free approach to inducing rank correlation among input variables. *Communications in Statistics-Simulation and Computation*, 11(3):311–334, 1982.
- [IDZ80] R. L. Iman, J.M. Davenport, and D.K. Zeigler. Latin hypercube sampling (program user’s guide). Technical report, Sandia Labs., Albuquerque, NM (USA), 1980.
- [IH90] R.L. Iman and S.C. Hora. A robust measure of uncertainty importance for use in fault tree system analysis. *Risk Analysis*, 10(3):401–406, 1990.
- [Ioo11] B. Iooss. Revue sur l’analyse de sensibilité globale de modèles numériques. *Journal de la Société Française de Statistique*, 152(1):1–23, 2011.
- [Jan12] A. Janon. *Analyse de sensibilité et réduction de dimension- Application à l’océanographie*. PhD thesis, Université de Grenoble, 2012.
- [JDHR10] P. Jaworski, F. Durante, W.K. Härdle, and T. Rychlik, editors. *Copula theory and its applications*, volume 198. Springer, 2010. Proceedings of the workshop held in Warsaw, 25-26 September 2009.
- [Jen06] J.L. Jensen. Sur les fonctions convexes et les inégalités entre les valeurs moyennes. *Acta Mathematica*, 10(1):175–193, 1906.
- [JKLR⁺12] A. Janon, T. Klein, A. Lagnoux-Renaudie, M. Nodet, and C. Prieur. Asymptotic normality and efficiency of two sobol index estimators. Disponible à <http://hal.inria.fr/hal-00665048>, 2012.
- [JLD06] J. Jacques, C. Lavergne, and N. Devictor. Sensitivity analysis in presence of model uncertainty and correlated inputs. *Reliability Engineering & System Safety*, 91:1126–1134, 2006.
- [JNP11] A. Janon, M. Nodet, and C. Prieur. Uncertainties assessment in global sensitivity indices estimation from metamodels. To appear in International Journal for Uncertainty Quantification, disponible à <http://hal.inria.fr/inria-00567977/en>, 2011.
- [Kel87] C.T. Kelley. *Iterative methods for optimization*, volume 18. Society for Industrial and Applied Mathematics, 1987.
- [Ken38] M.G. Kendall. A new measure of rank correlation. *Biometrika*, 30:81–93, 1938.

- [KN06] L. Kuipers and H. Niederreiter. *Uniform distribution of sequences*. Dover publications, New york, 2006.
- [KR08] C. Kharoubi-Rakotomalala. *Les fonctions copules en finance*, volume 7. Publications de la Sorbonne, 2008.
- [Kre98] R. Kress. *Numerical analysis*. Springer, New York, 1998.
- [Kru58] W.H. Kruskal. Ordinal measures of association. *Journal of the American Statistical Association*, 53(284):814–861, 1958.
- [KTA12] S. Kucherenko, S. Tarantola, and P. Annoni. Estimation of global sensitivity indices for models with dependent variables. *Computer Physics Communications*, 183:937–946, 2012.
- [KW70] G. Kimeldorf and G. Wahba. Spline functions and stochastic processes. *Sankhya*, 32(2):173–180, 1970.
- [LBW95] W.S. Lee, P. Bartlett, and R.C. Williamson. On efficient agnostic learning of linear combinations of basis functions. In *Proceedings of the Eighth Annual Conference on Computational Learning Theory*, pages 369–376. ACM Press, 1995.
- [LC09] H. Liu and X. Chen. Nonparametric greedy algorithms for the sparse learning problem. In *proceedings in advances in neural information processing systems*, volume 22, 2009.
- [LCT07] D. Lewandowski, R.M. Cooke, and R.J.D. Tebbens. Sample-based estimation of correlation ratio with polynomial approximation. *ACM Transactions on Modeling and Computer Simulation*, 18(1):1–17, 2007.
- [LD09] R. Lebrun and A. Dutfoy. A generalization of the nataf transformation to distributions with elliptical copula. *Probabilistic Engineering Mechanics*, 24(2):172–178, 2009.
- [Leb13] R. Lebrun. *Contributions à la modélisation de la dépendance stochastique*. PhD thesis, Université Paris VII-Denis Diderot, 2013.
- [Leh66] E.L. Lehmann. Some concepts of dependence. *The Annals of Mathematical Statistics*, 35(5):1139–1153, 1966.
- [Lem09] C. Lemieux. *Monte Carlo and Quasi-Monte Carlo sampling*. Series in Statistics. Springer, 2009.
- [LMM11] M. Lamboni, D. Makowski, and H. Monod. Indices de sensibilité, sélection de paramètres et erreur quadratique de prédiction: des liaisons dangereuses ? *Journal de la Société Française de Statistique*, 152(1):26–48, 2011.
- [LR10a] G. Li and H. Rabitz. D-morph regression: application to modeling with unknown parameters more than observation data. *Journal of mathematical chemistry*, 48(4):1010–1035, 2010.

- [LR10b] G. Li and H. Rabitz. D-morph regression: application to modeling with unknown parameters more than observation data. *Journal of Mathematical Chemistry*, 48:1010–1035, 2010.
- [LR12] G. Li and H. Rabitz. General formulation of hdmr component functions with independent and correlated variables. *Journal of Mathematical Chemistry*, 50(1):99–130, 2012.
- [LRH⁺08] G. Li, H. Rabitz, J. Hu, Z. Chen, and Y. Ju. Regularized random-sampling high dimensional model representation. *Journal of Mathematical Chemistry*, 43(3):6022–6032, 2008.
- [LRR01] G. Li, C. Rosenthal, and H. Rabitz. High dimensional model representations. *Journal of Physical Chemistry A*, 105(33):7765–7777, 2001.
- [LRY⁺10] G. Li, H. Rabitz, P.E. Yelvington, O. Oluwole, F. Bacon, Kolb C.E., and J. Schoendorf. Global sensitivity analysis with independent and/or correlated inputs. *Journal of Physical Chemistry A*, 114:6022–6032, 2010.
- [Lue97] D.G. Luenberger. *Optimization by vector space methods*. Wiley, New York, 1997.
- [LZ06] Y. Lin and H.H. Zhang. Component selection and smoothing in multivariate nonparametric regression. *The Annals of Statistics*, 34(5):2272–2297, 2006.
- [McK97] M.D. McKay. Nonparametric variance-based methods of assessing uncertainty importance. *Reliability Engineering & System Safety*, 57(3):267–279, 1997.
- [Mey00] C.D. Meyer. *Matrix analysis and applied linear algebra*. SIAM, Philadelphia, 2000.
- [Mey12] C. Meynet. *Sélection de variables pour la classification non supervisée en grande dimension*. PhD thesis, Université Paris-Sud XI, 2012.
- [MILR09] A. Marrel, B. Iooss, B. Laurent, and O. Roustant. Calculations of sobol indices for the gaussian process metamodel. *Reliability Engineering & System Safety*, 94(3):742–751, 2009.
- [MNM06] H. Monod, C. Naud, and D. Makowski. Uncertainty and sensitivity analysis for crop models. In *Working with Dynamic Crop Models: Evaluation, Analysis, Parameterization and Applications*, chapter 3, pages 55–100. Elsevier, 2006.
- [Mor56] D. Morgenstern. Einfache beispiele zweidimensionaler verteilungen. *Mitteilungsblatt für Mathematische Statistik*, 8:234–253, 1956.
- [MRL12] V.A.A. Marchi, F.A.R. Rojas, and F. Louzada. The chi-plot and its asymptotic confidence interval for analyzing bivariate

- dependence: An application to the average intelligence and atheism rates across nations data. *Journal of Data Science*, 10:711–722, 2012.
- [MT12] T. Mara and S. Tarantola. Variance-based sensitivity analysis of computer models with dependent inputs. *Reliability Engineering & System Safety*, 107:115–121, 2012.
- [MZ93] S.G. Mallat and Z. Zhang. Matching pursuits with time-frequency dictionaries. *IEEE transactions on signal processing*, 41(12):3397–3415, 1993.
- [Nat62] A. Nataf. Détermination des distributions de probabilité dont les marges sont données. Technical Report 225, Compte rendus de l’Académie des Sciences, 1962.
- [Nel65] J.A. Nelder. The analysis of randomized experiments with orthogonal block structure. i. block structure and the null analysis of variance. In *Proceedings of the Royal Society of London*, volume 283 of A, pages 147–162, 1965.
- [Nel06] R.B. Nelsen. *An introduction to copulas*. Springer, New York, 2006.
- [Oak82] D. Oakes. A model for association in bivariate survival data. *Journal of the Royal Statistical Society. Series B (Methodological)*, 44(3):414–422, 1982.
- [OO04] J.E. Oakley and A. O’Hagan. Probabilistic sensitivity analysis of complex models: a bayesian approach. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 66(3):751–769, 2004.
- [OPT00] M.R. Osborne, B. Presnell, and B.A. Turlach. On the lasso and its dual. *Journal of Computational and Graphical Statistics*, 9(2):319–337, 2000.
- [Osb00] M.R. Osborne. A new approach to variable selection in least squares problems. *Journal of numerical analysis*, 20:389–404, 2000.
- [Owe94] A.B. Owen. Lattice sampling revisited: monte carlo variance of means over randomized orthogonal arrays. *The Annals of Statistics*, 22(2):930–945, 1994.
- [Owe05] A.B. Owen. Multidimensional variation for quasi-Monte Carlo. In Jianqing Fan and Gang Li, editors, *International Conference on Statistics in honour of Professor Kai-Tai Fang’s 65th birthday*, pages 49–74, 2005.
- [Par62] Emanuel Parzen. On estimation of a probability density function and mode. *The annals of mathematical statistics*, 33(3):1065–1076, 1962.

- [Pea96] K. Pearson. Mathematical contributions to the theory of evolution. On a form of spurious correlation which may arise when indices are used in the measurement of organs. In *Proceedings of the Royal Society of London*, volume 60, pages 489–498. The Royal Society, 1896.
- [Pou11] D.-B. Pougaza. *Utilisation de la notion de copule en tomographie*. PhD thesis, Université Paris Sud-Paris XI, 2011.
- [RW06] C.E. Rasmussen and C.K.I. Williams. *Gaussian processes for machine learning*. MIT Press, Cambridge, 2006.
- [Sap06] G. Saporta. *Probabilités, analyse des données et statistique*. Editions Technip, 2006.
- [Sch90] R.E. Schapire. The strength of weak learnability. *Machine learning*, 5(2):197–227, 1990.
- [Sch02] R. Schmidt. Tail dependence for elliptically contoured distributions. *Mathematical Methods of Operations Research*, 55(2):301–327, 2002.
- [Sch05] M. Schmidt. Least squares optimization with l_1 -norm regularization. Technical report, Ecole Normale Supérieure, 2005.
- [Sch09] E. Schumann. Generating correlated uniform variates. <http://comisef.wikidot.com/tutorial:correlateduniformvariates>, 2009.
- [SCS00] A. Saltelli, K. Chan, and E.M. Scott. *Sensitivity Analysis*. Wiley, West Sussex, 2000.
- [SG04] C. Soize and R. Ghanem. Physical systems with random uncertainties: chaos representations with arbitrary probability measure. *SIAM Journal on Scientific Computing*, 26(2):395–410, 2004.
- [Skl71] A. Sklar. Random variables, joint distribution functions, and copulas. *Kybernetika*, 9(3):449–460, 1971.
- [SM50] J. Sherman and W.J. Morrison. Adjustment of an inverse matrix corresponding to a change in one element of a given matrix. *The Annals of Mathematical Statistics*, 21(1):124–127, 1950.
- [Sob67] I.M. Sobol. On the distribution of points in a cube and the approximate evaluation of integrals. *Zhurnal Vychislitel'noi Matematiki i Matematicheskoi Fiziki*, 7(4):784–802, 1967.
- [Sob76] I.M. Sobol. Uniformly distributed sequences with an addition uniform property. *USSR Computational Mathematics and Mathematical Physics*, 16(5):236–242, 1976.
- [Sob93] I.M. Sobol. Sensitivity estimates for nonlinear mathematical models. *Mathematical Modeling and Computational Experiment*, 1(4):407–414, 1993.

- [Sob01] I.M. Sobol. Global sensitivity indices for nonlinear mathematical models and their monte carlo estimates. *Mathematics and Computers in Simulations*, 55:271–280, 2001.
- [SRA⁺08] A. Saltelli, M. Ratto, T. Andres, F. Campolongo, J. Cariboni, D. Gatelli, M. Saisana, and S. Tarantola. *Global sensitivity analysis: The primer*. Wiley-Interscience, West Sussex, 2008.
- [ST02] A. Saltelli and S. Tarantola. On the relative importance of input factors in mathematical models: Safety assessment for nuclear waste disposal. *Journal of the American Statistical Association*, 97:2811–2828, 2002.
- [STC99] A. Saltelli, S. Tarantola, and K.S. Chan. A quantitative model-independent method for global sensitivity analysis of model output. *Technometrics*, 41(1):39–56, 1999.
- [STCR04] A. Saltelli, S. Tarantola, F. Campolongo, and M. Ratto. *Sensitivity analysis in practice: a guide to assessing scientific models*. Wiley, West Sussex, 2004.
- [Sto94] C.J. Stone. The use of polynomial splines and their tensor products in multivariate function estimation. *The Annals of Statistics*, 22(1):118–171, 1994.
- [Sud08] B. Sudret. Global sensitivity analysis using polynomial chaos expansion. *Reliability engineering and system safety*, 93(7):964–979, 2008.
- [SWMW89] J. Sacks, W.J. Welch, T.J. Mitchell, and H.P. Wynn. Design and analysis of computer experiments. *Statistical science*, 4(4):409–423, 1989.
- [SWN03] T.J. Santner, B.J. Williams, and W.I. Notz. *The design and analysis of computer experiments*. Springer Verlag, New York, 2003.
- [Tem11] V. Temlyakov. *Greedy approximation*. Cambridge monographs on applied and computational mathematics. Cambridge University Press, Cambridge, 2011.
- [TGM06] S. Tarantola, D. Gatelli, and T. Mara. Random balance designs for the estimation of first order global sensitivity indices. *Reliability Engineering & System Safety*, 91(6):717–727, 2006.
- [Tib96] R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society*, 58(1):267–288, 1996.
- [Tib12] R.J. Tibshirani. The Lasso problem and uniqueness. Disponible à <http://arxiv.org/abs/1206.0313>, 2012.
- [Tis12] J.Y. Tissot. *Sur la décomposition ANOVA et l’estimation des indices de Sobol’*. Application à un modèle d’écosystème marin. PhD thesis, Université de Grenoble, 2012.

- [Tou11] S. Touzani. *Méthodes de surface de réponse basées sur la décomposition de la variance fonctionnelle et application à l'analyse de sensibilité*. PhD thesis, Université de Grenoble, 2011.
- [TP12a] J.Y. Tissot and C. Prieur. A bias correction method for the estimation of sensitivity indices based on random balance designs. *Reliability Engineering & System Safety*, 107:205–213, 2012.
- [TP12b] J.Y. Tissot and C. Prieur. Variance-based sensitivity analysis using harmonic analysis. Disponible à <http://hal.inria.fr/hal-00680725/en>, 2012.
- [TZ05] P.K. Trivedi and D.M. Zimmer. Copula modeling: an introduction for practitioners. *Foundations and Trends in Econometrics*, 1:1–111, 2005.
- [VDV98] A.W. Van Der Vaart. *Asymptotic Statistics*. Cambridge University Press, Cambridge, 1998.
- [VdW52] B.L. Van der Waerden. Order tests for the two-sample problem and their power. *Indagationes Mathematicae*, 14(253):458, 1952.
- [vM13] R. von Mises. Mechanik der festen körper im plastisch deformablen zustand. *Göttingen Nachrichten. Math. Phys.*, 1:582–592, 1913.
- [Wei80] S. Weisberg. *Applied linear regression*. Wiley, New York, 1980.
- [Wey16] Hermann Weyl. Über die gleichverteilung von zahlen mod. eins. *Mathematische Annalen*, 77(3):313–352, 1916.
- [WHH06] N.A. Weiss, P.T. Holmes, and M. Hardy. *A course in probability*. Pearson Addison Wesley, 2006.
- [Wie38] N. Wiener. The homogeneous chaos. *American Journal of Mathematics*, 60(4):897–936, 1938.
- [XG07] C. Xu and G. Gertner. Extending a global sensitivity analysis technique to models with correlated parameters. *Computational Statistics and Data Analysis*, 51:5579–5590, 2007.
- [XG08] C. Xu and G. Gertner. Uncertainty and sensitivity analysis for models with correlated parameters. *Reliability Engineering & System Safety*, 93:1563–1573, 2008.
- [Zha11] T. Zhang. Adaptive forward-backward algorithm for learning sparse representations. *IEEE transactions on information theory*, 57(7):4689–4708, 2011.
- [ZKS13] M.M. Zuniga, S. Kucherenko, and N. Shah. Metamodelling with independent and dependent inputs. *Computer Physics Communications*, 184(6):1570–1580, 2013.