



Analyse de sensibilité en fiabilité des structures

Paul Lemaitre

► To cite this version:

Paul Lemaitre. Analyse de sensibilité en fiabilité des structures. Mécanique des structures [physics.class-ph]. Université de Bordeaux, 2014. Français. NNT : 2014BORD0061 . tel-01143779

HAL Id: tel-01143779

<https://theses.hal.science/tel-01143779>

Submitted on 20 Apr 2015

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 PRÉSENTÉE
POUR OBTENIR LE GRADE DE
DOCTEUR DE
L'UNIVERSITÉ DE BORDEAUX

ÉCOLE DOCTORALE DE MATHÉMATIQUES ET INFORMATIQUE

SPÉCIALITÉ : Mathématiques appliquées

Par Paul Lemaître

**Analyse de sensibilité en fiabilité des
Structures**

Directeur de thèse : M. Pierre Del Moral

Soutenue le 18 mars 2014 devant la commission d'examen composée de :

M. Josselin Garnier	Professeur	Université Paris VII	Rapporteur
M. Bertrand Iooss	Chercheur Senior	EDF R& D	Co-encadrant
M. François Legland	Dir. de Recherche	INRIA Rennes	Rapporteur
M. Emmanuel Remy	Chercheur Expert	EDF R& D	Examineur
M. Jérôme Saracco	Professeur	Inst. Poly. de Bordeaux	xamineur

Titre : Analyse de sensibilité en fiabilité des structures

Résumé : Cette thèse porte sur l'analyse de sensibilité dans le contexte des études de fiabilité des structures. On considère un modèle numérique déterministe permettant de représenter des phénomènes physiques complexes.

L'étude de fiabilité a pour objectif d'estimer la probabilité de défaillance du matériel à partir du modèle numérique et des incertitudes inhérentes aux variables d'entrée de ce modèle. Dans ce type d'étude, il est intéressant de hiérarchiser l'influence des variables d'entrée et de déterminer celles qui influencent le plus la sortie, ce qu'on appelle l'analyse de sensibilité. Ce sujet fait l'objet de nombreux travaux scientifiques mais dans des domaines d'application différents de celui de la fiabilité. Ce travail de thèse a pour but de tester la pertinence des méthodes existantes d'analyse de sensibilité et, le cas échéant, de proposer des solutions originales plus performantes. Plus précisément, une étape bibliographique sur l'analyse de sensibilité puis sur l'estimation de faibles probabilités de défaillance est proposée. Cette étape soulève le besoin de développer des techniques adaptées. Deux méthodes de hiérarchisation de sources d'incertitudes sont explorées. La première est basée sur la construction de modèle de type classifieurs binaires (forêts aléatoires). La seconde est basée sur la distance, à chaque étape d'une méthode de type subset, entre les fonctions de répartition originelle et modifiée. Une méthodologie originale plus globale, basée sur la quantification de l'impact de perturbations des lois d'entrée sur la probabilité de défaillance est ensuite explorée. Les méthodes proposées sont ensuite appliquées sur le cas industriel CWNR, qui motive cette thèse.

Mots clés : Analyse de sensibilité ; Fiabilité; Incertitudes ; Expériences numériques; Perturbation des lois

Title : Reliability sensitivity analysis

Abstract : This thesis' subject is sensitivity analysis in a structural reliability context. The general framework is the study of a deterministic numerical model that allows to reproduce a complex physical phenomenon. The aim of a reliability study is to estimate the failure probability of the system from the numerical model and the uncertainties of the inputs. In this context, the quantification of the impact of the uncertainty of each input parameter on the output might be of interest. This step is called sensitivity analysis. Many scientific works deal with this topic but not in the reliability scope. This thesis' aim is to test existing sensitivity analysis methods, and to propose more efficient original methods. A bibliographical step on sensitivity analysis on one hand and on the estimation of small failure probabilities on the other hand is first proposed. This step raises the need to develop appropriate techniques. Two variables ranking methods are then explored. The first one proposes to make use of binary classifiers (random forests). The second one measures the departure, at each step of a subset method, between each input original density and the density given the subset reached. A more general and original methodology reflecting the impact of the input density modification on the failure probability is then explored.

The proposed methods are then applied on the CWNR case, which motivates this thesis.

Keywords : Sensitivity Analysis; Reliability; Uncertainties; Computer experiments; Input perturbations

Cette thèse est une grande aventure qui commença en décembre 2008, à Lyngby, Danemark. A cette époque, jeune étudiant Erasmus en quête d'un stage de 4^{ème} année, j'ai vu arriver dans ma boîte mail une proposition de stage en analyse d'incertitudes au CEA de Cadarache. Un ami, que je dénonce plus tard, m'a aidé à passer le cap du "j'ai pas le niveau" et m'a poussé à candidater. Bien m'en a pris, car ce stage fut le début d'une fructueuse collaboration. C'est le moment de remercier toutes les personnes sans qui la thèse n'aurait tout simplement pas été possible.

Que ce soit à Tunis, Cadarache, Chatou ou à Villard de Lans, Bertrand Iooss a su faire montre d'une grande humanité et de compétences scientifiques très poussées sous couvert d'une carapace punk. Je lui suis très reconnaissant de m'avoir transmis une bonne partie de ses connaissances, je l'espère pas uniquement statistiques.

Merci à Fabrice Gamboa pour m'avoir permis de finir ma thèse dans de bonnes conditions, je lui suis infiniment reconnaissant de m'avoir accepté comme cobureau ("room service") pour les 5 derniers mois.

Je remercie Pierre Del Moral d'avoir accepté de m'encadrer au titre de directeur de thèse. De la même façon, je sais gré à Josselin Garnier et François Le Gland d'avoir patiemment relu cette thèse et d'avoir apporté des remarques pertinentes me permettant d'améliorer le résultat. Merci aussi à Jérôme Saracco pour avoir accepté de faire partie du jury.

Une bonne partie de ma gratitude va à Aurélie Arnaud pour son encadrement orienté applications. De la même façon, je suis particulièrement reconnaissant envers Emmanuel "Manu" Remy, pour m'avoir supporté dans son couloir à des heures que le code du travail réprouve, pour ses connaissances sur la physique des réacteurs et pour ses (très) patientes relectures.

Ma reconnaissance va également à Agnès Lagnoux pour m'avoir encadré lors d'un court mais efficace séjour de recherche à l'UPS fin 2012.

C'est l'occasion pour moi de dire la gratitude que j'ai envers Didier Larrauri pour m'avoir autorisé à prolonger ma thèse afin de la finir dans de bonnes conditions. J'ai pu apprécier pendant cette thèse d'être entouré de collègues à la fois sympathiques et efficaces, pointus en statistiques et en applicatif. Mes hommages à tout MRI, T-55 et T-57 en tête. En particulier, je remercie Mathieu pour son aide sur la partie logicielle et Michael pour son savoir des arcanes du calcul numérique ainsi que Merlin pour avoir patiemment écouté mes élucubrations, notamment sur la partie "perturbation des paramètres" des DMBRSI. À mon "oncle de thèse" Nicolas, j'adresse mes plus sincères remerciements. Merci de m'avoir rassuré à plusieurs occasions sur l'avenir de cette thèse. Et bien sûr, salutations sportives pour mes deux collègues de salle de sport Popi et Fafa.

Cette thèse n'aurait pas vu le jour sans le fort support dans tous les moments difficiles que j'ai eu de la part de mes amis. En particulier je remercie le trio parisien formé par Jean-Phi, Thierno et Nessie, pour les bons moments et les autres. Des poutoux à Morgane et Diday. Il me faut également tirer mon chapeau à l'équipe toulousaine, Draco et Maria en tête. Je remercie également Bébert (il sait pourquoi). Big up à Raphaël, pour m'avoir efficacement conseillé sur tout l'interfaçage avec l'Université de Bordeaux. Mention spéciale à Petit Paul, Agathe, Alex et Mélo pour l'anglais. Par ailleurs, je suis redevable à ceux qui prendront la peine de poser un RTT pour aller jusqu'à Bordeaux le jour de ma soutenance, Fly, Quentin, Jeb et tonton Vip - hop dénoncé. Enfin, je ne peux conclure sans une pensée pour Mathieu.

Pour finir, mes sentiments vont vers ma famille, en particulier mes parents et mes deux sœurs qui furent pendant cette épreuve la définition même d'indéfectible.

Résumé étendu

Introduction

Analyse d'incertitudes et expériences numériques

On présente ici brièvement le cadre général de cette thèse : l'exploitation d'un modèle numérique. Un modèle est ici une représentation mathématique d'un phénomène physique et son traitement est effectué au travers d'un système de calcul.

Ce modèle possède des entrées et des sorties (ou réponses). Ici, toutes ces quantités seront considérées scalaires mais d'autres types pourraient être envisagés, modales par exemple. En fonction d'un jeu de données d'entrée, le code de calcul va produire un jeu de réponses après un certain temps de calcul. Le cadre des codes déterministes est utilisé : un même jeu d'entrée produira toujours le même jeu de sortie. Dans ce rapport, il sera parfois fait un abus de langage en assimilant le code au modèle, pour des raisons de lisibilité.

Une notion essentielle est la quantité d'intérêt. Il est en effet possible que ce ne soit pas une valeur de sortie qui intéresse l'expérimentateur, mais plutôt une plage de valeurs ou une quantité définie à partir des sorties. Il est donc primordial avant toute étude de définir quelle est la quantité d'intérêt.

L'analyse de sensibilité est définie par Saltelli et al. [89] comme l'étude de la façon dont l'incertitude sur une quantité de sortie du modèle peut être attribuée aux différentes sources d'incertitudes dans les variables d'entrée.

L'analyse de sensibilité d'un modèle numérique peut servir à déterminer les variables d'entrée qui contribuent le plus à un certain comportement d'une sortie, déterminer celles sans influence ou celles qui vont interagir à travers le modèle. Le but peut être de comprendre le modèle, de le simplifier, ou encore de prioriser le recueil de données pour mieux modéliser une variable d'entrée. Une approche récente est l'approche dite globale. L'ensemble du domaine de variation des variables d'entrée est alors étudié. La plupart des techniques sont développées dans une approche indépendante du modèle ("model free"), c'est-à-dire sans émettre d'hypothèses sur le comportement du modèle comme par exemple la linéarité ou la monotonie.

Fiabilité des structures

On cherche à répondre au problème industriel de savoir si une structure ou un composant peut résister à des contraintes qui lui sont appliquées. L'approche basée sur des essais et mesures est possible, mais peut s'avérer difficile pour des raisons de coûts ou de risques. Parfois, l'expérimentation est impossible. Des modèles numériques sont alors utilisés comme représentation approchée de la réalité incluant certains mécanismes (comme par exemple ceux de la dégradation, de la propagation des fissures...).

Afin d'exploiter complètement le modèle, les incertitudes sur les paramètres d'entrées du code (essentiellement des grandeurs physiques) sont modélisées par des variables aléatoires. Le modèle

représente donc la structure, dotée d'une certaine résistance, et l'environnement, qui engendre une sollicitation. Le calcul pour un jeu d'entrées fixées permet d'obtenir un critère de défaillance qui amène à une réponse binaire : la structure est défaillante pour ces entrées ou non défaillante.

Le fait d'inclure les incertitudes comme des variables aléatoires permet de modéliser le risque comme une probabilité de défaillance. Cette approche est plus fine qu'une approche déterministe où les grandeurs sont fixées à des valeurs nominales.

Soit $\mathbf{X} = (X_1, \dots, X_d)$ le vecteur aléatoire d -dimensionnel (dont la densité $f_{\mathbf{X}}$ est connue) des variables d'entrée (scalaires) du modèle numérique. On s'intéresse à ce que la valeur scalaire $Y \in \mathbb{R}$ renvoyée par la fonction de défaillance G du modèle (ou fonction d'état-limite du modèle) soit plus faible qu'un certain seuil k (usuellement 0) : c'est le critère de défaillance. La structure est défaillante pour un jeu d'entrée $\mathbf{x}$ si $y = G(\mathbf{x}) \leq k$ (où $\mathbf{x} = (x_1, \dots, x_d) \in \mathbb{R}^d$ est une réalisation de $\mathbf{X}$ et k un seuil usuellement fixé à 0). L'ensemble de l'espace sur lequel cet événement se produit est appelé domaine de défaillance D_f . La surface définie par $\{\mathbf{x} \in \mathbb{R}^d, G(\mathbf{x}) = k\}$ est dite surface d'état-limite. La probabilité que l'évènement se produise est notée P_f , probabilité de défaillance. On a :

$$\begin{aligned} P_f &= \mathbb{P}(G(\mathbf{X}) \leq k) \\ &= \int_{D_f} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \\ &= \int_{\mathbb{R}^d} 1_{G(\mathbf{x}) \leq k} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \\ &= \mathbb{E}[1_{G(\mathbf{X}) \leq k}] \end{aligned}$$

La complexité des modèles et le possible grand nombre de variables d'entrée fait que, dans le cas général, on ne peut pas calculer la valeur exacte de la probabilité de défaillance. On peut cependant estimer cette quantité (qui est une espérance mathématique) à l'aide de diverses méthodes numériques. La base de la fiabilité des structures est de fournir une estimation de P_f et une incertitude autour de cette estimation. Cette estimation permet ensuite de répondre à la question initiale de la résistance de la structure.

Objectifs de la thèse

Le but de cette thèse est le développement de techniques d'analyse de sensibilité quand la quantité d'intérêt est une probabilité de dépassement de seuil (ce qui équivaut à une probabilité de défaillance dans le contexte de la fiabilité des structures). Les contraintes du code CWNR qui a motivé le travail de thèse doivent être prises en compte. La probabilité de défaillance dans le cas le moins pénalisant (7 variables) a un ordre de grandeur attendu de 10^{-5} . Si possible, les méthodes développées doivent être en relation avec l'estimation de P_f et doivent produire une estimation de l'erreur faite lors de l'estimation des indices de sensibilité et de P_f .

Organisation de la thèse

La thèse est divisée en quatre chapitres.

Le premier chapitre est une revue des stratégies existantes pour estimer des probabilités de défaillance et des techniques d'analyse de sensibilité.

Le second chapitre est consacré à la définition de mesures de sensibilité avec pour but la production d'un classement de variables (variable ranking).

Le troisième chapitre présente une méthode originale pour estimer l'importance de chacune des variables d'entrée sur une probabilité de défaillance. Cette méthode se concentre sur l'impact d'une modification de densité d'entrée sur la probabilité de défaillance produite en sortie.

Le quatrième chapitre présente une application des méthodes étudiées sur le cas CWNR, cas réel qui a motivé la thèse.

Méthodes de classement de variables

Le second chapitre présente deux méthodes permettant de classer les variables d'entrée en fonction de leur influence sur la sortie (binaire). De plus, ces méthodes sont des sous-produits de l'estimation de la probabilité de défaillance P_f .

En effet la première technique propose de faire usage de mesures dérivées de l'ajustement de forêts aléatoires sur un échantillon de type Monte-Carlo. Un rappel sur les arbres binaires puis sur les forêts aléatoires est proposé, puis l'étude de deux indices (Gini Importance et Mean Decrease Accuracy) mesurant l'importance des variables sur la quantité d'intérêt binaire est proposé.

La seconde technique mesure l'écart, à chaque étape d'une méthode de type subset simulation, entre les densités d'entrée et les densités sachant que le sous-ensemble est atteint.

La définition informelle est la suivante : l'indice de sensibilité est défini pour la variable i et l'étape du subset k comme la distance entre la fonction de répartition (f.d.r.) empirique et la f.d.r. théorique de la variable. Considérant M étapes de subset avec $k = 1 \dots M$; et en notant :

$$F_{n,i}^k = F_i(x|A_k),$$

la f.d.r. empirique de la $i^{\text{ème}}$ variable sachant que le seuil A_k a été dépassé. L'indice proposé s'écrit comme suit :

$$\delta_i^{SS}(A_k) = d(F_{n,i}^k, F_i),$$

où F_i est la f.d.r. de la $i^{\text{ème}}$ variable, et d est une distance. Une variable influente aura un grand écart en f.d.r. alors qu'une variable non-influente aura un faible écart en f.d.r., donc un faible indice. Des travaux sont menés sur le choix de la distance d en fonction du besoin de l'analyste.

Ces deux méthodes peuvent donc être vues comme des sous-produits de techniques d'estimation de la probabilité de défaillance.

Méthode basée sur une perturbation des densités (DMBRSI)

Dans le troisième chapitre, de nouveaux indices de sensibilité pour la fiabilité sont proposés. Cet indice de sensibilité est basé sur une modification des densités et est adapté aux probabilités de défaillance. Une méthode pour estimer de tels indices est proposée.

Ces indices reflètent l'impact d'une modification d'une densité d'entrée sur la probabilité de défaillance P_f . Ils sont indépendants de la perturbation dans le sens où l'utilisateur peut choisir la perturbation adaptée à son problème.

Pour des raisons de simplicité, un schéma d'échantillonnage Monte-Carlo classique est considéré par la suite, bien que le processus d'estimation a été étendu aux méthodes subset et tirages d'importance. Les indices de sensibilité peuvent être estimés en utilisant seulement le jeu de simulations déjà utilisé pour estimer la probabilité de défaillance P_f . Ceci limite le nombre d'appels au code de calcul, comme mentionné dans les contraintes du cas industriel CWNR.

Le chapitre est organisé de la façon suivante : en premier lieu, les indices et leurs propriétés théoriques sont présentées ainsi qu'une méthode d'estimation. En second lieu, plusieurs méthodes

de perturbation des densités sont présentées. Ces modifications peuvent être classées en deux grandes familles : minimisation de Kullback-Leibler et perturbation des paramètres. Le comportement des indices proposés est testé sur des cas tests, puis les avantages et problèmes restants sont finalement discutés.

Le chapitre 3 est une version étendue du papier par Lemaître et coauteurs [63].

Indice DMBRSI

Soit une entrée unidimensionnelle X_i de densité f_i , on appelle $X_{i\delta} \sim f_{i\delta}$ l'entrée perturbée correspondante.

La probabilité de défaillance modifiée devient :

$$P_{i\delta} = \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(\mathbf{x}) d\mathbf{x}$$

où x_i est la $i^{\text{ème}}$ composante du vecteur $\mathbf{x}$.

L'indice DMBRSI a la forme suivante.

Définition On définit les indices de sensibilité basés sur une modification des lois (Density Modification Based Reliability Sensitivity Indices - DMBRSI) comme la quantité $S_{i\delta}$:

$$S_{i\delta} = \left[\frac{P_{i\delta}}{P_f} - 1 \right] 1_{\{P_{i\delta} \geq P_f\}} + \left[1 - \frac{P_f}{P_{i\delta}} \right] 1_{\{P_{i\delta} < P_f\}} = \frac{P_{i\delta} - P_f}{P_f \cdot 1_{\{P_{i\delta} \geq P_f\}} + P_{i\delta} \cdot 1_{\{P_{i\delta} < P_f\}}}.$$

Estimation

Un estimateur $\hat{P}_N$ de P_f peut être calculé en utilisant un plan d'expérience de N points. Par la suite, N est considéré comme étant assez grand pour que le contexte de la théorie asymptotique s'applique. Par ailleurs, un échantillonnage de type Monte-Carlo standard est utilisé pour simplifier les calculs. On écrit alors

$$\hat{P}_N = \frac{1}{N} \sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}}$$

où $\mathbf{x}^1, \dots, \mathbf{x}^N$ sont des réalisations indépendantes de X . La loi forte des grands nombres et le théorème limite centrale (TLC) assurent que pour presque toutes les réalisations, $\hat{P}_N \xrightarrow[N \rightarrow \infty]{} P_f$ et

$$\sqrt{\frac{N}{P_f(1-P_f)}}(\hat{P}_N - P_f) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, 1).$$

Le cadre Monte-Carlo permet d'estimer $P_{i\delta}$ de façon consistante sans nouvel appel au code de calcul G , grâce à une technique de tirage d'importance "inverse" (reverse importance sampling):

$$\hat{P}_{i\delta N} = \frac{1}{N} \sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)}.$$

Ceci est très intéressant quand le code de calcul G est coûteux en temps de calcul (Beckman and McKey, Hesterberg [8, 45]).

Dans la thèse, les propriétés asymptotiques des estimateurs de P_f et $S_{i\delta}$ sont étudiées.

Stratégies de perturbation

La Section 3.3 propose plusieurs méthodes de perturbations. On insiste sur le fait que les DMBRSI et les techniques d'estimation présentées restent valides pour toute perturbation tant que des contraintes sur le support sont respectées. Ici on se focalise sur deux familles de méthodes. Dans la première, la densité perturbée est celle minimisant la divergence de Kullback-Leibler sous des contraintes fixées par l'utilisateur. Plusieurs contraintes sont proposées (perturbation de la moyenne, de la variance et des quantiles). L'usage de la seconde méthode est conseillé quand l'utilisateur veut tester la sensibilité de P_f aux paramètres des distributions. Chaque section est introduite par un exemple jouet.

Cette section illustre la capacité des DMBRSI à traiter des objectifs d'analyse de sensibilité différents. L'utilisateur est invité à proposer de nouvelles perturbations qui répondraient à ses objectifs.

Application au cas CWNR

Le quatrième chapitre présente l'application des méthodes développées au cas CWNR. Ce cas est présenté dans l'organisation de la thèse, page 24. On rappelle que ce modèle de type "boîte-noire" constitue la motivation initiale de ce travail.

Pour estimer P_f , la méthode FORM (voir Section 1.2.2.2) et un Monte-Carlo naïf (voir Section 1.2.1.1) ont été utilisées. Les résultats produits par la méthode Monte-Carlo sont considérés comme étant la référence dans ce chapitre.

La partie analyse de sensibilité est consacrée à la mise en œuvre de trois méthodes : premièrement, les facteurs d'importance FORM (voir Section 1.3.2.2). Ensuite, des forêts aléatoires (voir Section 2.2) sont construites sur l'échantillon Monte-Carlo et des mesures de sensibilité sont dérivées. Pour finir, les DMBRSI (voir Chapitre 3) sont utilisés. Plusieurs perturbations (moyenne, quantile et paramètres) sont testées.

Ce chapitre est divisé en trois sections principales, se concentrant chacune sur des cas de dimension croissante (3, 5 et 7 variables probabilisées), où plus la dimension est petite, plus le cas est pénalisant.

Les conclusions de ce chapitre sont les suivantes :

- en ce qui concerne la partie estimation de P_f , la méthode de Monte-Carlo reste la référence sur un code industriel. Le désavantage majeur est bien entendu le temps de calcul nécessaire.
- En ce qui concerne la partie analyse de sensibilité, les forêts aléatoires produisent des résultats contestables, car les modèles ajustés sont de mauvaise qualité. La méthode est donc peu concluante pour l'instant.
- Les DMBRSI semblent une méthode adaptée pour effectuer une analyse de sensibilité sur une probabilité de défaillance. Plusieurs ajustements et configurations ont été testées.

Axes de recherches futures

Les méthodes présentées dans le Chapitre 2 peuvent être améliorées. Plus spécifiquement, il y a un besoin d'améliorer les classifieurs binaires (forêts aléatoires). Les indices MDA couplés à la subset simulation doivent être implémentés. Une autre perspective d'amélioration, en utilisant les indices $\delta_i^{SS}(A_k)$, est de mener un travail incluant la théorie des copules.

Les DMBRSI introduits dans le Chapitre 3 présentent eux aussi plusieurs perspectives d'amélioration. La grande partie des travaux sera consacrée à l'amélioration des indices $S_{i\delta}$ en termes de réduction de variance et d'appels au code de calcul. Le couplage des estimateurs avec la subset simulation doit aussi être perfectionné. Une perturbation basée sur l'entropie pourrait également être proposée, mais des calculs plus poussés doivent être menés pour obtenir une solution du problème de minimisation de la divergence de Kullback-Leibler. Un autre axe serait de changer la métrique/divergence. Par ailleurs, une autre idée pourrait être la prise en compte des dépendances entre variables et de perturber cette dépendance entre marginales *via* la théorie des copules.

Des perspectives plus larges sont à considérer, en particulier l'utilisation de méthodes séquentielles couplées avec des méta-modèles (Bect *et al.* [9]) est à étudier.

Récemment, Fort *et al.* [35] ont introduit de nouveaux indices de sensibilité pouvant être considérés comme une généralisation des indices de Sobol'. La notion de fonction de contraste adaptée au besoin est introduite. Cet indice doit être testé et comparé avec les DMBRSI dans un travail futur.

Contents

Résumé étendu	5
Contents	11
List of Figures	14
List of Tables	18
Context, objectives and outline	21
On numerical simulation	21
Uncertainty quantification and sensitivity analysis	21
Structural reliability	23
Context: component within nuclear reactor (CWNr)	24
Objectives	24
Outline	25
1 State of the art for reliability and sensitivity analysis	27
1.1 Introduction	27
1.2 State of the art: reliability and failure probability estimation techniques	27
1.2.1 Monte-Carlo methods	28
1.2.2 Structural reliability methods	33
1.2.3 Subset simulation	36
1.3 Sensitivity analysis (SA)	39
1.3.1 Global sensitivity analysis	39
1.3.2 Reliability based sensitivity analysis	44
1.4 Functional decomposition of variance for reliability	45
1.4.1 First applications	45
1.4.2 Computational methods	47
1.4.3 Reliability test cases	51
1.4.4 Reducing the number of function calls: use of QMC methods	56
1.4.5 Reducing the number of function calls : use of importance sampling methods	57
1.4.6 Local polynomial estimation for first-order Sobol' indices in a reliability con-	
text	58
1.4.7 Conclusion on Sobol' indices for reliability	66
1.5 Moment independent measures for reliability	67
1.5.1 Application in the reliability case	67
1.5.2 Crude MC estimation of δ_i	67
1.5.3 Use of quadrature techniques	68

1.5.4	Use of subset sampling techniques	68
1.5.5	Hyperplane 6410 test case	68
1.5.6	Conclusion	69
1.6	Synthesis	69
1.7	Sensitivity analysis for failure probabilities (FPs)	70
2	Variable ranking in the reliability context	73
2.1	Introduction	73
2.2	Using classification trees and random forests in SA	73
2.2.1	State of the art for classification trees	73
2.2.2	Stabilisation methods	75
2.2.3	Variable importance - Sensitivity analysis	78
2.2.4	Applications	80
2.2.5	Discussion	85
2.3	Using input cumulative distribution function departure as a measure of importance	88
2.3.1	Introduction and reminders	88
2.3.2	Distances	89
2.3.3	Applications	90
2.3.4	Conclusion	104
2.4	Synthesis	104
3	Density Modification Based Reliability Sensitivity Indices	107
3.1	Introduction and overview	107
3.2	The indices: definition, properties and estimation	108
3.2.1	Definition	108
3.2.2	Properties	108
3.2.3	Estimation	108
3.2.4	Framework	110
3.3	Methodologies of input perturbation	112
3.3.1	Kullback-Leibler minimization	112
3.3.2	Parameters perturbation	119
3.3.3	Choice of the perturbation given the objectives	125
3.4	Numerical experiments	126
3.4.1	Testing methodology	126
3.4.2	Hyperplane 6410 test case	126
3.4.3	Hyperplane 11111 test case	133
3.4.4	Hyperplane with 15 variables test case	137
3.4.5	Hyperplane with same importance and different spreads test case	141
3.4.6	Tresholded Ishigami function	146
3.4.7	Flood test case	155
3.5	Improving the DMBRSI estimation	161
3.5.1	Coupling DMBRSI with importance sampling	161
3.5.2	Coupling DMBRSI with subset simulation	163
3.6	Discussion and conclusion	165
3.6.1	Conclusion on the DMBRSI method	165
3.6.2	Equivalent perturbation	165
3.6.3	Support perturbation	166
3.6.4	Further work	166

3.6.5	Acknowledgements	166
4	Application to the CWNR case	167
4.1	Introduction	167
4.2	Three variables case	167
4.2.1	Estimating P_f	168
4.2.2	Sensitivity Analysis	168
4.3	Five variables case	175
4.3.1	Estimating P_f	175
4.3.2	Sensitivity Analysis	175
4.4	Seven variables case	182
4.4.1	Estimating P_f	183
4.4.2	Sensitivity Analysis	183
4.5	Conclusion	189
	Conclusion	191
	Bibliography	195
A	Distributions formulas	203
B	Test cases	205
B.1	Hyperplane test case	205
B.2	Tresholded Ishigami function	206
B.3	Flood case	207
C	Isoprobabilistic transformations	209
C.1	Presentation of the copulas	209
C.2	Objectives, Rosenblatt transformation	210
D	Appendices for Chapter 3	211
D.1	Proofs of asymptotic properties	211
	Proof of Lemma 3.2.1	211
	Proof of Proposition 3.2.1	211
D.2	Computation of Lagrange multipliers	212
D.3	Proofs of the NEF properties	212
D.4	Numerical trick to work with truncated distribution	214

List of Figures

1	Uncertainty study reference framework	22
1.1	Space filling comparison: Sobol's sequence (left) and uniform random sampling (right).	29
1.2	2-dimensional illustration of directional sampling	32
1.3	Illustration of FORM/SORM	34
1.4	Conditional expectations for 2 variables	46
1.5	Boxplots of the estimated first order Sobol' indices with the Sobol' method	53
1.6	Boxplots of the estimated first order Sobol' indices with the Saltelli method	54
1.7	Comparison of first order and total indices, MC (left) and importance sampling (right), with 10^4 points for the hyperplane 6410 test case	59
1.8	Comparison of first order and total indices, MC (left) and importance sampling (right), with 10^3 points for the hyperplane 6410 test case	60
1.9	Boxplot of the estimated FOSIFD for the 6410 hyperplane case	62
1.10	Boxplot of the estimated FOSIFD for the 11111 hyperplane case	62
1.11	Boxplot of the estimated FOSIFD for the 15 variables hyperplane case	63
1.12	Boxplot of the estimated FOSIFD for the same importance different spread hyperplane case	64
1.13	Boxplot of the estimated FOSIFD for thresholded Ishigami case	64
1.14	Boxplot of the estimated FOSIFD for the flood case	65
1.15	Example surface	66
2.1	Binary tree	76
2.2	Boxplots of MDA indices (left) and GI indices (right) for the hyperplane 6410 test case .	81
2.3	Boxplots of MDA indices (left) and GI indices (right) for the hyperplane 11111 test case	82
2.4	Boxplots of MDA indices for the hyperplane 15 variables test case	82
2.5	Boxplots of GI indices for the hyperplane 15 variables test case	83
2.6	Boxplots of MDA indices (left) and GI indices (right) for the hyperplane different spreads test case	84
2.7	Boxplots of MDA indices (left) and GI indices (right) for the thresholded Ishigami test case	84
2.8	Boxplots of MDA indices (left) and GI indices (right) for the flood test case	85
2.9	Several c.d.f.	91
2.10	Hyperplane 6410 test case, Kolmogorov distance	92
2.11	Hyperplane 6410 test case, Cramer-Von Mises distance	92
2.12	Hyperplane 6410 test case, Anderson-Darling distance	93
2.13	Hyperplane 11111 test case, Kolmogorov distance	94
2.14	Hyperplane 11111 test case, Cramer-Von Mises distance	94
2.15	Hyperplane 11111 test case, Anderson-Darling distance	95

2.16	Hyperplane 15 variables test case, Kolmogorov distance	96
2.17	Hyperplane 15 variables test case, Cramer-Von Mises distance	96
2.18	Hyperplane 15 variables test case, Anderson-Darling distance	97
2.19	Hyperplane different spread test case, Kolmogorov distance	98
2.20	Hyperplane different spread test case, Cramer-Von Mises distance	98
2.21	Hyperplane different spread test case, Anderson-Darling distance	99
2.22	Thresholded Ishigami test case, Kolmogorov distance	100
2.23	Thresholded Ishigami test case, Cramer-Von Mises distance	100
2.24	Thresholded Ishigami test case, Anderson-Darling distance	101
2.25	Flood test case, Kolmogorov distance	102
2.26	Flood test case, Cramer-Von Mises distance	103
2.27	Flood test case, Anderson-Darling distance	103
3.1	General DMBRSI framework	111
3.2	The original density of mean 0 (full line) and several candidates densities of mean 2 . .	113
3.3	Mean shifting (left) and variance shifting (right) for Gaussian (upper) and Uniform (lower) distributions. The original distribution is plotted in solid line, the perturbed one is plotted in dashed line.	116
3.4	Standard Gaussian and perturbed density: quantile increase (left) and quantile decrease (right)	118
3.5	Uniform, Triangle and Truncated Gumbel pdf: quantile increase	119
3.6	Original and perturbed Weibulls pdfs	120
3.7	DMBRSI with parameters perturbations	122
3.8	Specific DMBRSI framework for parameters perturbations	124
3.9	Estimated indices $\widehat{S}_{i\delta}$ for the 6410 hyperplane function with a mean shifting	128
3.10	Estimated indices $\widehat{S}_{i,V_f}$ for hyperplane function with a variance shifting	128
3.11	5 th percentile perturbation on the hyperplane 6410 test case	129
3.12	1 st quartile perturbation on the hyperplane 6410 test case	130
3.13	Median perturbation on the hyperplane 6410 test case	130
3.14	3 rd quartile perturbation on the hyperplane 6410 test case	131
3.15	95 th percentile perturbation on the hyperplane 6410 test case	131
3.16	Parameters perturbation on the hyperplane 6410 test case. Dots are for means, triangle for the standard deviations. Green corresponds to X_1 , black to X_2 , red to X_3 and blue to X_4	132
3.17	Estimated indices $\widehat{S}_{i\delta}$ for the 11111 hyperplane function with a mean shifting	135
3.18	Estimated indices $\widehat{S}_{i\delta}$ for the 11111 hyperplane function with a variance shifting	135
3.19	Median perturbation on the hyperplane 11111 test case	136
3.20	Parameters perturbation on the hyperplane 11111 test case. Dots are for means, triangle for the standard deviations. A different color is used for each variable.	137
3.21	Estimated indices $\widehat{S}_{i\delta}$ for the 15 variables hyperplane function with a mean shifting . . .	139
3.22	Estimated indices $\widehat{S}_{i\delta}$ for the 15 variables hyperplane function with a variance shifting .	140
3.23	Median perturbation on the hyperplane with 15 variables test case	141
3.24	Parameters perturbation on the 15 variables hyperplane test case. Dots are for means, triangle for the standard deviations. Black is for the first group of influence, red is for the second and blue for the third.	142
3.25	Estimated indices $\widehat{S}_{i\delta}$ for the hyperplane with different spreads case with a mean shifting	143
3.26	Estimated indices $\widehat{S}_{i\delta}$ for the hyperplane with different spreads case with a median shifting	144

3.27	Parameters perturbation on the hyperplane with different spreads case. Dots are for means, triangle for the standard deviations. Black is for X_1 , red is for X_3 and blue is for X_5 .	145
3.28	Estimated indices $\widehat{S}_{i\delta}$ for the thresholded Ishigami function with a mean shifting	147
3.29	Estimated indices $\widehat{S}_{i,V_{\text{per}}}$ for the thresholded Ishigami function with a variance shifting	149
3.31	1 st quartile perturbation on the thresholded Ishigami test case	149
3.30	5 th percentile perturbation on the thresholded Ishigami test case	150
3.32	Median perturbation on the thresholded Ishigami test case	150
3.33	3 rd quartile perturbation on the thresholded Ishigami test case	151
3.34	95 th percentile perturbation on the thresholded Ishigami test case	152
3.35	Parameters perturbation on the thresholded Ishigami test cas. Triangles correspond to a minimum bound, dots to a maximum bound. X_1 is plotted in red, X_2 in black and X_3 in blue.	153
3.36	Parameters perturbation on the thresholded Ishigami test case. Triangles correspond to a minimum bound, dots to a maximum bound. X_1 is plotted in red, X_2 in black and X_3 in blue.	154
3.37	Estimated indices $\widehat{S}_{i\delta}$ for the flood case with a mean perturbation	156
3.38	5 th percentile perturbation on the flood case	157
3.39	1 st quartile perturbation on the flood case	157
3.40	Median perturbation on the flood case	158
3.41	3 rd quartile perturbation on the flood case	158
3.42	95 th percentile perturbation on the flood case	159
3.43	Parameters perturbation on the flood test case. The indices corresponding to Q are plotted in green: dark green for the location parameter and light green for the scale parameter. The indices corresponding to K_s are plotted as follows: black for the mean, dark grey for the standard deviation. The indices of the mode of Z_v are plotted in red while the ones corresponding to the mode of Z_m are plotted in blue.	160
4.1	Estimated indices $\widehat{S}_{i\delta}$ for the CWNR case with a mean perturbation - 3 variables	170
4.2	5 th percentile perturbation on the CWNR case - 3 variables	170
4.3	1 st quartile perturbation on the CWNR case - 3 variables	171
4.4	Median perturbation on the CWNR case - 3 variables	172
4.5	3 rd quartile perturbation on the CWNR case - 3 variables	172
4.6	95 th percentile perturbation on the CWNR case - 3 variables	173
4.7	Parameters perturbation on the CNWR case - 3 variables	174
4.8	Estimated indices $\widehat{S}_{i\delta}$ for the CWNR case with a mean perturbation - 5 variables	177
4.9	5 th percentile perturbation on the CWNR case - 5 variables	178
4.10	1 st quartile perturbation on the CWNR case - 5 variables	178
4.11	Median perturbation on the CWNR case- 5 variables	179
4.12	3 rd quartile perturbation on the CWNR case - 5 variables	180
4.13	95 th percentile perturbation on the CWNR case - 5 variables	180
4.14	Parameters perturbation on the CNWR case - 5 variables	182
4.15	Estimated indices $\widehat{S}_{i\delta}$ for the CWNR case with a mean perturbation - 7 variables	184
4.16	5 th percentile perturbation on the CWNR case - 7 variables	185
4.17	1 st quartile perturbation on the CWNR case - 7 variables	185
4.18	Median perturbation on the CWNR case- 7 variables	186
4.19	3 rd quartile perturbation on the CWNR case - 7 variables	186
4.20	95 th percentile perturbation on the CWNR case - 7 variables	187

4.21 Parameters perturbation on the CNWR case - 7 variables	188
B.1 Ishigami failure points from a MC sample	206

List of Tables

1	Distributions of the random physical variables of the CWNR model.	24
1.1	Sobol indices for the first failure rectangle	46
1.2	Sobol indices for the second failure rectangle	47
1.3	First order Sobol' indices for the hyperplane 6410 case	53
1.4	Estimated Sobol' indices for the hyperplane 6410 case	53
1.5	Estimated Sobol' indices for the hyperplane 11111 case	54
1.6	Estimated Sobol' indices for the hyperplane 15 variables case	55
1.7	Estimated Sobol' indices for the hyperplane "different spreads" case	55
1.8	Sobol' indices estimation for the thresholded Ishigami function	55
1.9	Estimated Sobol' indices for the flood case	56
1.10	4 dimensional points generated through Sobol' sequence	57
1.11	Estimation of Sobol indices using QMC for the 6410 hyperplane test case	57
1.12	True values of δ_i for the hyperplane 6410 case	69
1.13	Synthesis on the tested SA methods	70
1.14	Correspondence between the general SA objectives and the engineers' motivations	71
2.1	Data set	75
2.2	Confusion matrix of the forest with default parameters	85
2.3	MDA indices of the forest with default parameters	86
2.4	Confusion matrix of the forest with different weights	86
2.5	MDA indices of the forest with different weights	86
2.6	Confusion matrix of the forest with 2000 trees	86
2.7	MDA indices of the forest with 2000 trees	87
2.8	Confusion matrix of the forest built on an IS sample	87
2.9	MDA indices of the forest built on an IS sample	87
2.10	Synthesis on the presented SA methods	105
3.1	Hellinger distance in function of the parameter perturbation	122
3.2	Type of perturbation recommended given the objective or the motivation	126
3.3	Importance factors for hyperplane 6410 function	127
3.4	Estimated Sobol' indices for the hyperplane 6410 case	127
3.5	Hellinger distance in function of the parameter perturbation. The first value is an increase of the parameter (right hand of the graph) whereas the second is a decrease of the parameter (left hand of the graph). Both perturbation lead to the same H^2 departure.	132
3.6	Importance factors for hyperplane 11111 function	133
3.7	Estimated Sobol' indices for the hyperplane 11111 case	134
3.8	Importance factors for the hyperplane 15 variables	138
3.9	Estimated Sobol' indices for the hyperplane with 15 variables case	138

3.10	Importance factors for hyperplane with different spreads function	142
3.11	Estimated Sobol' indices for the hyperplane with different spreads case	143
3.12	Hellinger distance in function of the parameter perturbation	145
3.13	Importance factors for Ishigami function	146
3.14	Sobol' indices estimation for the thresholded Ishigami function	146
3.15	Hellinger distance in function of the parameter perturbation	154
3.16	Importance factors for the flood case	155
3.17	Estimated Sobol' indices for the flood case	155
3.18	Hellinger distance in function of the parameter perturbation	161
4.1	Distributions of the random physical variables of the CWNR model - 3 variables	168
4.2	MDA index - 3 variables	168
4.3	Gini importance - 3 variables	169
4.4	Confusion matrix of the forest - 3 variables	169
4.5	Distributions of the random physical variables of the CWNR model - 5 variables	175
4.6	MDA index - 5 variables	176
4.7	Gini importance - 5 variables	176
4.8	Confusion matrix of the forest - 5 variables	176
4.9	Distributions of the random physical variables of the CWNR model - 7 variables	182
A.1	Distributions of the random physical variables taken for the CWNR models.	203
B.1	Usual hyperplane test cases	205

Context, objectives and outline

On computer experiments

Numerical simulation is the process that allows to reproduce a physical phenomenon with a computer. This phenomenon is represented *via* a mathematical model, and this model is solved during a computation time.

The numerical simulation can be costly, due to the time needed to prepare the set of inputs or to the possibly large number of calculations needed. Moreover, the result of the simulation may be uncertain, thus this scientific topic is often referred to as *numerical experiments*. The use of simulation in conception and safety of an industrial system equipment - two applicative domains of interest in this thesis - has grown over the last decades.

Uncertainty quantification and sensitivity analysis

We briefly present the general framework of our work: the study of a deterministic numerical model. As explained before, a model is a mathematical representation of a complex physical phenomenon.

This model receives inputs and produces outputs (or responses). For the sake of simplicity, these quantities will be considered as scalar and continuous but other types could be considered, modal for instance. Given a certain input value, the model produces a certain output after computation. The deterministic framework is considered here, that is to say that a given set of input values always produces the same output values.

Consider the quantity of interest. It might be possible that the experimenter is interested in a quantity defined from one or several outputs. It is therefore of outmost importance to first define above all study the quantity of interest.

Some parameters (such as physical values) are not precisely characterized due to a lack of data or variability for instance, therefore these parameters can be seen as random variables. Some other inputs will be considered as known and modelled by deterministic values. Let us denote $\mathbf{X} = (X_1, \dots, X_d)$ the d -dimensional random vector (with known density $f_{\mathbf{X}}$) of random (scalar) input variables of the numerical model. Let us also denote by $\mathbf{t}$ the p -dimensional vector of deterministic input. Let us consider without loss of generality, a single output $Y \in \mathbb{R}$ defined as $Y = G(\mathbf{X}, \mathbf{t})$ where G is the deterministic model. The quantity of interest is Z or a function of it. In the following, we will denote $Y = G(\mathbf{X})$. Also, it is important to notice that in the whole thesis, independent inputs will be considered, although the study of models with dependent inputs is a major field of research.

Figure 1 summarizes the reference framework for uncertainty treatment (de Rocquigny *et al.* [30]). The breakdown of the study in several steps is done as follows:

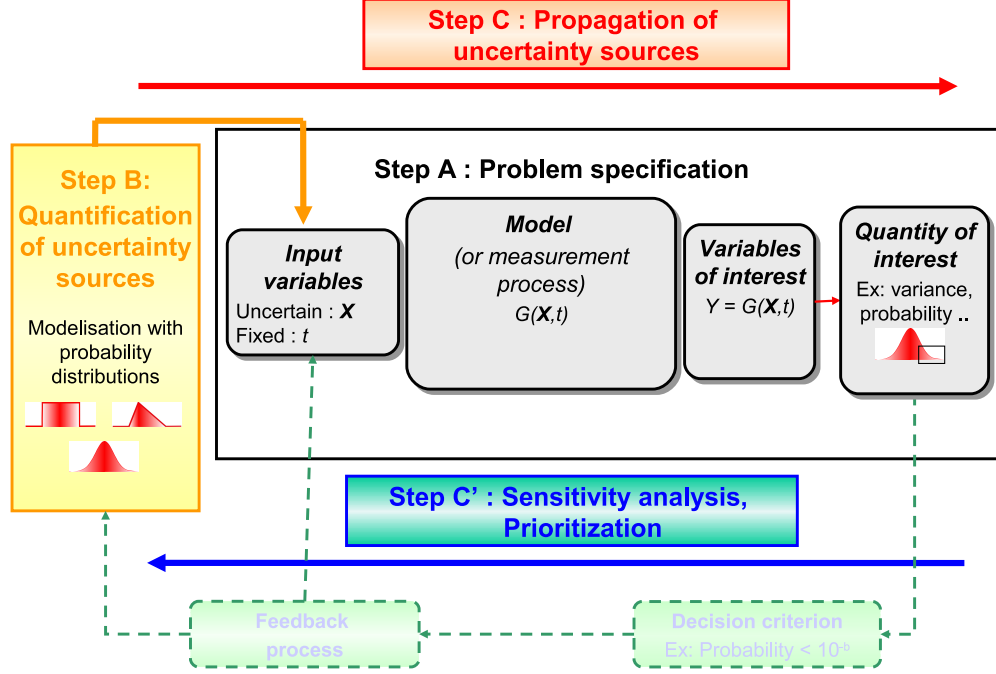


Figure 1: Uncertainty study reference framework

- Step A, problem specification: the objectives are defined, as well as the model used, the quantity of interest and the input variables (some of which are considered uncertain).
- Step B, quantification of uncertainty sources: the input variables considered uncertain are modelled by random distributions. This step is done collaborating with experts and collecting data points.
- Step C, propagation of uncertainty sources: the quantity of interest is evaluated according to the uncertainty on the input variables defined in step B.
- Step C', sensitivity analysis: the relative uncertainty contribution of each input on the output's uncertainty is evaluated.

The genericness allows this framework to address numerous problems. This thesis will mainly focus on Step C', even if this step cannot easily be separated from Step C.

Sensitivity analysis (SA) is defined by Saltelli *et al.* [89] as “the study of how the uncertainty in the output of a model can be apportioned to different sources of uncertainty in the model input”. It may be used to determine the most contributing input variables to an output behaviour. It can also be used to determine non-influential inputs, or ascertain some interaction effects within the model. The objectives of SA are numerous; one can mention model understanding, model simplifying or factor prioritisation.

There are many application examples, for instance Makowski *et al.* [67] analyse, for a crop model prediction, the contribution of 13 genetic parameters on the variance of two outputs. Another

example is given in the work of Varet [99] where the aim of SA is to determine the most influential inputs among a great number (around 60), for an aircraft infrared signature simulation model. In nuclear engineering field, Auder *et al.* [5] study the influential inputs on thermohydraulical phenomena occurring during an accidental scenario, while Iooss *et al.* [50] and Volkova *et al.* [100] consider the environmental assessment of industrial facilities.

The first historical approach to sensitivity analysis is known as the local approach. The impact of small perturbations of the inputs on the output is studied. These small perturbations occur around nominal values (the mean of a random variable for instance). This is a counterpart to the partial derivatives of the model in certain points of the input space. Most of these methods (some of them will be itemized in section 1.3.2) make strong assumptions on the model and/or on the inputs (in terms of linearity, normality, ...).

A second approach, more recent due to the development of computational power is known as the global approach. The whole variation range of the inputs is therein considered. An applicative introduction can be found in Iooss [49]. Most techniques (some of them will be defined in section 1.3.1 and tested in sections 1.4 and 1.5) are developed in an independent approach (“model free”), without making assumptions such as linearity or monotony.

Structural reliability

Consider the industrial problem of knowing if a structure, subject to physical loads or constraints, goes undamaged or goes to a state of failure. This will be referred as structural reliability. A “trial and measures” approach might be possible, but can be difficult to manage for safety or costs reason. Within this context, computer models are used in order to assess the safety of complex systems. These models are then used as an approximate representation of the reality, including some mechanisms such as flaw propagation, friction laws...

In order to completely use the model, uncertainties on the model inputs (essentially physical values) are modelled by random variables. The model is therefore representing the structure gifted with a certain toughness and the environment providing a load. Computation for a fixed set of inputs allows to obtain a failure criterion leading to a binary response: for this set of inputs, the structure fails or behaves soundly.

The fact that uncertainties are modelled by random variables enables risk modelling as a failure probability. This approach is more subtle than a deterministic approach where inputs are fixed to nominal values (generally penalized).

One is interested in the fact that the value $Y \in \mathbb{R}$ given by the failure function G is smaller than a given threshold k (usually 0): it is the failure criterion. The structure is failing for a given set of input $\mathbf{x}$ if $y = G(\mathbf{x}) \leq 0$, where $\mathbf{x} = (x_1, \dots, x_d) \in \mathbb{R}^d$ is a realization of $\mathbf{X}$. The part of space in which this event occurs is called failure domain, denoted D_f . The surface defined by $\{\mathbf{x} \in \mathbb{R}^d, G(\mathbf{x}) = 0\}$ is called limit-state surface. The probability for the event to occur is denoted P_f , failure probability. One has:

$$P_f = \mathbb{P}(G(\mathbf{X}) \leq 0) \quad (1)$$

$$= \int_{D_f} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (2)$$

$$= \int_{\mathbb{R}^d} 1_{G(\mathbf{x}) \leq 0} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (3)$$

$$= \mathbb{E}[1_{G(\mathbf{X}) \leq 0}] \quad (4)$$

The complexity of models and the possible great number of inputs make difficult, in a general case, to compute the exact value of P_f . However, it can be estimated (since written under the form of a mathematical expectation) with the help of several methods that will be itemized in section 1.2. The primer of structural safety is to provide an estimation of P_f and some uncertainty surrounding this estimation. It can be used to answer the original question of the structure supporting the loads.

Context: component within nuclear reactor (CWNR)

This case-study provided the initial motivation for this work. It focuses on the reliability and risk analysis of a nuclear power plant component. However the results of this thesis must be considered as textbook exercises, which can not be used to draw conclusions about the integrity or safety assessment of nuclear power plants.

During the normal operation of a nuclear power plant, the component within nuclear reactor (CWNR) is exposed to ageing mechanisms. In order to assess the integrity of the component, it has been demonstrated that a postulated manufacturing flaw can withstand severe mechanical loads.

The CWNR mechanical model includes three parts. Firstly, a simplified representation of the loading event, which analytically describes as functions of the time, the temperature T , the pressure and the heat transfer coefficient between the environment and the surface of the CWNR. Secondly, a thermo-mechanical model of the CWNR thickness, incorporating the CWNR material properties depending on the temperature. Lastly, an integrity model allowing to evaluate the nocivity of a manufacturing flaw, including different variables: (a) a variable, h , summarizing the dimension of the flaw, (b) a stress intensity factor, (c) the toughness depending on the temperature at the flaw and the level of deterioration, whose discrepancy with operation time is evaluated with some codified forecasting formulas. In practice, the modelling of the CWNR may assign probabilistic distributions to some physical sources of uncertainty. In this manuscript, a maximum of 7 input physical variables will be considered as random. Table 1 summarizes the distributions of the independent physical random inputs of the CWNR model. Table A.1 is a reminder of the inputs' densities.

Random var.	Distribution	Parameters
Thickness (m)	Uniform	$a = 0.0075, b = 0.009$
h (m)	Weibull	$a = 0.02, \text{scale} = 0.00309, \text{shape} = 1.8$
Ratio height/length	Lognormal	$a = 0.02, \ln(\mu) = -1.53, \ln(\sigma) = 0.55$
Azimuth flaw ($^\circ$)	Uniform	$a = 0, b = 360$
Altitude (mm)	Uniform	$a = -5096, b = -1438$
$\sigma \Delta T T$	Gaussian	$\mu = 0, \sigma = 1$
σRes	Gaussian	$\mu = 0, \sigma = 1$

Table 1: Distributions of the random physical variables of the CWNR model.

Also, for the numerical applications over the CWNR model, the random input will be considered as 3, 5 or 7 dimensional and will respectively correspond to the 3, 5 and 7 first random variables presented in Table 1.

Objectives

The aim of this dissertation is the development of sensitivity analysis techniques when the quantity of interest is a probability of exceedance of a given threshold (which is equivalent to a failure probability

in the field of structural reliability). The constraints of the CWNR code are to be taken into account. The expected magnitude of the failure probability is less than 10^{-5} . If possible, the methods must be related to the estimation of P_f and must provide an estimation of the error made when estimating sensitivity indices as well as an estimation of the error made when estimating P_f .

Outline

The following thesis is organised in four chapters.

The first chapter is an overview of both existing strategies for estimating failure probabilities and methods of sensitivity analysis. In this chapter, states of the art for reliability and sensitivity analysis (SA) techniques will be separately developed. More precisely, three main families of reliability techniques will be studied: Monte-Carlo methods, structural reliability methods and sequential Monte-Carlo methods. Finally, two families of well-known sensitivity analysis techniques will be put to the proof on reliability test cases (which are itemized in Appendix B). These techniques show some limitations, confirming the need to develop SA methods focused on failure probabilities. A table (Table 1.13) summarizing the presented methods is proposed, and a discussion on the meaning of sensitivity analysis in the reliability context is conducted.

The second chapter focuses on defining measures of sensitivity in order to produce a variable ranking. More specifically, the use of random forests on a Monte-Carlo sample is proposed in the first place. Two importance measures derived from the random forests predictors are tested on the usual cases. In the second place, a technique using a sample produced by sequential Monte-Carlo methods is elicited. This last method is based on the departure between the marginal distribution of an input and its equivalent given the step of the subset method.

The third chapter presents an original method to estimate the importance of each variable on a failure probability. This method focuses on the impact of perturbations upon the original input densities f_i . A general framework defining appropriate perturbations is elaborated, then sensitivity indices are presented. An estimation technique of these indices that makes no further calls to the model is given. The methodology is then tested on the usual cases.

The fourth chapter presents the application of the developed methods to the CWNR case. Several tunings will be studied to assess or infirm the ability of the different SA methods to identify influential variables.

Chapter 1

State of the art for reliability and sensitivity analysis

1.1 Introduction

The outline of the chapter is the following: in Section 1.2, a state of the art for reliability is proposed. Several techniques for estimating failure probabilities are presented. Then in Section 1.3, a review of Sensitivity Analysis (SA) is given. The application of a well-known SA method, Sobol' indices (1.3.1.3) on a failure probability, is tested on numerous application cases in Section 1.4. In Section 1.5, the so-called *moment independent sensitivity measures* (presented in Section 1.3.1.4) are tested within the reliability context. Next, Section 1.6 proposes a synthesis of these states of the art. Finally, Section 1.7 discusses the meaning and objectives of sensitivity analysis when dealing with failure probabilities.

1.2 State of the art: reliability and failure probability estimation techniques

This state of the art for reliability is widely inspired by the PhD thesis of Gille-Genest [41], Cannaméla [22] (in French) and Dubourg [33] (in English). In addition, monographs by Madsen *et al.* [66] and Lemaire [60] have been used. In this section, a state of the art for the estimation techniques of failure probabilities is detailed. Choice is set to present 3 families of methods.

- Monte-Carlo (MC) simulation methods: these techniques are standard in statistics. The MC methods are used to estimate an expectation. These are based upon an application of the Strong Law of Large Numbers for estimation and on the Limit Central Theorem for error control. Several variance-reduction techniques are available in the literature. The most appropriate of them will be itemised in 1.2.1.
- Reliability methods: historically these methods come from mechanical engineering. They provide answers based upon a linear (FORM) or quadratic (SORM) approximation of the failure surface. This approximation is then used to estimate the failure probability. As far as we know, error control is not easily made. These methods are presented in 1.2.2.
- Subset simulation methods: sometimes also referred as particle methods, sequential MC or splitting techniques, these methods have been more recently developed. They are based upon a decomposition of the objective probability as a product of conditional probabilities, that

are easier to estimate. These estimations are made running a large number of Monte-Carlo Markov Chains (MCMC). Some techniques will be presented in 1.2.3.

However, the partition must be qualified. In practice, methods can be associated; for instance one can first use FORM numerical approximation, then perform some importance sampling around the most probable failing point. In the same way, most of Munoz-Zuniga's works [72] are devoted to a stratified sampling technique (MC variance-reduction method) combined with directional simulation.

1.2.1 Monte-Carlo methods

These methods allow the estimation of an expectation of form:

$$I = \mathbb{E}[\varphi(\mathbf{X})] \quad (1.1)$$

or on the integral form:

$$I = \int_E \varphi(\mathbf{x}) f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (1.2)$$

where $\varphi(\cdot)$ is a function from $E \subset \mathbb{R}^d \rightarrow \mathbb{R}$ and $\mathbf{X}$ is a d -dimensional random vector (with known density $f_{\mathbf{X}}$). In a reliability framework, the function $\varphi(\cdot)$ is written as an indicator, $1_{G(\mathbf{X}) \leq k}$.

1.2.1.1 Crude Monte-Carlo method

Presentation of the estimator The main idea of this method is to generate a large number of i.i.d. vectors with density $f_{\mathbf{X}}$, then to estimate I with the empirical mean of the N values. The Strong Law of Large Numbers allows to get an unbiased estimator of I .

$$\hat{I} = \frac{1}{N} \sum_{i=1}^N \varphi(\mathbf{x}^i) \quad (1.3)$$

with given N and where $\mathbf{x}^i$ are i.i.d with $f_{\mathbf{X}}$. In the reliability case, an unbiased estimator of P_f is:

$$\hat{P} = \frac{1}{N} \sum_{i=1}^N 1_{\{G(\mathbf{x}^i) \leq k\}} \quad (1.4)$$

The variance of the estimator of $\mathbb{E}[\varphi(\mathbf{X})]$ is:

$$\text{Var} [\hat{I}] = \frac{1}{N} \text{Var}[\varphi(\mathbf{X})] \quad (1.5)$$

and it can be estimated by:

$$\widehat{\text{Var}} [\hat{I}] = \frac{1}{N-1} \left[\frac{1}{N} \sum_{i=1}^N \varphi^2(\mathbf{x}^i) - \hat{I}^2 \right] \quad (1.6)$$

When $\varphi(\cdot)$ is an indicator function, as usual in structural reliability studies, a simplified expression can be obtained:

$$\text{Var} [\hat{P}] = \frac{1}{N} P_f (1 - P_f). \quad (1.7)$$

Its classical estimator is:

$$\widehat{\text{Var}} [\hat{P}] = \frac{1}{N} \hat{P} (1 - \hat{P}) \quad (1.8)$$

Thanks to the Limit Central Theorem, one can build confidence intervals around the estimator.

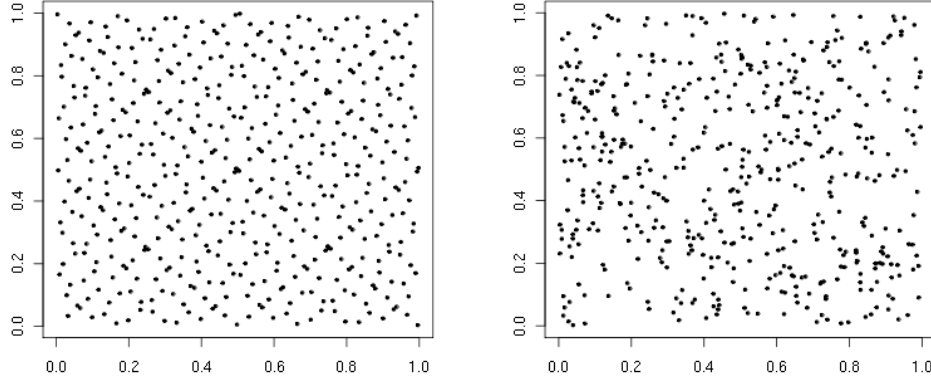


Figure 1.1: Space filling comparison: Sobol's sequence (left) and uniform random sampling (right).

Advantages and drawbacks of the MC method This method makes no hypothesis on the regularity of $\varphi(\cdot)$. The produced estimator is unbiased. Confidence intervals can be obtained around the estimator, which are useful to quantify the precision of the latter. Furthermore, quality of the estimation only depends on the sample size. This means that the MC method is independent of the dimension of the problem, unlike other integration methods.

However, this technique needs a fair number of function calls to reach sufficient precision. According to the rule of thumb, to obtain a variation coefficient of 10% on a 10^{-k} failure probability, $N = 10^{k+2}$ simulations are needed. This can be unrealistic in some applications when dealing with very low failure probabilities ($< 10^{-6}$). Furthermore, computer models can be complex and time-consuming.

Variance-reduction The variance of the estimator decreases in $\text{Var}[\varphi(\mathbf{X})]/N$. Therefore a large sample is needed to get a good estimation. Variance-reduction techniques consist in reducing the uncertainty involved by the numerical integration technique, thus diminishing fluctuations of estimations around the searched value.

In the reference books (see Rubinstein [85]), numerous variance-reduction techniques can be found. In a reliability context, such methods are based on focusing the exploration of the sample space around the limit state (ie, the failure) surface. In the following, we present three main methods.

1.2.1.2 Quasi Monte-Carlo Methods

Presentation of the method The idea beneath Quasi Monte-Carlo (QMC) method is to replace the random sampling by quasi-random sequences. These are deterministic sequences having good equirepartition properties. These sequences are called low-discrepancy sequences, or quasi-random sequences. Loosely speaking, discrepancy is a measure of departure from the uniform distribution. There exist a number of different definitions (L^∞, L^2 , modified $L^2, \dots$). Examples of pseudo-random sequences as well as theoretical developments are given in Niederreiter [75]. Figure 1.1 displays a two-dimensional example of “better” space filling by a low-discrepancy sequence (Sobol's sequence), compared with an uniform random sampling.

QMC estimation of the desired quantity is obtained substituting in the MC estimator the random samples by the pseudo-random samples. However, it is not possible to obtain a variance estimation of the QMC estimator. Koksma-Hlawka's inequality allows to bound the error made when integrating with QMC method, depending on the chosen sequence and on $\varphi(\cdot)$'s regularity.

Reliability case QMC methods are not well adapted for structural reliability. The main issue when estimating small failure probabilities by MC is to get “extreme” samples (within the distribution tail) leading to the failure event, rather than getting evenly distributed samples. However, these methods will be applied in Section 1.4 to decrease the number of function calls when estimating Sobol' indices (which are defined in Section 1.3.1.3).

1.2.1.3 Importance sampling

Presentation of the method The basic idea of importance sampling is to modify the sampling density. The estimator is then obtained by including a density ratio. The aim is to foster sampling in significant regions. In a reliability context, this is simply increasing the number of failure samples. Let us denote $f_{\tilde{\mathbf{X}}}$ a density selected by the practitioner. It will be referred to as the instrumental density. The problem rewrites as follows:

$$I = \int_E \varphi(\mathbf{x}) f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (1.9)$$

$$= \int_E \varphi(\mathbf{x}) \frac{f_{\mathbf{X}}(\mathbf{x})}{f_{\tilde{\mathbf{X}}}(\mathbf{x})} f_{\tilde{\mathbf{X}}}(\mathbf{x}) d\mathbf{x} \quad (1.10)$$

$$= \mathbb{E}_{\tilde{\mathbf{X}}} \left[\varphi(\mathbf{X}) \frac{f_{\mathbf{X}}(\mathbf{x})}{f_{\tilde{\mathbf{X}}}(\mathbf{x})} \right] \quad (1.11)$$

where $\mathbb{E}_{\tilde{\mathbf{X}}}$ is the expectation when $\mathbf{X}$ is of density $f_{\tilde{\mathbf{X}}}$. The estimation is then made by:

$$\hat{I}_{IS} = \frac{1}{N} \sum_{i=1}^N \varphi(\mathbf{x}^i) \frac{f_{\mathbf{X}}(\mathbf{x}^i)}{f_{\tilde{\mathbf{X}}}(\mathbf{x}^i)} \quad (1.12)$$

where $\mathbf{x}^i$ are i.i.d with density $f_{\tilde{\mathbf{X}}}$. One can also get the variance of the estimator:

$$\text{Var}(\hat{I}_{IS}) = \frac{1}{N} \text{Var}_{\tilde{\mathbf{X}}} \left[\varphi(\mathbf{X}) \frac{f_{\mathbf{X}}(\mathbf{X})}{f_{\tilde{\mathbf{X}}}(\mathbf{X})} \right] \quad (1.13)$$

where $\text{Var}_{\tilde{\mathbf{X}}}$ is the variance when $\mathbf{X}$ follows density $f_{\tilde{\mathbf{X}}}$. It should be noticed that the support of $f_{\tilde{\mathbf{X}}}$ must be included within the support of the initial density $f_{\mathbf{X}}$. Otherwise, the estimator is biased.

This technique does not consistently provide a variance reduction. A given instrumental density $f_{\tilde{\mathbf{X}}}$ useful only if:

$$\text{Var}_{\tilde{\mathbf{X}}} \left[\varphi(\mathbf{X}) \frac{f_{\mathbf{X}}(\mathbf{X})}{f_{\tilde{\mathbf{X}}}(\mathbf{X})} \right] < \text{Var}_{\mathbf{X}}[\varphi(\mathbf{X})] \quad (1.14)$$

Minimal variance is obtained with the following optimal density:

$$f_{\mathbf{X}^*}(\mathbf{x}) = \frac{|\varphi(\mathbf{x})| f_{\mathbf{X}}(\mathbf{x})}{\int |\varphi(\mathbf{y})| f_{\mathbf{X}}(\mathbf{y}) d\mathbf{y}} \quad (1.15)$$

However, the denominator on the latter is difficult to estimate as it boils down to I in the case of a positive function $\varphi(\cdot)$. Choosing of a well-fitted instrumental density is a problem in itself. Chapter 2 of Cannaméla [22] provides a state of the art of selecting a quasi-optimal instrumental density.

Reliability context The estimator of P_f is:

$$\hat{P}_{IS} = \frac{1}{N} \sum_{i=1}^N 1_{\{G(\mathbf{x}^i) \leq k\}} \frac{f_{\mathbf{X}}(\mathbf{x}^i)}{f_{\tilde{\mathbf{X}}}(\mathbf{x}^i)} \quad (1.16)$$

Thus the optimal density can be rewritten as:

$$f_{\mathbf{X}^*}(x) = \frac{1_{\{\mathbf{x} \in D_f\}} f_{\mathbf{X}}(\mathbf{x})}{\int_{D_f} f_{\mathbf{X}}(\mathbf{y}) d\mathbf{y}} = \frac{1_{\{\mathbf{x} \in D_f\}} f_{\mathbf{X}}(\mathbf{x})}{P_f} = f_{\mathbf{X}}(\mathbf{x}|D_f) \quad (1.17)$$

This density is intractable in practice, P_f being the quantity of interest. Choice of a good instrumental density is therefore a problem in reliability as well. One can quote chapter 5 of Munoz Zuniga [72] in which an adaptive and non parametric technique for instrumental density selection (adapted to this reliability context) is presented. Additionally, Pastel [79] developed an interesting non-parametric adaptive technique, still within the reliability framework.

1.2.1.4 Directional sampling

In practice this technique is specific to structural reliability studies.

Principle First, the random input vector is transformed into a random vector for which all components are standard Gaussian random variables. This is also referred as transforming the physical space into the standard Gaussian space (sometimes referred to as U-space). Such an isoprobabilistic transformation T which turns the random vector $\mathbf{X}$ of density $f_{\mathbf{X}}$ into a random vector whose all components are independent standard Gaussians. Given $\mathbf{X} = (X_1, \dots, X_d)$ the random input vector, one obtains $\mathbf{U} = (U_1, \dots, U_d) = T(\mathbf{X})$ where $U_i, i = 1, \dots, d$ are independent standard Gaussians. Let us denote:

$$H(\mathbf{u}) = G(T^{-1}(\mathbf{u})) = G(\mathbf{x}). \quad (1.18)$$

Several isoprobabilistic transformations exist. Nataf, generalized Nataf and Rosenblatt transformations (see Lebrun and Dutfoy [58, 59]) are the most adapted. The latter is developed in Appendix C. Once the transformation is done, the quantity of interest can be rewritten as:

$$P_f = \mathbb{P}(H(\mathbf{U}) \leq k). \quad (1.19)$$

The main idea of this method is to generate directions from the center of the standard Gaussian space in a uniform and independent way. Then, the failure function is computed along the directions. Given the direction, this allows a conditional estimation of the failure probability. Vector $\mathbf{U}$ can be rewritten as a product:

$$\mathbf{U} = R\mathbf{A}$$

with $R \geq 0$, R^2 following a χ^2 distribution with d degrees of freedom and $\mathbf{A}$ an uniform random variable on the unit sphere Ω_d , independent of R . Denoting $f_{\mathbf{A}}$ the uniform density on Ω_d , one can rewrite the failure probability conditionally to the directions:

$$P_f = \int_{\Omega_d} \mathbb{P}(H(R\mathbf{a}) \leq k) f_{\mathbf{A}}(\mathbf{a}) d\mathbf{a} \quad (1.20)$$

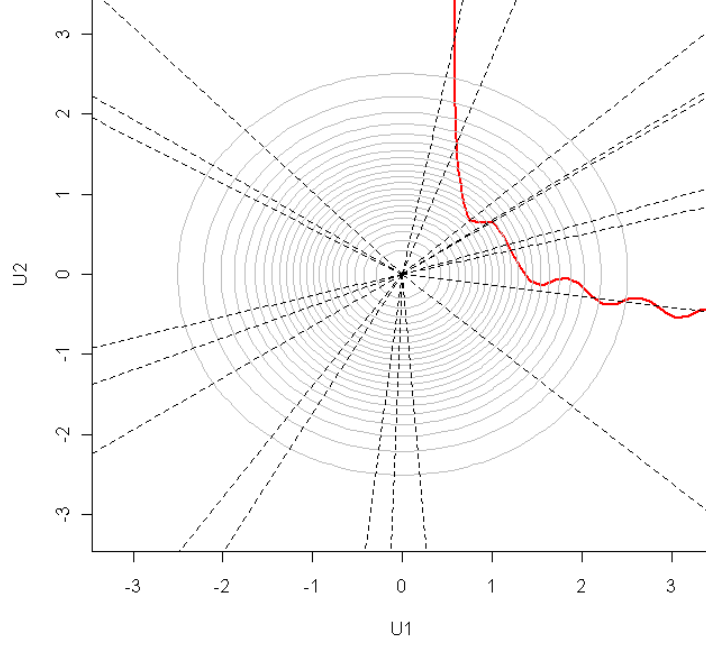


Figure 1.2: 2-dimensional illustration of directional sampling

The directional sampling probability failure estimator thus writes:

$$\hat{P}_{dir} = \frac{1}{N} \sum_{i=1}^N \mathbb{P}(H(R\mathbf{a}_i) \leq k) \quad (1.21)$$

where $\mathbf{a}_i$ are N random independent uniform directions on Ω_d . The variance of the estimator is:

$$\text{Var}[\hat{P}_{dir}] = \frac{1}{N} [\mathbb{E}[\mathbb{P}(H(R\mathbf{A}) \leq k)]^2] - P_f^2 \quad (1.22)$$

Computation of $\mathbb{P}(H(R\mathbf{a}_i) \leq k)$ In practice, one does not have an explicit expression for $H(\cdot)$. It is therefore necessary to use the $G(T^{-1}(\cdot))$ form to get the roots of equation $H(R\mathbf{a}_i) = k$. If r is the only root of the equation, then:

$$\mathbb{P}(H(R\mathbf{a}_i) \leq k) = 1 - \chi_d^2(r^2)G(T(0)) \geq 0. \quad (1.23)$$

If several roots exist ($r_i, i = 1, \dots, n$), one has:

$$\mathbb{P}(H(R\mathbf{a}_i) \leq k) = \sum_i (-1)^{i+1} (1 - \chi_d^2(r_i^2)) \text{ if } G(T(0)) \geq 0. \quad (1.24)$$

A root finding method must be used (the simplest being the dichotomic method). One can fix a bound beyond which the failure probability is considered to be negligible. Figure 1.2 illustrates directional simulation's principle in two dimensions. Isoprobability contours are plotted in grey, the limit-state surface is plotted in red. The dashed lines starting from the center are the directions $\mathbf{a}_i$.

1.2.2 Structural reliability methods

1.2.2.1 Reliability indices

Reliability indices give indications about the relative weights of input parameters in the whole reliability of the considered structure (they are also sometimes called safety index). They allow a comparison of several setups possible. The larger the index, the safer the structure. In the following, two indices are presented.

Hasofer-Lind index Proposed by Hasofer and Lind in 1974 [43], it is an exact geometric index, invariant with respect to the geometry of the limit state surface. It is defined in the Gaussian standard space. Let us define the most probable failure point as the closest failure point to the origin of the standard space (the origin of the standard space is considered outside of the failure domain). Such a point is also referred as a design point. Assuming the design point is unique, one can define the Hasofer-Lind index as the distance between the origin and the design point:

$$\beta_{HL} = \min_{H(u)=0} (u^T u)^{1/2} \quad (1.25)$$

Algorithms to find such design points are numerous, one can quote the Hasofer-Lind-Rackwitz-Fiessler algorithm [82] and its improved version iHLRF (Zhang and Der Kiureghian, [102]). One can also quote a work carried out at EDF R&D about testing the quality of a design point (Dutfoy and Lebrun, [34]). Further details on design point finding algorithms are given in section 1.2.2.2. One should note that an estimation of the Hasofer-Lind index does not require an estimation of P_f but only an estimated design point.

Generalized reliability index The generalized reliability index was proposed by Ditlevsen in 1979 [32] to take account of the curvature of the failure surface around the design point. Defining a reliability measure γ by integrating a weight function (in practice the d -dimensional standard Gaussian distribution) over the safe set S :

$$\gamma = \int_S \varphi_d dS. \quad (1.26)$$

The generalized reliability index is defined as a monotonically increasing function of γ :

$$\beta_G = \Phi^{-1}(\gamma) \quad (1.27)$$

where Φ^{-1} is the inverse cumulative distribution function of the standard Gaussian. One can estimate the index by:

$$\widehat{\beta}_G = \Phi^{-1}(1 - \hat{P}) \quad (1.28)$$

where $\hat{P}$ is an estimation of the failure probability (obtained for instance through MC integration or by FORM/SORM, see section 1.2.2.2). This index equals the Hasofer-Lind one if the failure surface is an hyperplane in the standard Gaussian space. Finally, the estimation of this index requires an estimation of the failure probability P_f .

1.2.2.2 FORM-SORM methods

The *First Order Reliability Method* (FORM) and *Second Order Reliability Method* (SORM) are estimation techniques for a failure probability based upon integration of an approximation of the failure surface. In practice, they are considered as a standard solution in structural reliability since they are not costly, easy to understand and to implement.

These methods proceed in four steps:

- transformation of the input space;
- design point search;
- approximation of the limit-state surface by an hyperplane (FORM) or a quadratic surface (SORM);
- failure probability estimation from the limit-state approximation.

Figure 1.3 graphically summarises the ideas of FORM/SORM.

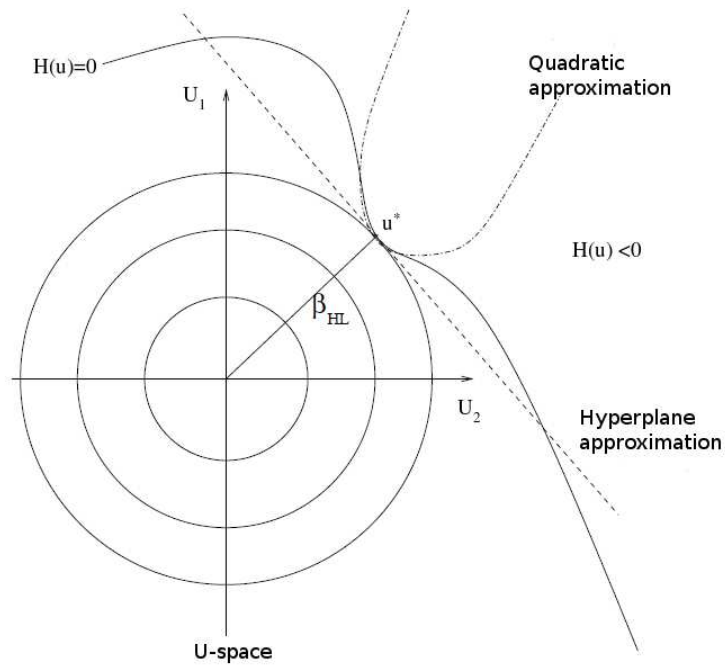


Figure 1.3: Illustration of FORM/SORM

Transformation of the input space It is an isoprobabilistic transformation as described in section 1.2.1.4. These techniques are reminded in appendix B.

Design point search Once within the standard Gaussian space, finding the design point requires to solve the following optimization problem:

$$u^* = \min_{H(u)=0} (u^t u) \quad (1.29)$$

This is a crucial step since it is needed to make as few function calls as possible, while it is required to find all the design points. The objective function is quadratic and convex, thus the minimization difficulties will come from the constraints ($H(u) = 0$). Let us make a distinction between local and global optimization methods.

Local methods Local minima search is efficient if one has an explicit expression of the gradient of H . This is seldom the case in industrial applications and one has to use approximations based upon finite differences. These approximations may be costly in terms of function calls, and they can lead to a loss of convergence of the algorithms. Most algorithms search for the optimum u^* in an iterative way. The idea is, starting from a given point $u^{(k)}$, to find the best descent direction $d^{(k)}$ and the best length of the step $\alpha^{(k)}$:

$$u^{(k+1)} = u^{(k)} + \alpha^{(k)} d^{(k)}. \quad (1.30)$$

The iteration can be followed by a projection.

Numerous methods are described in Lemaire [60], which are divided in 4 main categories : zero order methods, first order methods, second order methods and hybrid methods. Zero order methods (dichotomy for instance) does not require a computation of the gradient. However their convergence is slow. In the reliability case, this implies a large number of function calls. Thus these zero order methods are not adapted to reliability problems considered in this thesis

Here is presented the first order Hasofer-Lind-Rackwitz-Fiessler (HL-RF) algorithm. It has been developed specifically for reliability studies. Its convergence is not assured but the method is effective in many cases. It is worth noticing that the algorithm has been adapted to led to convergence improvements (Abdo and Rackwitz [1]). The iteration is as follows:

$$u^{(k+1)} = (u^{(k)t} \beta^{(k)}) \beta^{(k)} - \frac{H(u^{(k)})}{\|\nabla H(u^{(k)})\|} \beta^{(k)} \text{ with } \beta^{(k)} = \frac{\nabla H(u^{(k)})}{\|\nabla H(u^{(k)})\|} \quad (1.31)$$

Global methods If the limit-state surface presents several design points, the previously described algorithms may not identify these design points. Der Kiureghian and Dakessian [31] proposed to force the convergence of the HL-RF algorithm to a new design point by disturbing the vicinity of the previously found design point.

Approximation of the limit-state surface FORM method replaces the limit-state surface by a hyperplane tangent at the design point. A loss of precision depending on the form of the limit-state surface at the design point occurs. If the limit-state surface is close from the hyperplane, this method provides good precision compared to the needed number of function calls. The linear approximation writes as follows:

$$\nabla H(u) \big|_{u=u^*}^t (u - u^*) = 0 \quad (1.32)$$

The SORM method replaces the limit-state surface by a second-order (quadratic) hypersurface. Such a method requires the estimation of the curvature of the limit-state surface at the design point u^* . Several techniques are provided in Lemaire [60]. The key message is that the use of SORM over FORM is justified when it is known that the surface is almost quadratic.

Failure probability estimation In the FORM approximation, one uses the Hasofer-Lind reliability index presented in section 1.2.2.1 and estimates P_f with:

$$\hat{P}_{FORM} = 1 - \Phi(\beta_{HL}) \quad (1.33)$$

SORM approximation is a more complex problem, for which an asymptotic approximation was provided by Breitung [18].

On the geometric approximations FORM/SORM techniques are popular methods in the domain of structural reliability, because a few function calls are needed to get an estimation of P_f . Also, FORM is easy to understand and to implement. However, a significant error can be made when using FORM. Consequently, these methods should only be used when it is known that the limit-state surface has a given geometrical shape (almost hyperplane or almost quadratic). Such an information is not always available.

1.2.3 Subset simulation

1.2.3.1 Introduction

Subset simulation methods are based upon a division of the failure probability in a product of conditional probabilities. These are larger therefore easier to estimate. Let us consider a sequence of $M + 1$ thresholds T such as:

$$T = \{+\infty, t_1, \dots, t_M = 0\}$$

and let us also define the sequence of nested subsets (also sometimes referred to as intermediate failure events):

$$A_k = \{\mathbf{x} | G(\mathbf{x}) < t_k\}.$$

One has:

$$\mathbb{P}[\mathbf{x} \in A_k] = \prod_{i=1}^k \mathbb{P}[\mathbf{x} \in A_i | \mathbf{x} \in A_{i-1}]$$

and one can rewrite P_f as:

$$P_f = \mathbb{P}[\mathbf{x} \in A_M] = \prod_{i=1}^{M+1} \mathbb{P}[\mathbf{x} \in A_i | \mathbf{x} \in A_{i-1}] \quad (1.34)$$

thus the estimation of P_f is reduced to the estimation of the conditional failure probabilities. The name “subset simulation” has been introduced by Au and Beck [4]. For the sake of simplicity, let us denote:

$$\mathbb{P}(A_k) = \mathbb{P}[\mathbf{x} \in A_k]$$

and

$$\mathbb{P}(A_k | A_{k-1}) = \mathbb{P}[\mathbf{x} \in A_k | \mathbf{x} \in A_{k-1}].$$

The algorithm first step is to estimate $\mathbb{P}(A_1)$ by standard Monte Carlo simulation. One has:

$$\widehat{P(A_1)} = \frac{1}{N} \sum_{k=1}^N 1_{\{G(\mathbf{x}^i) < t_1\}}$$

where $\mathbf{x}^i$ are i.i.d. to f . MCMC techniques are thereafter used to estimate the conditional failure probabilities $\mathbb{P}(A_k | A_{k-1}), k = 2, \dots, M$. Let us denote:

$$f(\mathbf{x} | A_i) = \frac{f(\mathbf{x}) 1_{\{G(\mathbf{x}) < t_i\}}}{P(A_i)} \quad (1.35)$$

the conditional density of x given that the i -th threshold has been reached. The goal of the algorithms displayed in the following is to sample according to this objective distribution. As the denominator is an unknown quantity, indirect sampling of the objective distribution is needed, which is practically made using Monte Carlo Markov Chains (MCMC).

1.2.3.2 Algorithm

A so-called modified Metropolis algorithm is presented in the Au and Beck's [4] original article. The modification is operated to allow the practitioner to deal with high-dimensional densities. Let us first recall the Metropolis algorithm.

Metropolis algorithm Let us denote a proposal density $p^*(\epsilon|\mathbf{x})$, a joint d -dimensional density, centred in $\mathbf{x}$ with a symmetry property $p^*(\epsilon|\mathbf{x}) = p^*(\mathbf{x}|\epsilon)$. We are interested in the production of the sample $\mathbf{x}^{(i+1)}$, lying in the subset A_k . It is generated starting from the initial sample $\mathbf{x}^{(i)} \in A_k$ as follows:

- Sampling of the candidate sample $\tilde{\mathbf{x}}$: ϵ is simulated according to $p^*(\epsilon|\mathbf{x}^{(i)})$. Ratio $r = f(\epsilon)/f(\mathbf{x}^{(i)})$ is computed. The candidate sample is $\tilde{\mathbf{x}} = \epsilon$ with probability $\min(1, r)$ and stays $\tilde{\mathbf{x}} = \mathbf{x}^{(i)}$ with probability $1 - \min(1, r)$.
- Acceptance/rejection of the candidate $\tilde{\mathbf{x}}$: one checks that $\tilde{\mathbf{x}}$ lies within the interest zone A_k . If $G(\tilde{\mathbf{x}}) < s_k$ then $\mathbf{x}^{(i+1)} = \tilde{\mathbf{x}}$. Else, $\mathbf{x}^{(i+1)} = \mathbf{x}^{(i)}$.

According to the authors, this algorithm is not robust to the large dimension, given a high rejection rate. This rejection rate implies a high correlation within the produced samples, thus reducing the efficiency of the simulation process. The authors then propose a modified Metropolis algorithm to cope with the simulation of random vectors of high dimension.

Modified Metropolis algorithm For all dimensions $j = 1, \dots, d$ let us denote $p_j^*(\epsilon|x_j)$, a 1-dimensional proposal density, centred in x_j with a symmetry property $p_j^*(\epsilon|x_j) = p_j^*(x_j|\epsilon)$. The sample $\mathbf{x}^{(i+1)}$, lying in the subset A_k , is generated starting from the initial sample $\mathbf{x}^{(i)} \in A_k$ as follows:

- Sampling of the candidate sample $\tilde{\mathbf{x}}$: for each component j , let us sample ϵ_j according to $p_j^*(\epsilon|x_j^{(i)})$. Ratio $r_j = f_j(\epsilon)/f_j(x_j^{(i)})$ is computed. Candidate's j -th component is thus $\tilde{x}_j = \epsilon_j$ with probability $\min(1, r_j)$ and is $\tilde{x}_j = x_j^{(i)}$ with probability $1 - \min(1, r_j)$.
- Acceptance/rejection of the candidate $\tilde{\mathbf{x}}$: one checks that $\tilde{\mathbf{x}}$ lies within the interest zone A_k . If $G(\tilde{\mathbf{x}}) < s_k$ then $\mathbf{x}^{(i+1)} = \tilde{\mathbf{x}}$. Else, $\mathbf{x}^{(i+1)} = \mathbf{x}^{(i)}$.

The authors show that the Markov chain generated through this algorithm has stationary distribution $f(\mathbf{x}|A_k)$. The choice of proposal density is important, the authors state that the method is more sensible to the spread of the proposal densities than to their structural form (e.g., Gaussian, gamma, etc.). Based on this observation, the authors recommend to use uniform densities.

On the threshold choice The authors acknowledge that the choice of the threshold is essential in the simulation process. Thus, their advice is to choose an adaptive choice of the threshold so that the conditional probabilities $\mathbb{P}(A_k|A_{k-1})$ are fixed.

1.2.3.3 Theoretical results and strategies

Cérou *et al.* [23] present, from a theoretical point of view, two strategies to estimate small failure probabilities. The difference between these two methods lies in the adaptive selection of the threshold for the second.

Fixed levels algorithm The authors consider a transition Markov kernel K on $\mathbb{R}^d$ which is f -symmetric (thus f -invariant):

$$f(dx)K(x, dy) = f(dy)K(y, dx).$$

A Metropolis-Hasting kernel is proposed (as in Au and Beck [4]). The authors then consider a Markov chain $(X_k)_{k \geq 0}$ such that the initial density is f . The generation algorithm is as follows: a particle cloud of size N is sampled, one has $X_0^{(j)} \sim f$, $j = 1, \dots, N$. For each level $k = 1, \dots, M$, let us denote I_{k+1} the indices of the particles that reach the level of interest:

$$I_{k+1} = \{j | X_k^{(j)} \in A_{k+1}\}$$

conditional probability $\mathbb{P}(A_{k+1}|A_k)$ is estimated by $\hat{p}_{k+1} = \frac{|I_{k+1}|}{N}$. For the j of I_{k+1} , $\tilde{X}_{k+1}^{(j)} = X_k^{(j)}$ is proposed. For the j that are not in I_{k+1} , $\tilde{X}_{k+1}^{(j)}$ is randomly chosen (uniformly) as a copy of one of the particle in I_{k+1} . Thus each particle of $(\tilde{X}_{k+1})$ lies in A_{k+1} . Then for each particle indexed by $j = 1, \dots, N$, transition is twofold. First step is to mutate (or shake) the particle by applying (potentially several times) kernel K , producing the candidate particle Z :

$$Z \sim K(\tilde{X}_{k+1}^{(j)}, \cdot)$$

The second step is a post-mutation selection $X_{k+1}^{(j)} = Z$ if $Z \in A_{k+1}$, $X_{k+1}^{(j)} = \tilde{X}_{k+1}^{(j)}$ else. The particle cloud is then distributed according to $f(\mathbf{x}|A_{k+1})$. Failure probability P_f is then estimated by the product of the estimators of the conditional probabilities:

$$\hat{P} = \prod_{k=1}^M \hat{p}_k \tag{1.36}$$

The authors show the asymptotic normality of $\hat{P}$.

$$\sqrt{N} \frac{\hat{P} - P_f}{P_f} \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, \sigma^2) \tag{1.37}$$

where σ^2 has a complex expression, given in section 2.3 of Cérou *et al.* [23].

Adaptive levels algorithm The estimator produced by the fixed levels algorithm reaches minimal variance when the levels are evenly spaced (in probability), see Lagnoux [56]. The authors then propose another algorithm fixing the levels on the fly (adaptively). Let us consider a number $\alpha \in [0, 1]$, success rate between two levels. At each step, the threshold set is the α -quantile (or the αN particles which $G(\cdot)$ values are the smallest) of the current sample. The algorithm stops when the α -quantile of the sample is lower than 0. One notices that the number of steps is a random variable. However, for a cloud size N large enough, the number of steps is:

$$n_s = \lfloor \frac{\log P_f}{\log \alpha} \rfloor \tag{1.38}$$

The authors also show the asymptotic normality of $\hat{P}$.

$$\sqrt{N} (\hat{P} - P_f) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, \sigma^2) \tag{1.39}$$

where $\sigma^2 = P^2 \left(n_s \frac{1-\alpha}{\alpha} + \frac{1-r_0}{r_0} \right)$ with $r_0 = P\alpha^{-n_s}$. The estimator $\hat{P}$ is biased. This bias is positive and decreases with a $\frac{1}{N}$ rate. However, the adaptive algorithm is more efficient than the fixed levels algorithm, in terms of mean square error (MSE).

On the tuning of parameters In the following, the adaptive algorithm presented in Cérou *et al.* [23] will be used. Several parameters are yet to be tuned: N , α and the Markov kernel (or proposal density) choice. Balesdent *et al.* [7] also recommend to tune the number of application of the kernel.

- For the α , authors of Cérou *et al.* [23] recommend to take α of order 0.75. On the other hand, authors of Au and Beck [4] propose to take α of order 0.1. Unless otherwise mentioned, we have chosen to take $\alpha = 0.75$.
- The choice of N depends on the studied problem and on the complexity of the studied numerical model. Unless otherwise mentioned, we have chosen to take $N = 10^4$.
- The choice of the Markov kernel (or proposal density) is the most crucial point. Both articles [4] and [23] let the practitioner choose the parameter according to the problem. The chosen density will be given for each example.

1.3 Sensitivity analysis (SA)

In this section, the main methods of SA will be developed. The motivations have been presented in page 22. Additionally, a deeper discussion of these motivations, that proposes new guidelines for conducting SA for failure probabilities is provided in section 1.7.

1.3.1 Global sensitivity analysis

Global SA methods are used to identify the inputs contributing to the output variability, considering the whole input support. The methods presented in this subsection, which is inspired by Iooss [49], are divided into four main classes. The first will be the screening methods, designed to deal with a large number of inputs. The second class is composed of the methods based on the analysis of linear models, where a linear model is fitted and its by-products are used to perform SA. The third class contains methods based on a variance decomposition of the output. Finally, some moment-independent methods will be presented in the fourth class.

1.3.1.1 Screening methods

Screening methods are based on a discretisation of the inputs in levels, allowing a quick exploration of the code behaviour. These methods are adapted to a fair number of inputs; practice has often shown that only a small number of inputs are influential. The choice has been made to present Morris method [71]. The aim of this type of method is to identify the non-influential inputs in a small number of model calls. The model is therefore simplified before using other SA methods, more subtle but more costly.

The method of Morris allows to classify the inputs in three groups: inputs having negligible effects; inputs having linear effects without interactions and inputs having non-linear effects and/or with interactions. The method consists of discretising the input space for each variable, then performing a given number of OAT designs (one-at-a-time design of experiments, in which only one input varies). Such designs of experiments are randomly chosen in the input space, and the variation direction is also random. The repetition of these steps allows the estimation of elementary effects for each input. From these effects are derived sensitivity indices.

Let us denote r the number of OAT designs (Saltelli *et al.* [89] propose to set parameter r between 4 and 10). Let us discretise the input space in a d -dimensional grid with n levels per

input. Let us denote $E_j^{(i)}$ the elementary effect of the j -th variable obtained at the i -th repetition, defined as:

$$E_j^{(i)} = \frac{G(\mathbf{X}^{(i)} + \Delta e_j) - G(\mathbf{X}^{(i)})}{\Delta} \quad (1.40)$$

where Δ is a predetermined multiple of $\frac{1}{(n-1)}$ and e_j a vector of the canonical base. Indices are obtained as follows:

- $\mu_j^* = \frac{1}{r} \sum_{i=1}^r |E_j^{(i)}|$ (mean of the absolute value of the elementary effects),
- $\sigma_j = \sqrt{\frac{1}{r} \sum_{i=1}^r \left(E_j^{(i)} - \frac{1}{r} \sum_{i=1}^r E_j^{(i)} \right)^2}$ (standard deviation of the elementary effects).

The interpretation of the indices is the following:

- μ_j^* is a measure of influence of the j -th input on the output. The larger μ_j^* is, the more the j -th input contributes to the dispersion of the output.
- σ_j is a measure of non-linear and/or interaction effects of the j -th input. If σ_j is small, elementary effects have low variations on the support of the input. Thus the effect of a perturbation is the same all along the support, suggesting a linear relationship between the studied input and the output. On the other hand, the larger σ_j is, the less likely the linearity hypothesis is. Thus a variable with a large σ_j will be considered having non-linear effects, or being implied in an interaction with at least one other variable.

Then, a graph linking μ_j^* and σ_j allows to distinguish the 3 groups.

1.3.1.2 Methods based on the analysis of linear models

If a sample of inputs and outputs large enough is available, it is possible to fit a linear model explaining the behaviour of Y given the values of the random vector $\mathbf{X}$. Global sensitivity measures defined through the study of the fitted model are available and presented in the following. Statistical techniques allow to confirm the linear hypothesis. If the hypothesis is rejected, but that the monotony of the model is confirmed, one can use the same measures using a rank transformation. Main indices are:

- Pearson correlation coefficient:

$$\rho(X_j, Y) = \frac{\sum_{i=1}^N (X_j^{(i)} - \mathbb{E}(X_j))(Y_i - \mathbb{E}(Y))}{\sqrt{\sum_{i=1}^N (X_j^{(i)} - \mathbb{E}(X_j))^2} \sqrt{\sum_{i=1}^N (Y_i - \mathbb{E}(Y))^2}}. \quad (1.41)$$

It can be seen as a linearity measure between variable X_j and output Y . It equals 1 or -1 if the tested input variable has a linear relationship with the output. If X_j and Y are independent, the index equals 0.

- Standard Regression Coefficient (SRC):

$$\text{SRC}_j = \beta_j \sqrt{\frac{\text{Var}(X_j)}{\text{Var}(Y)}} \quad (1.42)$$

where β_j is the linear regression coefficient associated to X_j . SRC_j^2 represents a share of variance if the linearity hypothesis is confirmed.

- Partial Correlation Coefficient (PCC):

$$\text{PCC}_j = \rho(X_j - \widehat{X}_{-j}, Y - \widehat{Y}_{-j}) \quad (1.43)$$

where $\widehat{X}_{-j}$ is the prediction of the linear model, expressing X_j with respect to the other inputs and $\widehat{Y}_{-j}$ is the prediction of the linear model where X_j is absent. PCC measures the sensitivity of Y to X_j when the effects of the other inputs have been cancelled.

1.3.1.3 Functional decomposition of variance : Sobol' indices

When the model is non-linear and non-monotonic, the decomposition of the output variance is still defined and can be used for SA. Let us have $f(\cdot)$ a square-integrable function, defined on the unit hypercube $[0, 1]^d$. It is possible to represent this function as a sum of elementary functions (Hoeffding [46]):

$$G(\mathbf{X}) = G_0 + \sum_{i=1}^d G_i(X_i) + \sum_{i < j}^d G_{ij}(X_i, X_j) + \cdots + G_{12\dots d}(\mathbf{X}) \quad (1.44)$$

This expansion is unique under condition (Sobol' [92]):

$$\int_0^1 G_{i_1 \dots i_s}(x_{i_1}, \dots, x_{i_s}) dx_{i_k} = 0 \quad , 1 \leq k \leq s, \quad \{i_1, \dots, i_s\} \subseteq \{1, \dots, d\}.$$

This implies that G_0 is a constant.

In the SA framework, let us have $\mathbf{X} = (X_1, \dots, X_d)$, a random vector where the variables are mutually independent and $Y = G(\mathbf{X})$, output of a deterministic code $G(\cdot)$. Thus a functional decomposition of the variance is available, often referred as functional ANOVA:

$$\text{Var}[Y] = \sum_{i=1}^d D_i(Y) + \sum_{i < j}^d D_{ij}(Y) + \cdots + D_{12\dots d}(Y) \quad (1.45)$$

where $D_i(Y) = \text{Var}[\mathbb{E}(Y|X_i)]$, $D_{ij}(Y) = \text{Var}[\mathbb{E}(Y|X_i, X_j)] - D_i(Y) - D_j(Y)$ and so on for higher order interactions. The so-called ‘‘Sobol' indices’’ or ‘‘sensitivity indices’’ (Sobol' [92]) are obtained as follows:

$$S_i = \frac{D_i(Y)}{\text{Var}[Y]}, \quad S_{ij} = \frac{D_{ij}(Y)}{\text{Var}[Y]}, \quad \dots$$

These indices express the share of variance of Y that is due to a given input or input combination. The number of indices grows in an exponential way with the number d of dimension: there are $2^d - 1$ indices. For computational time and interpretation reasons, the practitioner should not

estimate indices of order higher than two. Homma and Saltelli [47] introduced the so-called “total indices” or “total effects” that writes as follows:

$$S_{T_i} = S_i + \sum_{i < j} S_{ij} + \sum_{j \neq i, k \neq i, j < k} S_{ijk} + \dots = \sum_{l \in \#i} S_l \quad (1.46)$$

where $\#i$ are all the subsets of $(1\dots d)$ including i . In practice, when d is large, only the main effects and the total effects are computed, thus giving a good information on the model sensitivities. Main methods for the estimation of such indices are presented in section 1.4. These indices will be tested in the reliability framework.

1.3.1.4 Moment independent importance measure

In this part, 4 indices that have a moment independence property are presented. Most of them are based on the idea that the importance measure is a distance or a divergence between the distribution of the output (denoted f_{Y_0}) and the distribution of the output given a condition on one or several inputs. Such measures are moment independent, meaning they do not require any computation of the moments of the output. Furthermore, such indices might be suited when the variance poorly represents the variability of the distribution (for instance for multimodal distributions)

Kullback-Leibler divergence index (Park and Ahn) In order to assess the importance of a variable, Park and Ahn [78] proposed to use the Kullback-Leibler (KL) divergence between the distribution of the output, and another distribution f_{Y_i} . Recall that between two pdf p and q the KL divergence is defined as:

$$KL(p, q) = \int_{-\infty}^{+\infty} p(y) \log \frac{p(y)}{q(y)} dy \text{ if } \log \frac{p(y)}{q(y)} \in L^1(p(y)dy). \quad (1.47)$$

The proposed sensitivity index reads as follows:

$$I(i; 0) = \int_{\mathbb{R}} f_{Y_i}(y) \log \left[\frac{f_{Y_i}(y)}{f_{Y_0}(y)} \right] dy \quad (1.48)$$

and can be interpreted as “*the mean information for discrimination in favor of f_{Y_i} against f_{Y_0}* ”. It is clear that the larger the index, the more important the variable. The authors then propose some input distributional changes.

Entropy index (Krzykacz-Hausmann) Krzykacz-Hausmann [55] proposes a sensitivity index based on entropy arguments that is defined as follows. First recall the entropy of an output:

$$H(Y) = - \int_{\mathbb{R}} f_{Y_0}(y) \log f_{Y_0}(y) dy \quad (1.49)$$

that can be interpreted as “*the measure of the total uncertainty of Y* ”. Then, one can define the expectation of the conditional entropy of Y given X_i :

$$H(Y|X_i) = \mathbb{E}_{X_i} [H(Y|X_i)] \quad (1.50)$$

Given these two quantities, the author defines the following sensitivity index:

$$\eta_i = \frac{H(Y) - H(Y|X_i)}{H(Y)} = 1 - \frac{H(Y|X_i)}{H(Y)} \quad (1.51)$$

which is “a representation of the information learnt on Y based on the knowledge of X_i ” (Auder and Iooss [6]).

Relative entropy index (Liu *et al.*) Liu *et al.* [65] introduce an index representing how much the output varies in distribution when an input is fixed to its mean. Recall that the distribution of the output Y is denoted f_{Y_0} . Then, one can fix one input X_i to its mean, namely $\bar{x}_i$. The pdf of Y after such a change is denoted f_{Y_i} . The sensitivity index can then be defined as follows, using a modified version of KL divergence:

$$KL_i(f_{Y_i}|f_{Y_0}) = \int_{\mathbb{R}} f_{Y_0}(y(x_1, \dots, x_i, \dots, x_n)) \left| \log \frac{f_{Y_i}(y(x_1, \dots, \bar{x}_i, \dots, x_n))}{f_{Y_0}(y(x_1, \dots, x_i, \dots, x_n))} \right| dy \quad (1.52)$$

The larger the index is, the more influential the input is. The authors present their index as a total effect of X_i . Another measure of importance is obtained by setting all the input but X_i to their mean, but will not be presented here. It is worth noticing that the authors derived their index in the reliability case, where the quantity of interest is a failure probability. Denoting P_f the original failure probability and $\bar{P}_f$ the failure probability when X_i is fixed at $\bar{x}_i$, the index becomes:

$$KL_i(\bar{P}_f|P_f) = \bar{P}_f \log \frac{\bar{P}_f}{P_f} + (1 - \bar{P}_f) \log \frac{1 - \bar{P}_f}{1 - P_f} \quad (1.53)$$

A moment free importance measure (Borgonovo) The objective of the work of Borgonovo [13] was to propose an importance measure without reference to any particular moment of the output. Recall that the distribution of the output Y is denoted f_{Y_0} and denote f_{Y/X_i} the conditional density of the output given that one of the inputs (X_i) is fixed to a given value, say x_i^* , one can define the density shift between these two densities:

$$s(X_i) = \int |f_{Y_0}(y) - f_{Y/X_i}(y)| dy. \quad (1.54)$$

This quantity can be seen as the area between the two pdfs. In order to take the whole range of variation of X_i , one defines the expected shift as follows:

$$\mathbb{E}_{X_i}[s(X_i)] = \int f_{X_i}(x_i) \left[\int |f_{Y_0}(y) - f_{Y/X_i}(y)| dy \right] dx_i. \quad (1.55)$$

Thus the moment independent measure is defined as:

$$\delta_i = \frac{1}{2} \mathbb{E}_{X_i}[s(X_i)] \quad (1.56)$$

and it represents the normalised expected shift in the distribution of Y due to X_i . It is worth noticing that the author extends the definition of the sensitivity index to any group of inputs. Such an index is denoted $\delta_{i_1, \dots, i_r}$. The sense of δ_i proposed by the author is to determine “the model input that, if determined, would lead to the greatest expected modification in the distribution of Y ”. Additionally, one can present the sensitivity measure of Y to X_j conditionally to X_i as follows:

$$\delta_{j|i} = \frac{1}{2} \int f_{X_i}(x_i) f_{X_j}(x_j) \left[\int |f_{Y/X_i}(y) - f_{Y/X_i, X_j}(y)| dy \right] dx_i dx_j$$

which represents the sensitivity of Y to X_j when X_i is determined.

The properties of δ_i are presented:

- $0 \leq \delta_i \leq 1$.
- If Y does not depend on X_i , then $\delta_i = 0$.
- $\delta_{1,\dots,d} = 1$.
- If Y depends on X_i but independent of X_j , then $\delta_{ij} = \delta_i$.
- Any bidimensional index is bounded: $\delta_i \leq \delta_{ij} \leq \delta_i + \delta_{j|i}$.
- The indices are invariant to any monotonic transformation of the output (scale invariant).

This index, having useful properties, will be tested in 1.5. The importance measure defined in Borgonovo [13] has been extended in Borgonovo *et al.* [14], where a new computation procedure is proposed. Additionally, Caniou [21] proposes an index estimation procedure, based on kernel smoothing estimation of the conditional pdfs then on a quadrature estimation of the shift. Another procedure based on kernel smoothing of the cdfs is tested as well. This index will be tested in the reliability context in section 1.5.

1.3.2 Reliability based sensitivity analysis

The reliability community produced specific methods to estimate a failure probability, as seen for instance in section 1.2.2. The question of the sensitivity of the failure probability to the input parameters arose in this context. Specific SA methods have been produced to meet these expectations. In this subsection methods based on partial derivatives are presented, as well as methods based on the search of a design point in the standard space. The reference here is chapter 6 of Lemaire [60].

1.3.2.1 Sensitivity measure based on partial derivatives

The main idea of this measure is to estimate the sensitivity of the probability of failure to a parameter. From the formulation of the Hasofer-Lind index (see section 1.2.2.1), one has:

$$P_f \simeq 1 - \phi(\beta_{HL})$$

Denoting by p_i the parameter (mean, standard deviation, ...) of an input distribution, then the index is:

$$\frac{\partial P_f}{\partial p_i} = \frac{\partial P_f}{\partial \beta_{HL}} \frac{\partial \beta_{HL}}{\partial p_i} = -\phi(\beta_{HL}) \frac{\partial \beta_{HL}}{\partial p_i} |_{u^*} \quad (1.57)$$

Such an index cannot be used to compare parameters. Indeed, the value of the derivative depends on the way to express the parameter p_i (it depends for instance of its unit), leading to some scale effect. To allow a comparison between parameters, one introduces the *elasticity* Lemaire [60], which is a dimensionless quantity:

$$e_{p_i} = \frac{p_i}{P_f} \frac{\partial P_f}{\partial p_i} \quad (1.58)$$

However, this quantity is non-informative when dealing with parameters of value 0. Moreover both of the presented methods are very dependent on the quality of the founded design point. Since they consider the impact of a variation in the vicinity of the design point, these can be qualified as local SA methods.

1.3.2.2 Global sensitivity measures

The importance factors are by-products of the FORM/SORM methods. These sensitivity measures aim at quantifying the importance of a variable on the failure probability. Since they quantify the impact of a variable on the failure probability, they can be qualified as global SA methods, but since they strongly depend on the approximation of the design point, they can also be qualified as local SA methods.

From the design point u^* one writes:

$$u^* = \beta_{HL} \alpha^* \quad (1.59)$$

where β_{HL} is the distance between the origin of the standard space and u^* ; and α^* is the normalised vector of direction. Then for each variable U_i , one can obtain

- the importance factor: α_i^{*2} , which sums to one and are therefore sometimes plotted as a pie chart.
- the direction cosine: α_i^* . One gets $\alpha_i^* = \frac{\partial \beta_{HL}}{\partial u_i} |_{u^*}$, this formula justifies the use of α_i^* as a sensitivity index.

However, these measures depend on the founded design point in the standard space, therefore they are not related to the variables in the physical space. Consequently their interpretation in the physical space might be complicated. Furthermore, they do not take the shape of the limit-state surface into account.

1.4 Functional decomposition of variance for reliability

In the context of global SA, a widespread technique is based upon the functional decomposition of variance, as presented in section 1.3.1.3. This section presents some works on the application of such a method for reliability problems. At first, really simple toy models will be used in 1.4.1 to provide an intuition about the meaning of Sobol' indices applied to reliability. Then in 1.4.2, some estimation techniques for the Sobol' indices are presented and their properties are discussed. The application on the presented test cases (Appendix B) is done in 1.4.3. Two techniques of variance-reduction are tested in 1.4.4 and in 1.4.5, respectively Quasi Monte-Carlo techniques (QMC) and Importance Sampling (IS). An original work on the first-order indices within the failure domain is proposed in 1.4.6. Finally, a conclusion about the use of Sobol' indices in the reliability context is proposed in 1.4.7.

1.4.1 First applications

Let us recall that:

$$P_f = \mathbb{E}[1_{G(\mathbf{X}) \leq 0}]$$

This failure probability depends on the distribution of $\mathbf{X}$. We will then consider the function from $\mathbb{R}^d$ to $\mathbb{R}$, so that $\mathbf{X}$ maps to $1_{G(\mathbf{X}) \leq 0}$ as the studied function $f(\cdot)$ defined in section 1.3.1.3. Therefore the functional decomposition of variance can be applied, provided that the components of $\mathbf{X}$ are independent. In the following, a toy example where the indices can easily be computed is studied. The aim is to verify if the indices are adapted to the objective.

Failure rectangle For this first toy example, the failure function has expression:

$$1_{G(\mathbf{X}) \leq k} = 1_{\{0.1 < X_1 < 0.2\} \{0 < X_2 < 0.8\}}(\mathbf{X})$$

where $\mathbf{X} = (X_1, X_2)$ with $X_1, X_2 \sim \mathcal{U}[0, 1]$, the two inputs being independent. The failure probability is $P_f = \mathbb{E}[1_{G(\mathbf{X}) \leq k}] = 8 \times 10^{-2}$ and the variance is $\text{Var}[1_{G(\mathbf{X}) \leq k}] = P_f(1 - P_f) = 3.6 \times 10^{-3}$. The conditional expectation of the output given the input is plotted on figure 1.4.

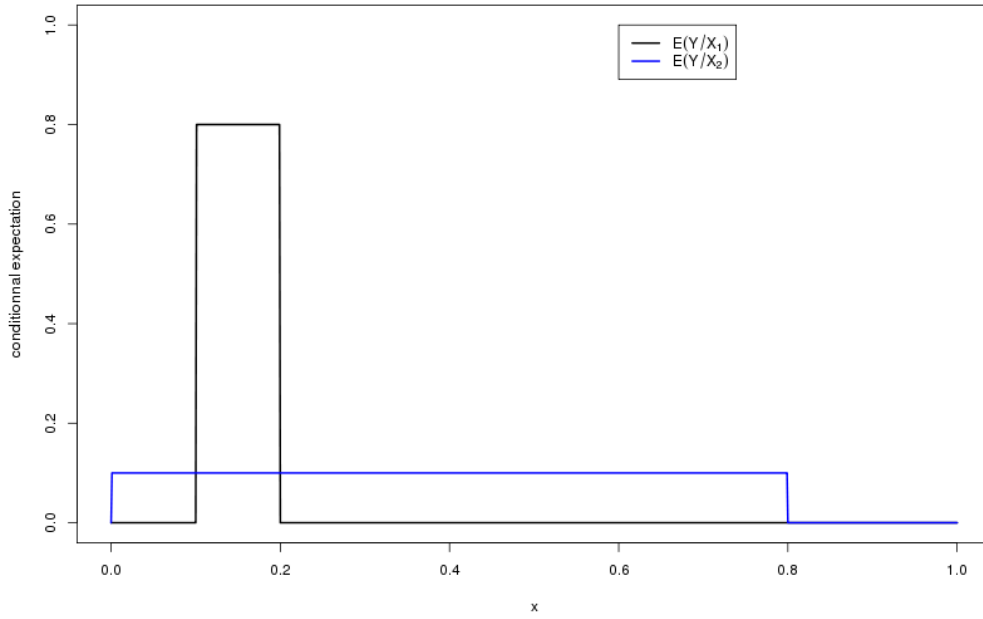


Figure 1.4: Conditional expectations for 2 variables

This figure provides information on the local features of the considered quantity. The (exact) Sobol' indices appear in table 1.1.

Variable or group	X_1	X_2	X_1 and X_2	Total eff. of X_1	Total eff. of X_2
Sobol indice	$S_1 = 0.783$	$S_2 = 0.022$	$S_{12} = 0.196$	$S_{T1} = 0.979$	$S_{T2} = 0.218$

Table 1.1: Sobol indices for the first failure rectangle

The values of the index reads as follows: X_1 explains on its own 78% of the output variance, while X_2 explains only 2%. The total effects confirm that X_1 is of first importance (98% of the output variance explained), and show that X_2 has a medium impact (22% of the output variance explained).

These values appear to be consistent with figure 1.4 and with the expression of the failure function. Indeed, the first order indices are the variance of the conditional expectations. The black curve associated to variable X_1 varies on its support with more amplitude than the blue curve associated to variable X_2 . It seems consistent to have an index S_1 superior to S_2 . Similarly, when looking at the expression of the failure function, one sees that variable X_1 impacts the failure probability on a small fraction of its support. On the opposite, variable X_2 impacts the failure probability on a broader fraction of its support. The information gained by the knowledge of the first variable value is then larger than the one gained by the knowledge of the second variable value. This toy example

draws attention to the relatively high value of the index associated to the interaction between the two variables (around 20%). This interaction is important: to get a failure event, both variables need to have a critical value jointly.

For the second example, the failure function has expression:

$$1_{G(\mathbf{X}) \leq k} = 1_{\{0,15 < X_1 < 0,2\} \{0,4 < X_2 < 0,8\}}(\mathbf{X})$$

where $\mathbf{X} = (X_1, X_2)$ with $X_1, X_2 \sim \mathcal{U}[0, 1]$. The failure probability is $P_f = 0.02$. Sobol' indices appear in table 1.2.

Variable or group	X_1	X_2	X_1 and X_2	Total eff. of X_1	Total eff. of X_2
Sobol indice	$S_1 = 0.388$	$S_2 = 0.031$	$S_{12} = 0.582$	$S_{T1} = 0.970$	$S_{T2} = 0.613$

Table 1.2: Sobol indices for the second failure rectangle

It can be seen that the impact of the interaction is much larger (58% of the share of variance), despite the similarity of the failure function. The total effects show that both variables are important.

On the failure hypercubes More generally, one can show that for a d -dimensional failure hypercube where the inputs are independent uniforms that:

- Sobol' indices associated to a variable decays with the width of its associated failure indicator.
- The indices corresponding to interactions grow as the failure probability diminishes.
- A variable has interaction effect with all the others, unless its associated failure indicator is as wide as the support of the variable. In this last case, the first order index associated with this variable is null.

This basic example shows how Sobol' indices can be used to rank the impact on the failure probability, using the total effects rather than the first order effects. Based on this conclusion, we will pursue the study of Sobol' indices applied to a failure indicator.

1.4.2 Computational methods

The following sections are dedicated to several estimation techniques of the Sobol' indices. To do so, consistent estimators of the following quantities are required:

- $\text{Var}(Y)$,
- $D_i(Y) = \text{Var}[\mathbb{E}(Y|X_i)]$,
- $D_{ij}(Y) = \text{Var}[\mathbb{E}(Y|X_i, X_j)] - D_i(Y) - D_j(Y)$,
- and so on.

The organization is the following: first we will present the techniques based upon MC sampling, namely Sobol'; Saltelli, Mauntz, Jansen and Janon-Monod. Secondly, the techniques based upon Fourier transformation -namely FAST, E-FAST, RBD- will be presented.

1.4.2.1 MC based estimation techniques

Sobol' - Presentation of the method This method is presented in the founder article by Sobol' [92]. Let us denote G the d -dimensional model. Sobol' method principle is the following: consider two independent matrices of N realisations of the vector of d inputs; representing two sets of inputs. In those matrices, a realisation of the d inputs is figured linewise. Those matrices are the following:

$$\xi_1 = \begin{pmatrix} X_{1,1}^{(1)} & X_{1,2}^{(1)} & \cdots & X_{1,d}^{(1)} \\ X_{2,1}^{(1)} & X_{2,2}^{(1)} & \cdots & X_{2,d}^{(1)} \\ \vdots & \vdots & \ddots & \vdots \\ X_{N,1}^{(1)} & X_{N,2}^{(1)} & \cdots & X_{N,d}^{(1)} \end{pmatrix} \quad \text{and} \quad \xi_2 = \begin{pmatrix} X_{1,1}^{(2)} & X_{1,2}^{(2)} & \cdots & X_{1,d}^{(2)} \\ X_{2,1}^{(2)} & X_{2,2}^{(2)} & \cdots & X_{2,d}^{(2)} \\ \vdots & \vdots & \ddots & \vdots \\ X_{N,1}^{(2)} & X_{N,2}^{(2)} & \cdots & X_{N,d}^{(2)} \end{pmatrix} \quad (1.60)$$

In Sobol' method, the mean of the output Y is estimated by:

$$\hat{D}_0 = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \quad (1.61)$$

Conversely, the variance of the output is computed as follows:

$$\hat{D} = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)})^2 - \hat{D}_0^2 \quad (1.62)$$

To compute the D_i quantities, the two data sets are considered, yet one column (i.e. i -th input) in the second data-set is replaced by the corresponding values of the first data-set. This writes:

$$\hat{D}_i = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \times G(X_{k,1}^{(2)}, \dots, X_{k,i-1}^{(2)}, X_{k,i}^{(1)}, X_{k,i+1}^{(2)}, \dots, X_{k,d}^{(2)}) - \hat{D}_0^2 \quad (1.63)$$

In the same order of ideas, the quantities D_{ij} are estimated by "fixing" two columns of the second matrix to the corresponding values of the first matrix. This writes:

$$\hat{D}_{ij} = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \times G(X_{k,1}^{(2)}, \dots, X_{k,i}^{(1)}, X_{k,i+1}^{(2)}, \dots, X_{k,j}^{(1)}, X_{k,j+1}^{(2)}, \dots, X_{k,d}^{(2)}) - \hat{D}_i - \hat{D}_j - \hat{D}_0^2 \quad (1.64)$$

Thus an estimation of the first, second,... order Sobol' indices can be made:

$$\hat{S}_i = \frac{\hat{D}_i}{\hat{D}}, \quad \hat{S}_{ij} = \frac{\hat{D}_{ij}}{\hat{D}} \quad (1.65)$$

and so on. Thus the total indices S_{T_i} can be estimated by summing all the indices containing i . However, this technique has a prohibitive cost: to get all the first order sensitivity indices, one must perform $N \times (d+1)$ function calls. To get all the indices (thus estimate the total indices) one must perform $N \times (2^d)$ function calls. Additionally, this method is known for needing a fair N to get precise estimations, of order 10000 to get a 10% error on the indices, much more for low value indices.

Saltelli - Presentation of the method Saltelli [87] proposed an efficient method to compute the sensitivity indices. This method is popular within the engineering fields since it allows estimation for each input the first and total order indices, for a smaller cost than the Sobol' method.

The estimation of the quantities $D_i, D_{ij}, \dots$ are realised in the same way as in the Sobol' method. The total indices are estimated as follows: consider the quantity $D_{\sim i}$ defined as the total share of variance that does not come from variable X_i . Then the total indices rewrite:

$$S_{T_i} = 1 - \frac{D_{\sim i}}{\text{Var}(Y)} \quad (1.66)$$

Thus total sensitivity indices are computed by estimating:

$$\hat{D}_{\sim i} = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \times G(X_{k,1}^{(1)}, \dots, X_{k,i-1}^{(1)}, X_{k,i}^{(2)}, X_{k,i+1}^{(1)}, \dots, X_{k,d}^{(1)}) - \hat{D}_0^2 \quad (1.67)$$

To minimize the number of function calls, the estimation of D_i is made as in Sobol' method, but switching the samples:

$$\hat{D}_i = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(2)}, \dots, X_{k,d}^{(2)}) \times G(X_{k,1}^{(1)}, \dots, X_{k,i-1}^{(1)}, X_{k,i}^{(2)}, X_{k,i+1}^{(1)}, \dots, X_{k,d}^{(1)}) - \hat{D}_0^2 \quad (1.68)$$

The number of function calls to estimate the first-order and totals sensitivity indices is $N \times (d+2)$

Mauntz - Presentation of the method In order to improve the estimation of indices S_i with small values, Mauntz (Sobol' et al. [94]) proposed an estimator of D_i that writes:

$$\hat{D}_i = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(2)}, \dots, X_{k,d}^{(2)}) \times \left[G(X_{k,1}^{(1)}, \dots, X_{k,i-1}^{(1)}, X_{k,i}^{(2)}, X_{k,i+1}^{(1)}, \dots, X_{k,d}^{(1)}) - G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \right] \quad (1.69)$$

and the numerator of S_{T_i} writes:

$$\text{Var}(Y) - \hat{D}_{\sim i} = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \times \left[G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) - G(X_{k,1}^{(1)}, \dots, X_{k,i}^{(2)}, \dots, X_{k,d}^{(1)}) \right] \quad (1.70)$$

For the indices close to 0, one or two decades are gained on the indices' uncertainty. The number of function calls for the method of Mauntz (first-order and totals sensitivity indices) is $N \times (d+2)$.

Jansen - Presentation of the method Jansen [54] proposed alternative estimators for S_i and S_{T_i} .

$$\hat{D}_i = \text{Var}(Y) - \frac{1}{2N} \sum_{k=1}^N \left[G(X_{k,1}^{(2)}, \dots, X_{k,d}^{(2)}) - G(X_{k,1}^{(1)}, \dots, X_{k,i}^{(2)}, \dots, X_{k,d}^{(1)}) \right]^2 \quad (1.71)$$

and the numerator of S_{T_i} writes:

$$\text{Var}(Y) - \hat{D}_{\sim i} = \frac{1}{2N} \sum_{k=1}^N \left[G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) - G(X_{k,1}^{(1)}, \dots, X_{k,i}^{(2)}, \dots, X_{k,d}^{(1)}) \right]^2 \quad (1.72)$$

The number of function calls for the Jansen's method (first-order and totals sensitivity indices) is $N \times (d+2)$.

Janon-Monod - Presentation of the method In order to improve the estimation of the first-order indices in Sobol' method, Monod *et al.* [70] have proposed new estimators for the sensitivity indices. Janon *et al.* [53] proved the asymptotic efficiency of these estimators.

$$\hat{D}_i = \frac{1}{N} \sum_{k=1}^N G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) \times G(X_{k,1}^{(2)}, \dots, X_{k,i}^{(1)}, \dots, X_{k,d}^{(2)}) - \left[\frac{1}{N} \sum_{k=1}^N \frac{G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) + G(X_{k,1}^{(2)}, \dots, X_{k,i}^{(1)}, \dots, X_{k,d}^{(2)})}{2} \right]^2 \quad (1.73)$$

The estimator of the variance of Y ($\hat{D}$) reads:

$$\hat{D} = \frac{1}{N} \sum_{k=1}^N \left[\frac{[G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)})]^2 + [G(X_{k,1}^{(2)}, \dots, X_{k,i}^{(1)}, \dots, X_{k,d}^{(2)})]^2}{2} \right] - \left[\frac{1}{N} \sum_{k=1}^N \frac{G(X_{k,1}^{(1)}, \dots, X_{k,d}^{(1)}) + G(X_{k,1}^{(2)}, \dots, X_{k,i}^{(1)}, \dots, X_{k,d}^{(2)})}{2} \right]^2 \quad (1.74)$$

Reliability case The estimation methods based upon the principles of MC estimation will present the drawbacks of such methods. Practically, the small failure probability implies that the simulation sets will include few failure points. The estimation of the indices will be imprecise at best, impossible in the worst case (no failure point in the data set). Tests provided in section 1.4.3 (where a large data set is needed) will confirm these reflections.

1.4.2.2 Fourier analysis based techniques

Presentation of the methods The *Fourier Amplitude Sensivity Test* (FAST) method was first presented by Cukier *et al.* [27]; and is based upon a Fourier transformation. It allows an estimation of the indices at a smaller cost than the Sobol' method. Saltelli *et al.* [90] extended this method for the estimation of total indices, thus giving the Extended-FAST (E-FAST) method.

Classical FAST method is based on a selection of N points (i.e. sampling) on a specific curve constructed in such a way that it explores each dimension (associated to an input variable) with a preset frequency (different for each input). Let us assume that the input domain is the unit hypercube. The curve is then defined by:

$$x_i(s) = G_i(\sin \omega_i s), \forall i = 1, \dots, d$$

where s is a scalar such that $-\infty < s < \infty$. G_i is a function from $[-1 : 1]$ to $[0, 1]$ and defines the search-curve - it is not related to the numerical model G . ω_i the frequency associated to the i -th input.

Based on the approximation of Weyl's theorem ([101]); one has, for any d-dimensional function f and for the $x_i(s)$ defined as previously:

$$\int_{[0,1]^d} G(x) dx \approx \frac{1}{2\pi} \int G(x(s)) ds \quad (1.75)$$

where $x(s) = (x_1(s), \dots, x_d(s))$. Equation (1.75) is only true when the frequencies are linearly independent. This cannot be the case in practice. Therefore the algorithm requires that the practitioner sets a maximal interaction order M and selects the frequencies free of interferences up to M .

The function is then computed on each of the N points, then a Fourier decomposition is performed on the sample to estimate its spectrum. Decomposing the spectrum with respect to the frequencies allows to estimate the estimators of the parts of variance. Indeed, denoting A_j and B_j the following Fourier coefficients:

$$A_j = \frac{1}{2\pi} \int_{-\pi}^{\pi} G(x(s)) \cos(js) ds$$

$$B_j = \frac{1}{2\pi} \int_{-\pi}^{\pi} G(x(s)) \sin(js) ds$$

Main results from Cuckier *et al.* [27] is that

$$\text{Var}[Y] \approx 2 \sum_{k=1}^{+\infty} (A_k^2 + B_k^2) \quad (1.76)$$

$$D_i = \text{Var}[\mathbb{E}(Y/X_i)] \approx 2 \sum_{k=1}^{+\infty} (A_{k\omega_i}^2 + B_{k\omega_i}^2) \quad (1.77)$$

The complexity of such an algorithm comes from the way to generate the sampling curve, that needs to explore each dimension with preset frequencies avoiding interactions.

Random Balance Design (RBD) method, proposed by Tarantola *et al.* [96] is a modification of the FAST technique. The algorithm starts exploring the input space via a search curve, but unlike in FAST, each dimension is explored with the same frequency. Then a random permutation of the coordinates of the sample points is performed. The function is called on each point of the new sample, then the Fourier decomposition is carried out for the sampling frequency and its harmonics, up to order M of supposed maximal interaction order. This allows an estimation of the indices associated to each input. Tissot *et al.* [97] proposed a way to correct the biais produced in such estimates.

Reliability case It can be expected that the FAST/E-FAST/RBD methods will not perform well in the reliability case. Indeed, the indices cannot be computed easily on a discontinuous function, especially on the indicator of a small set. Numerical tests have shown that a correct estimation of the indices for a discontinuous function is possible, provided a high maximal interaction order M is selected. Unfortunately, increasing this order leads to frequency selection problems. Therefore the FAST and derived methods will not be tested in the following.

1.4.3 Reliability test cases

This applicative subsection have the following objectives:

- The first objective is to check the consistency of the estimators, to verify that the estimator of the indices converges to the true value as the sample size grows. This will be performed on test cases for which one can easily compute or approximate closely the indices.
- Another objective is to perform the sensitivity analysis on the numerical examples defined in Appendix B.

1.4.3.1 Numerical results: convergence to the true value

In this part, we will focus on the hyperplane test case, described in Appendix B.1. Let us first remind the formulation of the first order Sobol' indices:

$$S_i = \frac{\text{Var}(\mathbb{E}[Y|X_i])}{\text{Var}(Y)} = \frac{D_i}{\text{Var}(Y)}.$$

In the reliability case, the expression of $\text{Var}(Y)$ is straightforward:

$$\begin{aligned} \text{Var}(Y) &= \mathbb{E}(Y^2) - \mathbb{E}(Y)^2 \\ &= \mathbb{E}(1_{G(\mathbf{X}) \leq 0}^2) - \mathbb{E}(1_{G(\mathbf{X}) \leq 0})^2 \\ &= P_f(1 - P_f). \end{aligned}$$

In the hyperplane case, the failure probability is known: it equals $P = \phi \left(-k / \sqrt{\sum_{i=1}^d a_i^2} \right)$. This allows an exact computation of the variance. Let us denote $T_{X_i}(x) = \mathbb{E}[Y|X_i = x]$, the function depending solely on X_i that explains best the output Y . In the hyperplane case with Gaussian inputs, one has:

$$T_{X_i}(x) = \mathbb{E}[Y|X_i = x] = \mathbb{P} \left(\sum_{j=1; j \neq i}^d a_j X_j \leq k - a_i x \right) = \phi \left(\frac{k - a_i x}{\sqrt{\sum_{j=1; j \neq i}^d a_j^2}} \right). \quad (1.78)$$

Then by definition:

$$D_i = \mathbb{E}[T_{X_i}^2] - \mathbb{E}[T_{X_i}]^2$$

with:

$$\mathbb{E}[T_{X_i}] = \int_{\mathbb{R}} T_{X_i}(x) f_{X_i}(x) dx = P_f$$

and $f_{X_i}(x)$ is the pdf of a standard Gaussian. In the same way,

$$\mathbb{E}[T_{X_i}^2] = \int_{\mathbb{R}} T_{X_i}^2(x) f_{X_i}(x) dx.$$

The last mono-dimensional integral does not have a simple expression, but one can estimate it using the quadrature method. This, associated with the exact knowledge of $\text{Var}(Y)$ allows to get precise estimations of first order Sobol' indices. This estimation will be used to control the quality of the estimations.

Let us verify for the hyperplane 6410 case (described in Appendix B.1), where $\mathbf{a} = (1, -6, 4, 0)$ that the estimations of the indices converge to the “real” values of the indices. First, we estimate the “real” indices with the procedure described above, and the results are displayed in table 1.3.

Variable	X_1	X_2	X_3	X_4
Indice S_i	0.002	0.259	0.055	0

Table 1.3: First order Sobol' indices for the hyperplane 6410 case

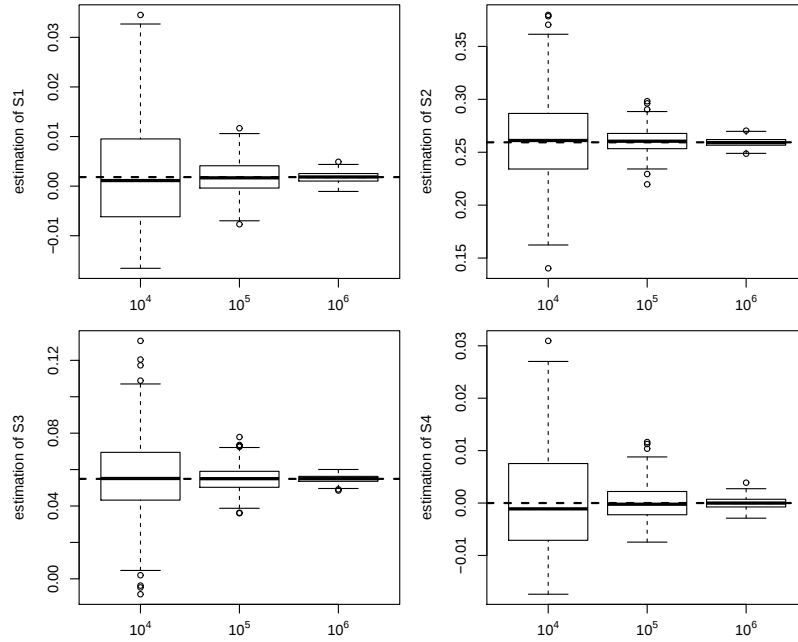


Figure 1.5: Boxplots of the estimated first order Sobol' indices with the Sobol' method

We repeat the following operation 500 times: generating two samples of size N , with N varying from 10^4 to 10^6 , and estimating the indices on all these samples. The results of the estimation of the first order Sobol' indices are shown in figure 1.5 for the Sobol' method and in figure 1.6 for the Saltelli method.

The graphics show that the estimator converges to the true value when the sample size increases. Additionally, it shows that the estimations of a null index (S_4) with the Sobol' method can provide results with a wider spread than the ones provided with the Saltelli method. For this reason, we will use the Saltelli method in the following. Concerning the good sample size to correctly estimate the Sobol' indices, the results show that obviously the larger the sample is, the better the estimation is. For our test cases, we will use samples of size 10^6 , since our toy-models are not costly. However it should be noticed that this number of function calls might be unrealistic for real models.

1.4.3.2 Hyperplane 6410 case

We present on table 1.4 the estimated Sobol' indices with 2 samples of size 10^6 , using the Saltelli method. The total number of function evaluations is 6×10^6 .

Index	S_1	S_2	S_3	S_4	S_{T1}	S_{T2}	S_{T3}	S_{T4}
Estimation	0.002	0.254	0.054	0	0.200	0.940	0.720	0

Table 1.4: Estimated Sobol' indices for the hyperplane 6410 case

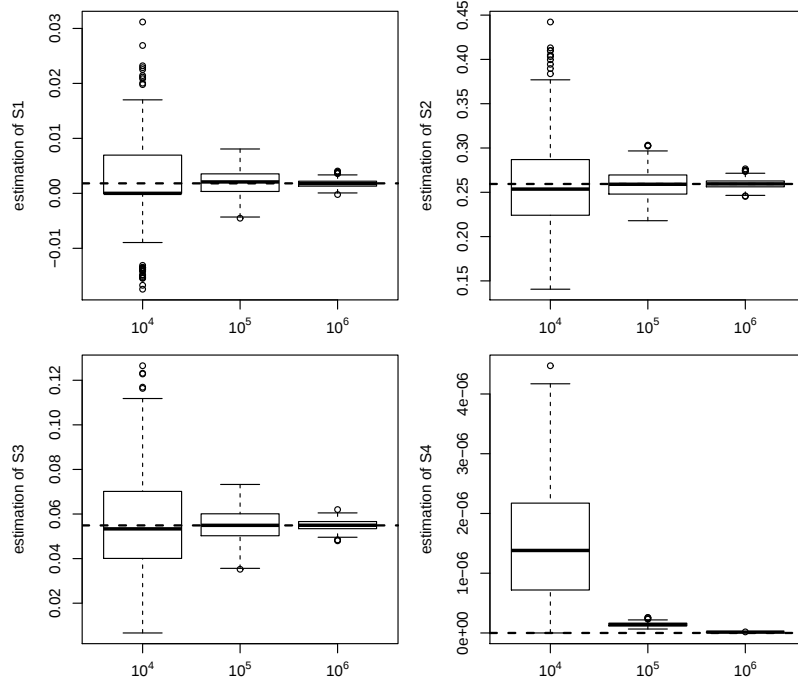


Figure 1.6: Boxplots of the estimated first order Sobol' indices with the Saltelli method

The total indices assess that X_2 is extremely influential, and that X_3 is highly influential. X_1 has a moderate influence and X_4 has a null influence. This last point is interesting: it shows that this SA method can detect the non-influential variables.

1.4.3.3 Hyperplane 11111 case

This numerical example has been described in Appendix B.1. We present on table 1.5 the estimated Sobol' indices with 2 samples of size 10^6 , using the Saltelli method. The total number of function evaluations is 7×10^6 .

Index	S_1	S_2	S_3	S_4	S_5	S_{T1}	S_{T2}	S_{T3}	S_{T4}	S_{T5}
Estimation	0.015	0.013	0.014	0.009	0.015	0.677	0.673	0.695	0.674	0.685

Table 1.5: Estimated Sobol' indices for the hyperplane 11111 case

The weak first order indices (less than 2% of the variance explained) and the high total indices assess that all the variables are influential in interaction with the others. All the total indices are approximatively the same showing that this SA method can give the same importance to each equally contributing input.

1.4.3.4 Hyperplane 15 variables case

This numerical example has been described in Appendix B.1. We present on table 1.6 the estimated Sobol' indices with 2 samples of size 10^6 , using the Saltelli method. The total number of function evaluations is 17×10^6 .

Index	S_1 to S_5	S_6 to S_{10}	S_{11} to S_{15}
Estimation	0.014 to 0.018	0.001 to 0.002	0
Total Index	S_{T1} to S_{T5}	S_{T6} to S_{T10}	S_{T11} to S_{T15}
Estimation	0.655 to 0.673	0.141 to 0.150	0

Table 1.6: Estimated Sobol' indices for the hyperplane 15 variables case

The first order indices are all weak, yet separated in three groups. The total indices give a good separation between the influential, weakly influential and non influential variables. The Sobol' indices SA method is able to deal with problems of medium dimension; however it has an heavy computational cost in this case.

1.4.3.5 Hyperplane with same importance and different spreads

This numerical example has been described in Appendix B.1. We present on table 1.7 the estimated Sobol' with 2 samples of size 10^6 , using the Saltelli method. The total number of function evaluations is 7×10^6 .

Index	S_1	S_2	S_3	S_4	S_5	S_{T1}	S_{T2}	S_{T3}	S_{T4}	S_{T5}
Estimation	0.027	0.028	0.025	0.025	0.028	0.611	0.622	0.618	0.618	0.624

Table 1.7: Estimated Sobol' indices for the hyperplane "different spreads" case

The weak first order indices (less than 3% of the variance explained) and the high total indices assess that all variables are influential in interaction with the others, and that no variable is influential on its own. All the total indices are approximatively equal showing that this SA gives to each equally contributing variable the same importance, despite their different spread.

1.4.3.6 Thresholded Ishigami function

We use the example defined in Appendix B.2, the thresholded Ishigami function. The estimated Sobol' with 2 samples of size 10^6 , using the Saltelli method, are given in table 1.8. The total number of function evaluations is 5×10^6 .

Index	S_1	S_2	S_3	S_{T1}	S_{T2}	S_{T3}
Estimation	0.018	0.007	0.072	0.831	0.670	0.919

Table 1.8: Sobol' indices estimation for the thresholded Ishigami function

The first order indices are close to 0. The variable with the most influence on its own is X_3 , explaining 7% of the output variance. Total indices state that all the variable are of high influence. A variable ranking can be made using the total indices, ranking X_3 with the highest influence, then X_1 and then X_2 . Figure B.1 allows to understand the meaning of the total indices. Each variable "causes" the failure event on a restricted portion of its support. On the other hand, the knowledge of a single variable does not allow to explain the variance of the indicator, thus the weak first-order indices. The fact that the failure points are grouped in narrow strips can only be explained by the 3 variables together, thus the high 3-order index.

1.4.3.7 Flood case

This test case has been described in Appendix B.3. The estimated Sobol' with 2 samples of size 10^6 , using the Saltelli method, are given in table 1.9. The total number of function evaluations is 6×10^6 .

Index	S_Q	S_{K_s}	S_{Z_v}	S_{Z_m}	S_{TQ}	S_{TK_s}	S_{TZ_v}	S_{TZ_m}
Estimation	0.019	0.251	0	0	0.746	0.976	0.248	0.115

Table 1.9: Estimated Sobol' indices for the flood case

Most first order indices are small, except the one associated to K_s that explains 25% of the variance on its own. The total indices state that K_s and Q are extremely influential, Z_v is influential and Z_m is little influential. One can see that S_{TZ_v} and S_{TZ_m} differ from 0, meaning these variables have an impact on the failure probability when interacting with other variables.

1.4.3.8 Conclusion

In most tested cases, Sobol' indices allow distinguishing the influential and the non-influential variables. However, their evaluation is costly. The objective of the two next subsections is to study methods that allow a reduction of function calls.

1.4.4 Reducing the number of function calls: use of QMC methods

This subsection focuses on the use of Quasi Monte-Carlo methods (presented in section 1.2.1.2) to estimate Sobol' indices. This technique is presented in Sobol' [93].

1.4.4.1 Estimation of Sobol' indices through QMC

The main idea when using pseudo-random sequences is to use the estimators presented in section 1.4.2.1, replacing the random samples by samples coming from a low-discrepancy sequence. In the following, Sobol' sequence is used (see Niederreiter [75]).

When estimating the indices with the Sobol' method, 2 samples of size N and of dimension d i.i.d. to $\mathbf{X}$ are generated. These samples are then separated in complementary sets. A generation of two samples from the pseudo-random sequence is meaningless, since it is a deterministic sequence. The trick is to generate a sample of size N and of dimension $2d$, then to split this sample. Such a separation allows to get two samples of dimension d . Sobol' sequence produces orthogonal columns, these pseudo-random samples can be considered as independent. As an example on the pseudo-random sample generation, table 1.10 displays the 8 first points generated by Sobol' sequence in dimension 4.

1.4.4.2 Illustration on the hyperplane test case

In this part, the focus will be set on the hyperplane 6410 test case, described in Appendix B.1. The aim of this part is to assess the capability of QMC sampling to get a good estimation of Sobol' indices at a smaller computational cost.

First, two QMC samples of size 10^4 and of dimension 4 are generated (using the trick given above). The same is done for size 10^5 . Let us notice that the sample of size 10^4 is included in the one of size 10^5 , due to the determinism of the Sobol' sequence. Then the Sobol' indices are estimated

V_1	V_2	V_3	V_4
0,5	0,5	0,5	0,5
0,75	0,25	0,75	0,25
0,25	0,75	0,25	0,75
0,375	0,375	0,625	0,125
0,875	0,875	0,125	0,6250
0,625	0,125	0,375	0,375
0,125	0,625	0,875	0,875
0,1875	0,3125	0,3125	0,6875

Table 1.10: 4 dimensional points generated through Sobol' sequence

on the samples, using resp. 6×10^5 and 6×10^4 function calls. The results are displayed in table 1.11 and compared to a large sample size MC.

Index	S_1	S_2	S_3	S_4	S_{T1}	S_{T2}	S_{T3}	S_{T4}
MC, size 10^6	0.002	0.254	0.054	0	0.200	0.940	0.720	0
QMC, size 10^4	0.007	0.270	0.051	0	0.175	0.934	0.730	0
QMC, size 10^5	0.002	0.266	0.059	0	0.195	0.944	0.720	0

Table 1.11: Estimation of Sobol indices using QMC for the 6410 hyperplane test case

From these results, we conclude that the use of QMC for sampling allows to gain a factor 10 in the number of function calls. Indeed, one can see that the estimation with 10^4 QMC points is less accurate than the estimation with 10^5 QMC points, assuming the “true” values are the ones obtained with a MC sample of size 10^6 . Despite this loss of precision, the variable ranking is not changed when using a “small” QMC sample.

1.4.4.3 Conclusion on using QMC sampling to estimate Sobol' indices

This method as presented here does not provide an estimation of the error made, due to the determinism of the sampling. However, scrambling techniques have been developed (Jakubowicz *et al.* [52]) to add randomness in the sampling, thus allowing the computation of confidence intervals. This might be an avenue for future researches. As a conclusion on the use of QMC sampling to estimate Sobol' indices, this method might be used to identify the non influential variables at a smaller computational cost.

1.4.5 Reducing the number of function calls : use of importance sampling methods

The main idea in this part is to use importance sampling methods to estimate the Sobol' indices. This is the same as to run the simulations with a modified sampling density, then weight the estimations to take this density into account. Importance sampling is not used when estimating Sobol' indices for a continuous variable, there is no sense in fostering sampling in a particular zone. But it makes some sense in the reliability case: we want to obtain more failure samples. The numerical simulations presented in this section shows that this technique is effective if the sampling density is well chosen. To the best of our knowledge, this is an original contribution.

1.4.5.1 Rewriting the estimators with an importance density

The estimators of the Sobol' method (presented in section 1.4.2.1) are used. The aim is to estimate the index associated to variables $X_{i_1}, \dots, X_{i_s}$. The set of inputs $X_1, \dots, X_d$ is separated, like in the Sobol' method, into two data sets, of respective sizes s and $d - s$. Let us denote these data sets v and t , where v includes the inputs of interest $X_{i_1}, \dots, X_{i_s}$. Inputs are independent, therefore we can rewrite the input density as a product of two margins:

$$f_{\mathbf{X}}(\mathbf{x}) = f_v(v)f_t(t).$$

Two sets of N points are sampled with density $f_{\tilde{\mathbf{X}}}$, chosen by the practitioner. Each is separated into two data sets, $(\tilde{v}_1, \tilde{t}_1)$ $(\tilde{v}_2, \tilde{t}_2)$. The estimators in the reliability case writes:

$$\widehat{G_{0TI}} = \frac{1}{N} \sum_{j=1}^N 1_{G(\tilde{v}_{1j}, \tilde{t}_{1j}) < 0} \frac{f_{\mathbf{X}}(\tilde{v}_{1j}, \tilde{t}_{1j})}{f_{\tilde{\mathbf{X}}}(\tilde{v}_{1j}, \tilde{t}_{1j})} \quad (1.79)$$

$$\widehat{D_{TI}} = \widehat{G_{0TI}} - \widehat{G_{0TI}}^2 \quad (1.80)$$

$$\widehat{D_1 + G_{0TI}^2} = \frac{1}{N} \sum_{j=1}^N 1_{g(\tilde{v}_{1j}, \tilde{t}_{1j}) < 0} 1_{g(\tilde{v}_{1j}, \tilde{t}_{2j}) < 0} \frac{f_{\mathbf{X}}(\tilde{v}_{1j}, \tilde{t}_{2j})}{f_{\tilde{\mathbf{X}}}(\tilde{v}_{1j}, \tilde{t}_{2j})} \frac{f_t(\tilde{t}_{1j})}{f_{\tilde{t}}(\tilde{t}_{1j})} \quad (1.81)$$

$$\widehat{D_2 + G_{0TI}^2} = \frac{1}{N} \sum_{j=1}^N 1_{g(\tilde{v}_{1j}, \tilde{t}_{1j}) < 0} 1_{g(\tilde{v}_{2j}, \tilde{t}_{1j}) < 0} \frac{f_{\mathbf{X}}(\tilde{v}_{2j}, \tilde{t}_{1j})}{f_{\tilde{\mathbf{X}}}(\tilde{v}_{2j}, \tilde{t}_{1j})} \frac{f_v(\tilde{v}_{1j})}{f_{\tilde{v}}(\tilde{v}_{1j})} \quad (1.82)$$

1.4.5.2 Numerical applications

As a numerical test case, the hyperplane 6410 defined in Appendix B.1 is used. Let us first notice that the design point of such a failure surface has coordinates $u^* = (0.302, -1.811, 1.207, 0)$. The sampling density will thus consist in an independent Gaussian vector centred in the design point.

Let us then estimate, with samples size 10^4 the first order and total indices, with MC and with importance sampling. We repeat this estimation 100 times, the results are boxplotted in figure 1.7. The dashed lines represent the “theoretical” values obtained with a MC sample of size 10^6 .

One can see that the dispersion of the indices estimated with importance sampling is much smaller than the one associated with the indices estimated by MC.

The same procedure is applied with only 10^3 points and the results are displayed in figure 1.8.

The MC estimators are too dispersed to conclude anything, whereas the indices estimated with importance sampling are centred around the theoretical value.

1.4.5.3 Conclusion on using importance sampling to estimate Sobol' indices

Results are very good provided that the practitioner sets an adapted importance density. This might be much more complicated than in the example. For instance an adapted importance density might be hard to find for the thresholded Ishigami function.

1.4.6 Local polynomial estimation for first-order Sobol' indices in a reliability context

In the context of reliability analysis, we study the technique proposed by Da Veiga *et al.* [28] to deal with Sobol' indices estimation when inputs are correlated. The variance of the failure function

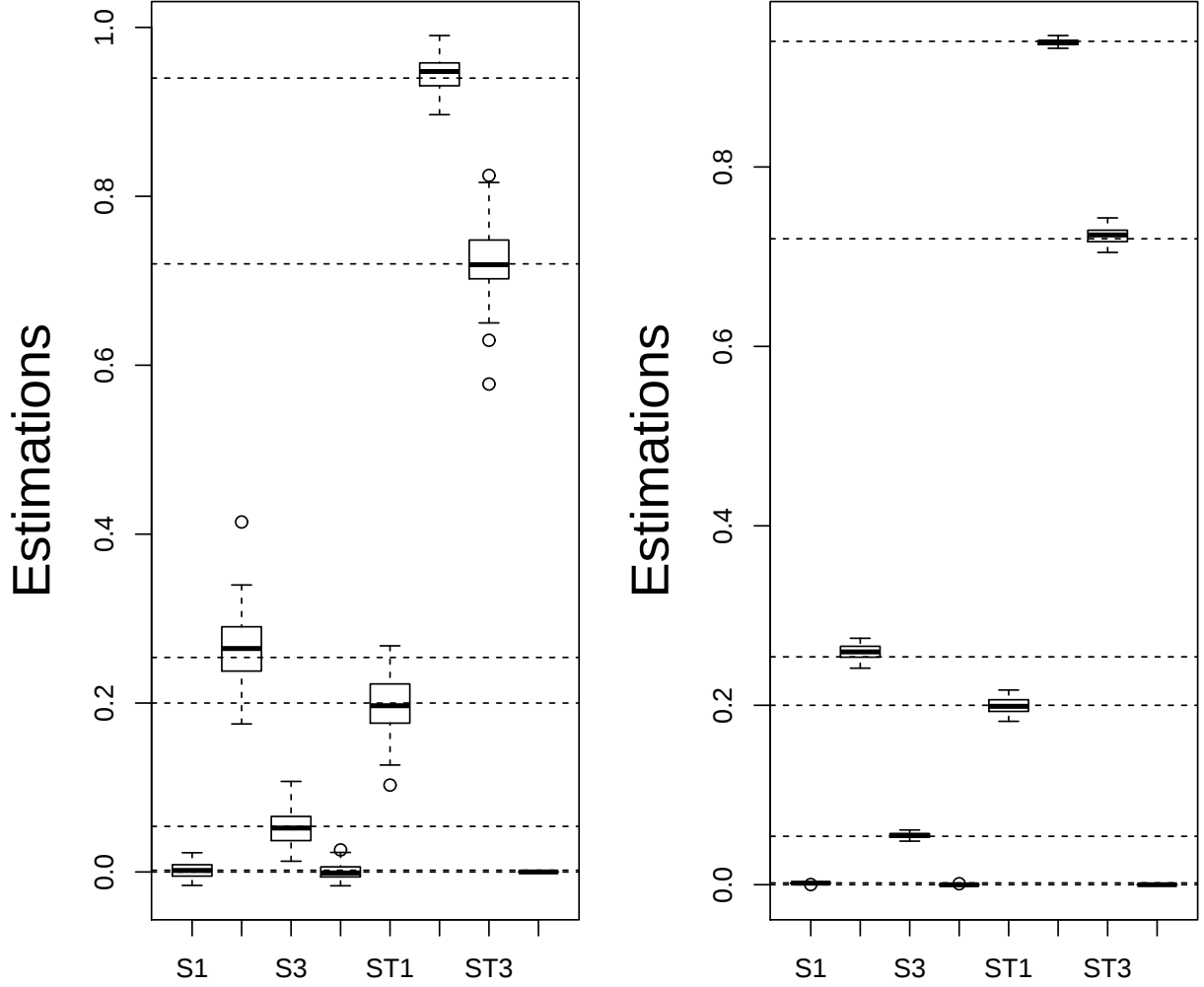


Figure 1.7: Comparison of first order and total indices, MC (left) and importance sampling (right), with 10^4 points for the hyperplane 6410 test case

within the failure domain is of interest here. The presented method is used to find out the first-order contribution of each variable to this variance. The question asked in this subsection is “*How each variable contributes to the variance of the failure function G within the failure domain?*”.

1.4.6.1 Sobol’ indices estimation by local polynomial smoothing

Let us recall that for a mathematical model denoted $G : \mathbb{R}^d \rightarrow \mathbb{R}$ with random inputs $\mathbf{X} \sim f$ and random output Y , first order Sobol’ indices are given by:

$$S_k = \frac{\text{Var}(\mathbb{E}(Y/X^k))}{\text{Var}(Y)}, \quad \forall k = 1, \dots, d. \quad (1.83)$$

In the case of independent inputs, one can quote Sobol’ and FAST estimation techniques, as presented in section 1.4.2. These methods cannot be applied when the inputs are no longer independent. Nevertheless there is a need for sensitivity analysis methods when inputs are non-independently distributed. Several recent works deal with this kind of problems.

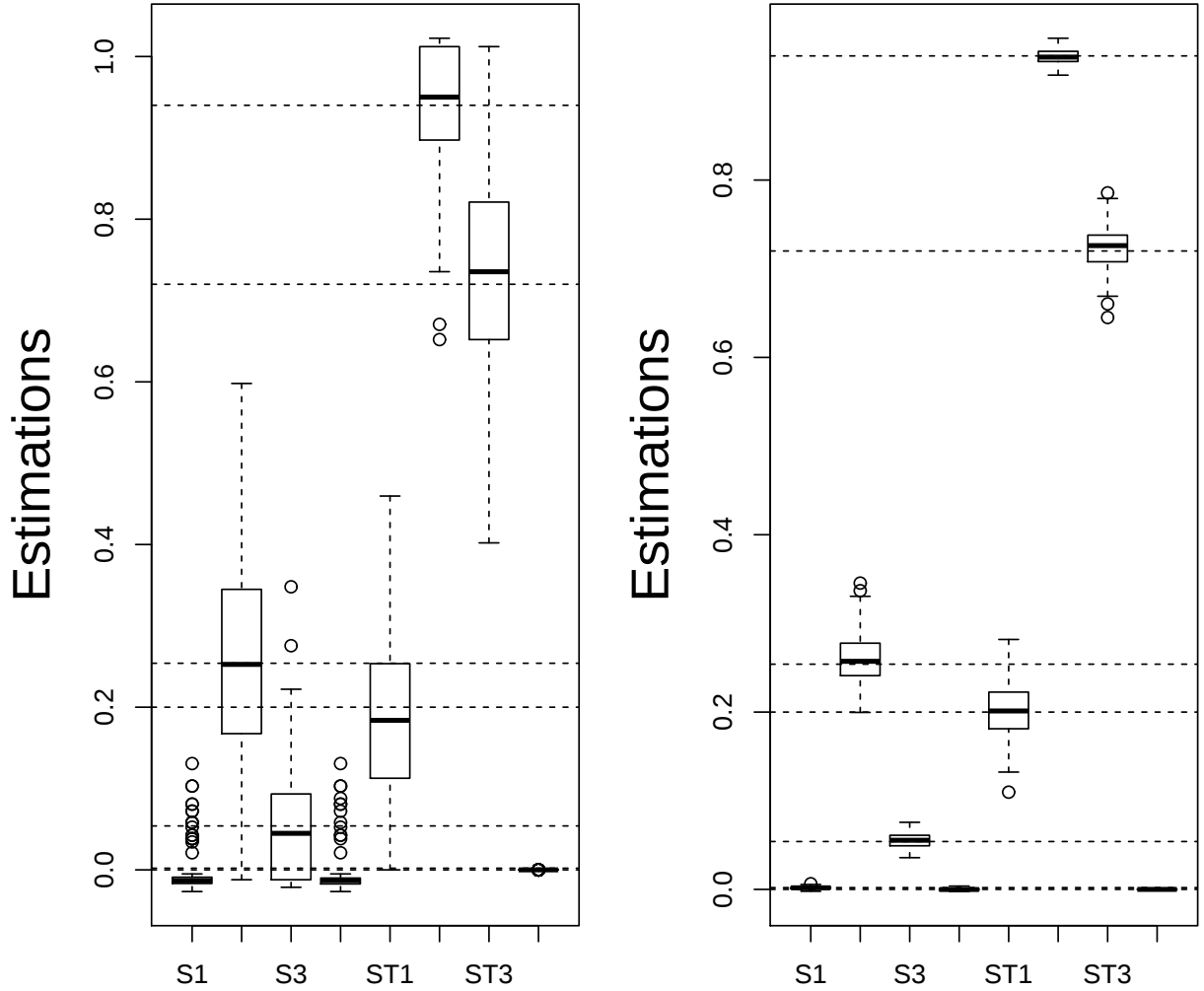


Figure 1.8: Comparison of first order and total indices, MC (left) and importance sampling (right), with 10^3 points for the hyperplane 6410 test case

The original technique proposed by Da Veiga *et al.* [28] to estimate S_k is based on local polynomial approximation of the conditional moments. More precisely the authors use a first sample $(\mathbf{X}_i, Y_i)_{i=1, \dots, N}$ to fit d local polynomial response surface to explain the following relationship for each given input k :

$$Y_i = m_k(X_i^k) + \sigma_k(X_i^k)\epsilon_i^k \quad (1.84)$$

where $m_k(x) = \mathbb{E}(Y/X^k = x)$ and $\sigma_k^2(x) = \text{Var}(Y/X^k = x)(x \in \mathbb{R})$. $\epsilon_i^k \forall i = 1, \dots, N$ are independent errors satisfying $\mathbb{E}(\epsilon_i^k/X^k) = 0$ and $\text{Var}(\epsilon_i^k/X^k) = 1$. The local polynomial (LP) smoothing provides estimators for $m_k(\cdot)$ and $\sigma_k^2(\cdot)$. Two formulations for Sobol' first order indices are given in the article, we choose to focus on the one involving $m_k(\cdot)$. Given another sample of i.i.d. inputs $(\tilde{\mathbf{X}}_i)_{i=1, \dots, N'}$ with same distribution as $\mathbf{X}$, one can use a plug-in estimation as follow. Denoting $\hat{m}(\cdot)$ the LP estimator of the conditional expectation, fitted on the first sample; denoting as well

$\bar{\hat{m}} = \frac{1}{N'} \sum_{i=1}^{N'} \hat{m}(\tilde{X}_i^k)$, one has:

$$\hat{T}_k = \frac{1}{N' - 1} \sum_{i=1}^{N'} \left(\hat{m}(\tilde{X}_i^k) - \bar{\hat{m}} \right)^2. \quad (1.85)$$

$\hat{T}_k$ is an empirical estimator of the variance of the expectation of Y given X^k . Dividing $\hat{T}_k$ by the estimated variance of Y , one has an estimator of S_k .

1.4.6.2 Reliability context

When dealing with the reliability context, the event $G(\mathbf{X}) < 0$ (system failure) and the complementary event $G(\mathbf{X}) \geq 0$ (system safe mode) are of interest. To quantify the impact of each input X^k on the failure probability $P = \int \mathbf{1}_{G(\mathbf{x}) < 0} f(\mathbf{x}) d\mathbf{x}$, we propose to study the first order Sobol' indices in the failure domain (FOSIFD).

It is obvious that given the failure event, the inputs in the failure domain are no longer independent. Thus the methodology proposed in Da Veiga *et al.* [28] is of interest here. It will be studied in the following part. One should be cautious with one point: sampling from the conditional joint distribution has a strong computational cost, since the second sample must be distributed as the first one; that is to say according to $f_{G(\mathbf{x}) < 0}(\mathbf{x}) = \frac{\mathbf{1}_{G(\mathbf{x}) < 0} f(\mathbf{x})}{P}$. This sampling operation can be performed by running new calls of the model G . Da Veiga *et al.* [28] propose two options in this case : splitting the original sample or performing a leave one out procedure. As our models are toy functions, our sample sizes can be large.

1.4.6.3 Hyperplane 6410 case

This numerical example has been described in Appendix B.1. We perform 100 runs of the following experiment: through simulation and function calls, we obtain two samples of size $N = N' = 10^6$. Only one out of a hundred of these points are of interest, since we study the FOSI in the failure domain. From the first sample failure points, we build a LP response surface and its mean is predicted through the second sample failure points. The variance of the expectation of the LP response surface is estimated and divided by the variance of the first sample failure points; as describe in section 1.4.6.1. The results are boxplotted in figure 1.9.

According to the first order sensitivity indices, the second variable contributes for 20% of the failure domain variance whereas the third variable contributes for 5% of the failure domain variance. The two other variables provide a negligible effect on their own. Since the inputs are no longer independent in the failure domain, one cannot assess that the sum of all the Sobol' indices is one. However in this case, we strongly suspect that most of the variance in the failure domain is caused by a higher-level interaction between variables.

1.4.6.4 Hyperplane 11111 case

This numerical example has been described in Appendix B.1. The aim of this example is to assess or infirm the capability of the FOSIFD to give to each equally contributing input the same importance. The results of the experiment with the same global parameters (100 runs, two samples of size 10^6) are boxplotted in figure 1.10.

The indices assess the same importance value for all the variables. However, one can see that each variable is said to contribute approximatively for 2% of the failure domain variance on its own.

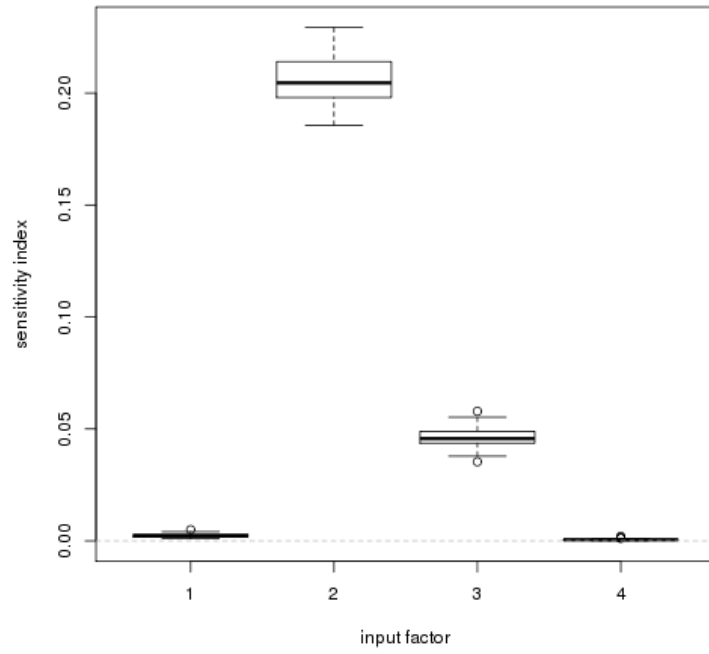


Figure 1.9: Boxplot of the estimated FOSIFD for the 6410 hyperplane case

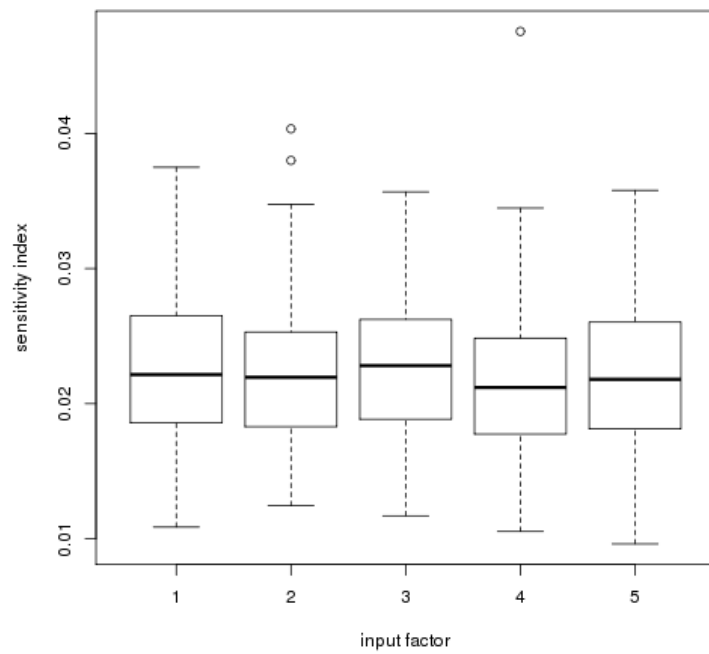


Figure 1.10: Boxplot of the estimated FOSIFD for the 11111 hyperplane case

Therefore, as in the previous case, we suspect that there is a higher-order interaction that causes most of the variance in the failure domain.

1.4.6.5 Hyperplane 15 variables case

This numerical example has been described in Appendix B.1. The results of the experiment with the same global parameters (100 runs, two samples of size 10^6) are boxplotted in figure 1.11.

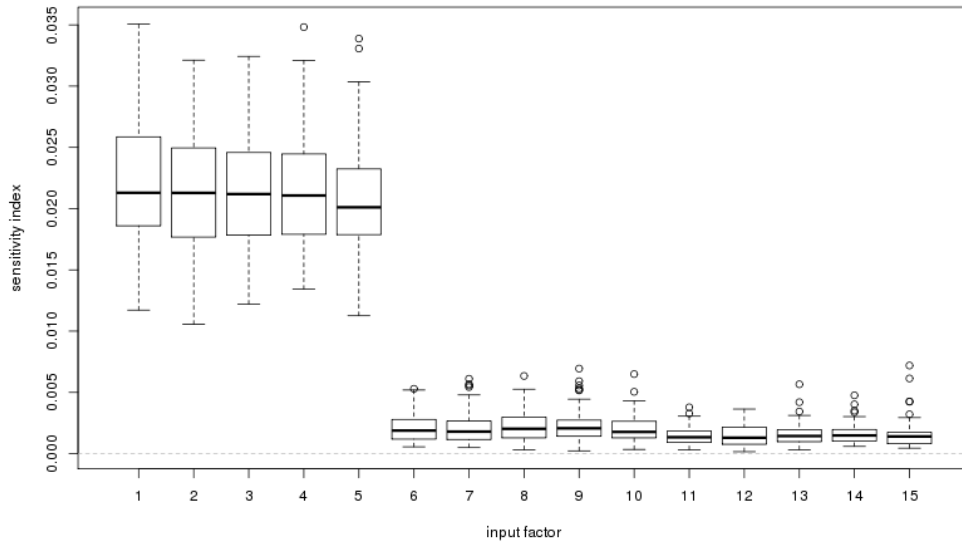


Figure 1.11: Boxplot of the estimated FOSIFD for the 15 variables hyperplane case

As one can see two groups of importance variables, one can conclude that the FOSIFD fails to separate variables with a low contribution and variables with a null contribution. However, the influential variables are detected and contribute for approximatively 2% of the failure domain variance.

1.4.6.6 Hyperplane with same importance and different spreads

This numerical example has been described in Appendix B.1. The aim of this test is to assess or infirm the capability of the FOSIFD to give to each equally contributing variable the same importance, despite their different spread. The results of the experiment with the same global parameters (100 runs, two samples of size 10^6) are boxplotted in figure 1.12.

One can see that the values of the FOSIFD are approximatively equal for each variable. Thus, each variable explain on its own 2% of the failure domain variance. These results are the same as in section 1.4.6.4. Thus one can think that the spread of the variable has no impact, at least on this test case.

1.4.6.7 Tresholded Ishigami function

This numerical example has been described in Appendix B.2. The results of the experiment with the same global parameters (100 runs, two samples of size 10^6) are boxplotted in figure 1.13.

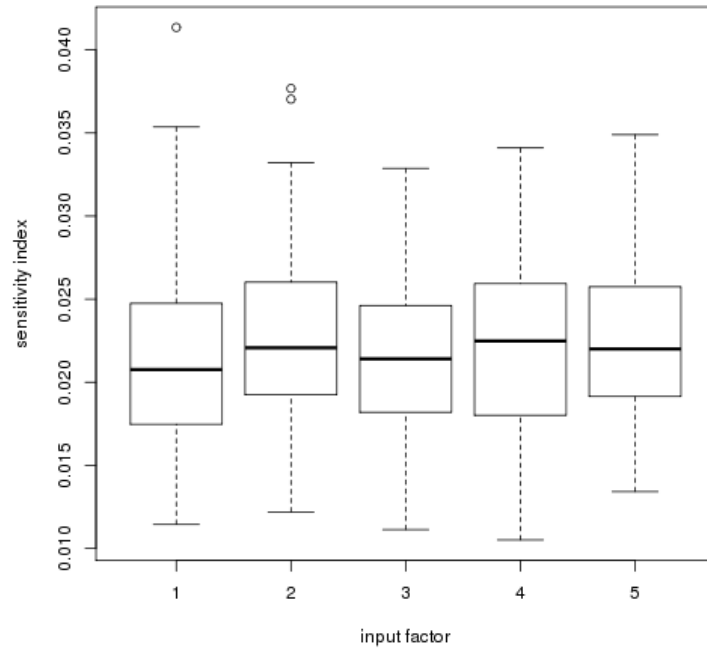


Figure 1.12: Boxplot of the estimated FOSIFD for the same importance different spread hyperplane case

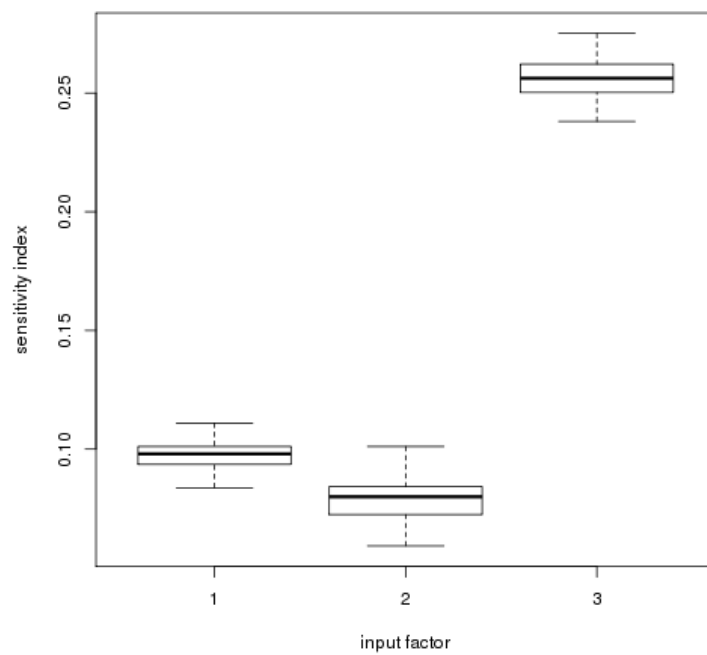


Figure 1.13: Boxplot of the estimated FOSIFD for thresholded Ishigami case

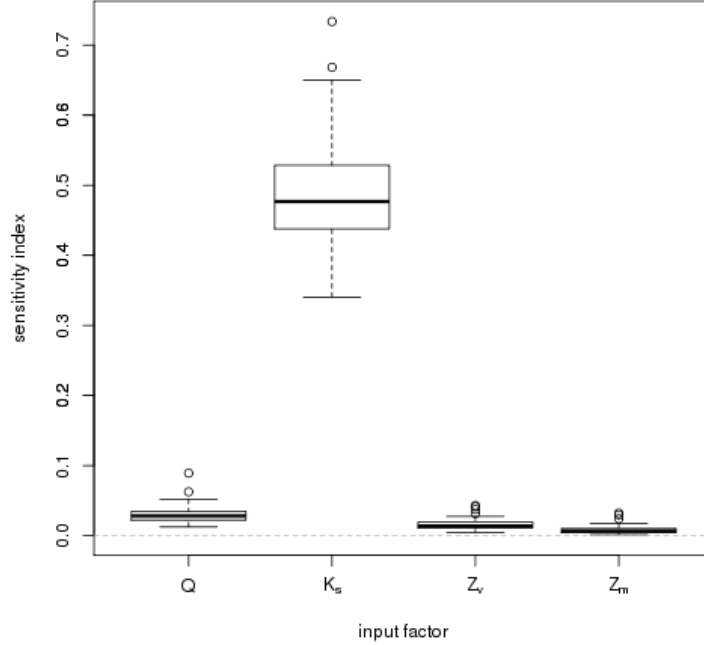


Figure 1.14: Boxplot of the estimated FOSIFD for the flood case

One can see from the boxplot that the FOSIFD is around 10% for variable 1, 8% for variable 2 and 25% for variable 3. The conclusion of such a result is that fixing variable 3 would provide the greatest variance reduction in the failure domain.

1.4.6.8 Flood case

This numerical example has been described in Appendix B.3. The results of the experiment with the same global parameters (100 runs, two samples of size 10^6) are boxplotted in figure 1.14.

The FOSIFD assess that the variable K_s is of first importance to explain the variations of the failure function within the failure domain, with almost 50% of the variance explained. All the other variables have a weak influence, and the ranking is as follows: Q then Z_v and finally Z_m .

1.4.6.9 Conclusion on FOSIFD

The FOSIFD method can be considered as a by-product of MC technique, since the computational cost of the FOSIFD is negligible compared with the time needed to obtain the samples/responses. This method has shown a capacity to assess which variable needs to be fixed to get a reduction of variance within the failure domain, see for instance section 1.4.6.7.

However, this method focuses on how does the failure domain behaves, and not on what causes the failure. One could possibly imagine an example in which the variables that cause the most variation within the failure domain are not the ones leading to failure.

This example might be the following:

$$G(\mathbf{X}) = 1_{X_1 < .5} + 0.2 \times \sin(10X_2) \quad (1.86)$$

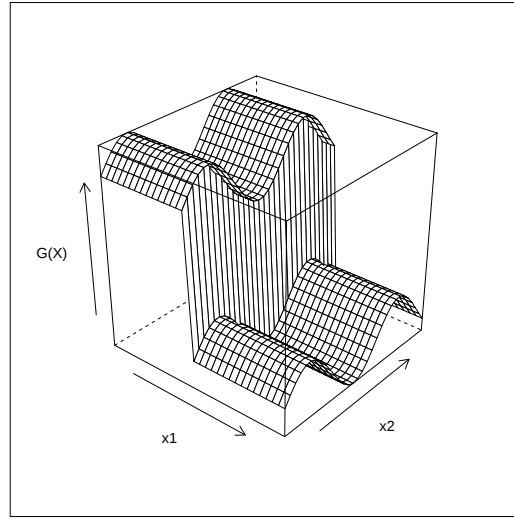


Figure 1.15: Example surface

where $X_1, X_2 \sim \mathcal{U}[0, 1]$ and the failure event is when $G(\mathbf{x}) < 0$. The surface picturing such a function is displayed in Figure 1.15 .

It can be seen that the failure event is only caused by variable X_1 whereas the variation within the failure domain is only caused by variable X_2 . The Sobol' indices of the indicator function are $S_1 = 1$ and $S_2 = 0$ whereas the FOSIFD worth respectively 0 and 0.91 for variables X_1 and X_2 . Consequently, if the objective is the variance reduction within the failure domain, one should focus on variable X_2 but if the objective is to understand what causes the failure event, one should focus on variable X_1 .

As in our study we are more interested in the failure event, we will not pursue the testing of the FOSIFD method.

1.4.7 Conclusion on Sobol' indices for reliability

Sobol' indices applied directly on the indicator function have shown a capacity to separate the influential and non-influential variables. Based on this observation, it seems an adapted method for sensitivity analysis in the reliability context. However, in most tested cases, Sobol' indices behave as follows: weak first order indices, strong total indices. This assesses that no variable is influential on its own, and that most variables contribute to the failure probability when interacting with the others. Unfortunately in most structural reliability cases, this is an already known information: it is when all the variable takes extreme values at the same time that the equipment fails. However, Sobol' total indices convey a strong information if the objective is the discrimination of the influential and non-influential variables.

One can observe that this useful information is obtained at a strong computational cost. As a rule of thumb we suggest to use samples of size 10^6 for failure probabilities of order 10^{-3} : with smaller sample sizes the estimations might be too noisy. Variance reduction technique have been studied, QMC and importance sampling. QMC allows a reduction of function calls of order 10. Importance sampling might be used if the goal of the SA is to rank the variable (i.e. obtain a qualitative information) and can lead to a reduction of function calls of order 100. However, such a reduction is possible only if a good importance density is available.

If the model is not costly we would recommend the use of such indices, using the Saltelli [87] method that allows an estimation of the first order and total indices. Other methods can be quoted and are compared in Saltelli *et al.* [88]. However if the model is costly, other methods than the Sobol' indices need to be found.

1.5 Moment independent measures for reliability

Let us study, in the reliability case, the indices defined in Borgonovo [13] that have been presented in section 1.3.1.4.

1.5.1 Application in the reliability case

For the reliability case, one has:

$$f_Y \sim B(P_f) \text{ with } P_f = \int 1_{G(\mathbf{x}) \leq 0} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (1.87)$$

where $B(p)$ denotes the Bernoulli distribution of parameter p . When fixing the i^{th} output to a given value x_i , one denotes:

$$f_{Y|X_i=x_i} \sim B(P_{x_i}) \text{ with } P_{x_i} = \int 1_{G(x_1, \dots, x_i, \dots, x_n) \leq 0} f_{\mathbf{X}_{-i}}(\mathbf{x}_{-i}) d\mathbf{x}_{-i} \quad (1.88)$$

Then, the shift defined in Equation (1.54) rewrites as follows:

$$\begin{aligned} s(x_i) &= \sum_{y=0}^1 |f_Y(y) - f_{Y|X_i=x_i}(y)| = |(1 - P_f) - (1 - P_{x_i})| + |P_f - P_{x_i}| \\ &= 2|P_f - P_{x_i}|. \end{aligned} \quad (1.89)$$

Thus the sensitivity index defined in Equation (1.56) rewrites:

$$\delta_i = \frac{1}{2} \mathbb{E}_{X_i} [s(X_i)] = \int f_{X_i}(x_i) |P_f - P_{x_i}| dx_i. \quad (1.90)$$

If the quantities P_f and P_{x_i} are known, this is a one-dimensional integral.

1.5.2 Crude MC estimation of δ_i

Let us explicit here the methodology to use in order to estimate the indices δ_i by crude MC. First of all, an estimation of P_f is made with N_1 points:

$$\hat{P} = \frac{1}{N_1} \sum_{j=1}^{N_1} 1_{G(\mathbf{x}^{(j)}) < 0} \quad (1.91)$$

where $\mathbf{x}^{(j)}$, $j = 1, \dots, N_1$ are i.i.d. realisations of $f_{\mathbf{X}}$. Then, for a given x_i that lies in the support of f_{X_i} , let us estimate P_{x_i} with a crude MC:

$$\hat{P}_{x_i} = \frac{1}{N_2} \sum_{j=1}^{N_2} 1_{G(\mathbf{x}_{-i}^{(j)}, x_i) < 0} f_{\mathbf{X}_{-i}}(\mathbf{x}_{-i}) \quad (1.92)$$

where $f_{\mathbf{X}_{-i}}(\mathbf{x}_{-i})$ is the joint pdf of $\mathbf{X}$ bereft of its i^{th} component and $(\mathbf{x}_{-i}^{(j)}, x_i)$ is a realisation of $f_{\mathbf{X}}$ where the i^{th} component is fixed at the value x_i . The cost for estimating P_{x_i} is N_2 function calls. Denoting $\hat{s}(x_i) = 2|\hat{P} - \hat{P}_{x_i}|$, one can estimate the first order index δ_i by:

$$\hat{\delta}_i = \frac{1}{2N_3} \sum_{k=1}^{N_3} f_{X_i}(x_i^{(k)}) \hat{s}(x_i^{(k)}) \quad (1.93)$$

Therefore, for d inputs the total estimation cost of all the first order indices is $d(N_3.N_2) + N_1$. This cost is prohibitive in our cases where at least 10^5 function calls are needed to get a correct estimation of the quantities.

1.5.3 Use of quadrature techniques

This technique is inspired by Caniou [21], who proposes to reduce the number of function calls by using a quadrature method, namely the Gauss-Legendre integration rule. Rewriting the equations for our problem, one has:

$$\tilde{\delta}_i = \frac{1}{2M} \sum_{k=1}^M w(k) f_{X_i}(x_i^{(k)}) \hat{s}(x_i^{(k)}) \quad (1.94)$$

for M quadrature points, and where the $w(k)$ are the weights associated to each point. The computational cost of the first order indices becomes $d(M.N_2) + N_1$, where $M \ll N_3$. According to Caniou [21], 30 quadrature points are sufficient to reach a good precision.

1.5.4 Use of subset sampling techniques

One can remark that the computational cost of the indices comes from the estimation of the conditional and unconditional failure probabilities, namely $\hat{P}_{x_i}$ and $\hat{P}$. To reduce the number of function calls, we can use subset sampling methods to estimate these probabilities, as presented in section 1.2.3. Assuming that we use the adaptive-levels algorithm, the number of function calls becomes a random variable, which is expected to take a value around $N.n_s = N.\lfloor \frac{\log P_f}{\log \alpha} \rfloor$, as described in Equation (1.38). One can expect that $N.n_s \ll N_2$ and $N.n_s \ll N_1$. Accordingly, the number of function calls to estimate all the first order indices should be around $d(M+1).N.n_s$ which is expected to be much smaller than $d(N_3.N_2) + N_1$.

1.5.5 Hyperplane 6410 test case

Let us focus on the hyperplane 6410 test case (Appendix B.1). One can rewrite an analytical expression of $s(\cdot)$, as presented in Equation (1.89). One has, for input X_i set at value x_j :

$$s_i(x_j) = 2|P_f - P_{i,x_j}| \quad (1.95)$$

where $P_{i,x_j} = P(G(\mathbf{X}) < 0 | X_i = x_j)$. This rewrites:

$$s_i(x_j) = 2\left| \phi\left(-k/\sqrt{\sum_{p=1}^d a_p^2}\right) - \phi\left((-k+x_j)/\sqrt{\sum_{p=1;p \neq i}^d a_p^2}\right) \right| \quad (1.96)$$

Consequently, one can estimate in a very precise way these quantities and thus δ_i . This goes the same for indices δ_{ij} and the higher order terms. These “true” values are displayed in table 1.12.

	Variable	X_1	X_2	X_3	X_4	
	δ_i	0.0039	0.0228	0.0154	0	
Group	X_1X_2	X_1X_3	X_1X_4	X_2X_3	X_2X_4	X_3X_4
δ_{ij}	0.0230	0.0159	0.0039	0.0271	0.0228	0.0154
Group	$X_1X_2X_3$	$X_1X_2X_4$	$X_2X_3X_4$			
δ_{ijk}	1	0.0230	0.0271			

Table 1.12: True values of δ_i for the hyperplane 6410 case

One can see that all the first order indices are rather small. According to Borgonovo [13], this result suggests that the effects of the variable on the failure event are non separable. This means that following the indices δ , interactions play a large role in the failure event. Indeed, one can see that most, if not all, shift in distribution is determined by an interaction between the three first variables. Unfortunately, that information is already known. Additionally, the first order indices can provide a variable ranking of the influence.

1.5.6 Conclusion

According to Table 1.12 and to complementary numerical tests, one can conclude the following on these moment-independent sensitivity measures. At first glance, the theoretical values shows that they are adapted for the discrimination of influential and non influential variables. On the other hand, the first order indices are all small and the estimation suffers from a positive bias. This drawback means that those indices are poorly adapted for sensitivity analysis in the reliability case, despite their sound properties.

1.6 Synthesis

This chapter has presented an overview of existing strategies for estimating failure probabilities and of sensitivity analysis methods.

First, the mathematical context for estimating failure probabilities has been set. We presented three classes of methods; yet it has been seen that theses classes are not partitioned. Approaches based on numerical approximation of the failure (limit state) surface have not been considered in this chapter. Dubourg [33] focuses on replacing in an adaptive way the failure surface by a meta-model. Li [64] focuses on the estimation of failure probabilities using sequential design of experiments and surrogate models.

Then, the main existing sensitivity analysis (SA) methods have been presented. Two of these methods (Sobol' indices and Borgonovo indices) have been tested on reliability toy examples. We conclude the following: the moment independent techniques are not adapted for the reliability case, due to a positive bias in the estimations. On the contrary, Sobol' indices applied to a failure indicator have highlighted a capacity to distinguish the non-influential from the influential variables. However, tests have shown that the following configuration -low first-order indices, high total order indices- is often present. Therefore the information provided by such indices is limited and may only confirm that all the variables interact to cause the failure event.

Table 1.13 is a short synthesis on the presented SA methods. In particular are itemized the available evaluation methods altogether with the pros and cons of the methods.

Indice	Sensitivity type	Evaluation method	Pros/Cons
Importance factors and direction cosine (1.3.2.2) α_i^* ; α_i^{*2}	Global/local First order indices	•Every design point finding algorithm	+ Potentially a very small number of function calls – Measure depending on the foudned design point – complicated interpretation in the physical space
Sobol' indices applied on the indicator (1.4) $S_i; S_{T_i}$	Global Every order indices, use of total indices.	• Sobol' (with QMC and/ or Importance Sampling) • Saltelli, Mauntz, Jansen, Janon-Monod • FAST/E-FAST/RBD • Use of meta-models (not treated here)	+ Every order indices allowing to quantify the influence of interactions – Total indices make more sense and their computation is costly – Limited information provided
Borgonovo indices (1.5) $\delta_i; \delta_{ij} \dots$	Global Every order indices	•Crude Monte-Carlo •Quadrature techniques •Subset sampling techniques	+ Good properties – Limited information provided – Positive bias in the estimation

Table 1.13: Synthesis on the tested SA methods

In the next section, we extend our thoughts on SA for failure probabilities.

1.7 Sensivity analysis for failure probabilities (FPs)

A common point of view on SA is that it is the *art of determining the model inputs the most influential on the output*. But what does exactly "influential" mean, especially in the reliability field where an input can be "influential" on the model output but can have a small "influence" on P_f ? The present paragraph focuses on the meaning of SA for FPs. This is motivated by a practitioner-friendly point of view.

Let us ask the question: what are the reliability engineer's motivations when he/she performs a SA on his/her black-box model that produces a binary response? In the global introduction, we provided an overview of the "general objectives" of SA: variable ranking, model simplification, model understanding. But from our discussions with EDF practitioners, we have identified three "Reliability Engineer Motivations" (REM):

- **REM1**: the practitioner wants to determine which are the inputs that impact the most the failure event - the inputs distributions being set and supposed to be perfectly known. This amounts to an absolute ranking objective.
- **REM2**: P_f will be impacted by the choice of the input distributions; the reliability engineer wants to assess the influence of this choice on P_f . Therefore the objective here is to quantify the sensitivity of the model output to the family or shape of the inputs, making the assumption that the parameters of the underlying distribution are perfectly known (thus set to fixed given values).

- **REM3**: in practice, input distributions are estimated from data, thus leading to uncertainty on the values of the distribution parameters. The practitioner wants to assess the influence of the distribution parameters on P_f . Therefore the objective here is to test the sensitivity of the model to the parameters of the inputs

Conversely, we present here what we meant by "general use" of SA.

- Variable ranking (**objective 1**) is to assess which input "most needs better determination" (Saltelli *et al.* [89]). This means that after the SA, a variable ranking is wanted in order to know how the uncertainty relative to each input (often assimilated to the inputs' variance) is reverberated on the output uncertainty (variance). The effects of such an analysis is then research prioritization, to collect new data allowing to reduce the uncertainty on the selected inputs thus on the output. A typical tool for such a need is Sobol' indices. But what exactly is the uncertainty of the output in the reliability case? The output is a Bernoulli random variable with parameter P_f , but does its variance ($P_f(1 - P_f)$) reflects well the uncertainty on the quantity of interest P_f ?
- Model simplification (**objective 2**) would rather be determining which inputs can be set to a reference value or to any value of its support without affecting the model precision. This amounts to determine non-influential inputs. The use of such a result can be model dimension reduction. In the case of reliability, it is known (Pastel [79]) that not all P_f estimation methods resist well to a large dimensional problem. The aim of SA in this case is then allowing the use of sharper P_f estimation methods.
- Model understanding (**objective 3**) includes all information gained after the SA, for instance which particular values of some inputs leads to some behaviour of the output. In the reliability case, this amounts to determining which inputs/groups of inputs/specific zones of the support of specific inputs lead to the failure event. After such an analysis, the practitioner might take actions to avoid this specific input behaviour (by replacing an equipment, warming injection water, raising a dam among others corrective actions).
- Let us add a new item: calibration sensitivity (**objective 4**). In practice the inputs of the model are not fully determined and are calibrated with the following procedure: the family of the input is given by the physic laws (for instance the Weibull distribution which historically comes from the field of fracture mechanics) whereas the parameters of the distribution are data-driven. But given the lack of data/knowledge, the modelled input can be far from the "real" (physical) input. In this case and in the reliability context, the practitioner might want to know how this distributions/parameters errors impact P_f .

Let us explicit in Table 1.14 the correspondence between the general objectives and the engineers' motivations.

	Objective 1	Objective 2	Objective 3	Objective 4
REM1	×	×	×	
REM2				×
REM3				×

Table 1.14: Correspondence between the general SA objectives and the engineers' motivations

As noticed in Section 1.4, the direct application of Sobol' indices on the failure indicator provides the following pattern: very small first order indices, very large and similarly equal total indices. The interpretation of this pattern is that all/most of the variables play an active role in the failure event (objective 3) and we can use the total indices to provide a variable ranking (objective 1). However the answer to both these questions is in practice already known (the practitioner knows that the equipment fails when all variables take extreme values at the same time). Accordingly, this method can in some cases detects non-influential inputs (objective 2). But from the practitioner point of view, Sobol' indices only fulfills REM1.

In the following of this thesis, we propose 3 specific methods allowing to answer the different objectives.

The two first methods are itemized in Chapter 2 and provide a variable ranking (objective 1, REM1). Specifically, the first method makes use of sensitivity indices produced by a classification method (random forests). The second method measures the departure, at each step of a subset method, between each input original density and the density given the subset reached.

The method presented in Chapter 3 will be referred to as Density Modification Based Reliability Sensitivity Indices (DMBSRI). These indices altogether with their estimation methods have been initially presented in Lemaître and Arnaud [62] then in Lemaître *et al.* [63]. They are based upon an input pdf modification, and quantify the impact of such a modification on the FP. We argue that with an adapted perturbation, this method can fulfill the four presented general uses (objective 1 to 4), altogether with the three engineers' motivations (REM1 to 3). This will be developed further in section 3.3.3.

Chapter 2

Variable ranking in the reliability context

2.1 Introduction

As stated in Section 1.7, there is a need in SA for techniques producing a variable ranking (REM1, objective 1). This chapter presents two methods allowing to rank the random inputs by their influence on the output. Furthermore, these methods are thought as by-products of the estimation of the failure probability P_f . Indeed the first technique (Section 2.2) proposes to make use of classification trees and random forests built on a MC sample. The second technique (Section 2.3) measures the departure, at each step of a subset method, between each input original density and the density given the subset reached. Thus both of these methods are by-products of two sampling techniques. Section 2.4 summarises the chapter and proposes a conclusion.

2.2 Using classification trees and random forests in SA

Classification trees and random forests are two well-known classification techniques. Additionally, sensitivity measures can be derived. This section aims at introducing these techniques. A state of the art on classification trees is proposed in 2.2.1. A subsection introducing the main stabilisation methods (such as random forests) is then studied in 2.2.2. Variable ranking techniques are derived in 2.2.3. The variable ranking is then tested on the usual cases in 2.2.4. A discussion is then proposed in 2.2.5, where the main theme is the improvement of models.

2.2.1 State of the art for classification trees

This section is widely inspired by Besse [11]; parts 3 and 4 of Briand [19]; but also parts 1 and 2 of Genuer [39] (in French). All those contributions are inspired by the founding monograph by Breiman *et al.* [17]. An introduction on statistical learning and the growing of classification tree can also be found in Hastie *et al.* [44].

Sample Let us assume that we have an input sample of $j = 1, \dots, N$ observations from d explanatory variables (or inputs) considered as quantitative, denoted by X_i^j , $i = 1, \dots, d$. A quantitative variable Y^j with two modalities is associated with these realisations of the inputs. Let us assume that the values taken by Y are in $\{0, 1\}$. In the considered framework, this sample might be the result of a Monte-Carlo experiment for a computer model where the quantity of interest is a probability of exceeding a given threshold (the events are failure/non-failure of the system). A sample aggregating the inputs and output of a subset simulation might also be used - this case is discussed

in Section 2.2.5. The sample is divided in two parts: a training set and a test set. The training set is used to fit the model (in the next section, the classification tree). The test set is used to assess the generalization error of the model (Hastie *et al.* [44]).

Growing a binary tree A classification tree is built by recursive partitioning of the input space. Focus will be set on the CART (Classification And Regression Tree) method, Breiman *et al.* [17]. Moreover, the regression case will not be treated here.

The growth (or fitting) of a classification tree is done in selecting a sequence of nodes (binary partition of the input space) then in determining a subsequence (pruning) that will be optimal according to a given criterion. A node is defined by an input variable (splitting variable) and a division, allowing the separation of the sample in two subsamples. A division is defined by a value (split point). At the first node (also referred to as root of the tree) corresponds the whole sample; then iterations are made on the produced subsamples.

The algorithm requires :

- the definition of a criterion allowing to select the best node (variable+division);
- a rule to end the algorithm and decide that a node is terminal (also referred to as leaf);
- a rule to assign a terminal node to a class.

Division criterion Each variable $(1, \dots, d)$ produces $N - 1$ allowed splits (that is to say creating a non-empty node). There are $d \times (m - 1)$ allowed splits in which the optimal division must be chosen. The division criterion is related to a node impurity measure: the aim is to obtain nodes as homogeneous as possible with respect to the output Y . The impurity measure considers the mixture of Y 's modality in a node. It is null if and only if all the individuals of the same node share the same value of Y . It is maximal when the modalities of Y are equally present in the node.

The deviance (or heterogeneity) of a node k is denoted D_k . The reduction of deviance (or impurity reduction) from splitting this node into descending nodes t and s would then be:

$$\Delta D = D_k - D_t - D_s$$

The tree is built by taking the maximum reduction in deviance over the allowed splits:

$$\max_{\text{allowed splits } \delta} D_k - (D_t + D_s)$$

Stopping rule The algorithm stops for a given node when it is homogeneous (it contains a single class and therefore cannot be divided no more). The algorithm can also be calibrated to avoid useless splits: the division process is stopped when the number of values in the node is less than a fixed size (for instance 5 individuals).

Affectation rule If the terminal node (leaf) is homogeneous, it is affected to the represented class. If not, a majority rule is applied. If wrong-classification costs are given, the less costly class is chosen.

Heterogeneity criteria Let us propose two heterogeneity measures: the entropy criterion and Gini index (in practice this choice is less influential than the pruning criterion, Besse [11]).

Define p_{lk} the probability that an element of node k belongs to class l ($l = \{0, 1\}$ in our case). This quantity is estimated by $\frac{n_l(k)}{n_k}$ where $n_l(k)$ represents the number of individuals in node k presenting class l and n_k the number of individuals in node k .

The impurity of node k in the entropy sense is defined by:

$$D_k = -2 \sum_{l=0}^1 n_k p_{lk} \log(p_{lk}.)$$

The impurity of node k in the Gini index sense is:

$$D_k = \sum_{l=0}^1 p_{lk}(1 - p_{lk}).$$

Pruning A maximal tree might overfit the data (the training set) while a small tree might not explain the structure of the data. The pruning step is a model selection step. Breiman *et al.* [17] propose to select an optimal tree in a sequence of sub-trees.

Let us define the discrimination quality of a tree A : $D(A)$ as the sum of misclassified individuals. Let us define as well a cost-complexity measure $C(A) = D(A) + \gamma \times K$ where K is the number of leaves in the tree. The pruning algorithm starts with $\gamma = 0$ then increases the value of γ , allowing the building of a sequence of nested trees. It is straightforward that $D(A)$ will rise as K decreases. The selection of the final tree is done through cross-validation; or with a validation sample (or pruning sample) if the data size N is sufficient.

Example We propose in this paragraph a simple example of binary classification tree, coming from Mishra *et al.* [68]. The data set is presented in Table 2.1. There are two inputs and one binary output, taking the values "Safe" and "Failure".

X_1	4	3	1	5	9	11	2	6	9	8	6	7
X_2	5	1	3	4	2	6	7	8	9	10	11	12
Y	Safe	Safe	Safe	Failure	Failure	Failure	Safe	Safe	Safe	Safe	Safe	Safe

Table 2.1: Data set

The following tree can be constructed (Figure 2.1), where it can be noticed that all the leaves are pure (containing only one category). On the R environment, library **rpart** was used to build this tree.

2.2.2 Stabilisation methods

A classification method is said to be unstable if a small perturbation in the training set generates a large perturbation in the final predictor. Tree-based methods (such as CART method) have been identified as unstable. A review of classification tree stabilisation methods is proposed.

2.2.2.1 Principals overall strategies (Genuer [39], Section 1.1.3)

The principle of this family of methods is to build a collection of predictors then aggregate their predictions. These overall strategies might be applied with CART as predictors. In the classification

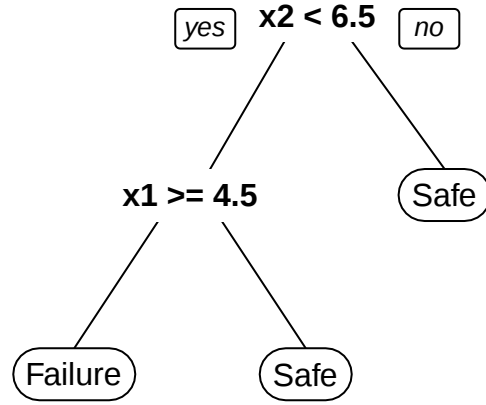


Figure 2.1: Binary tree

case, the aggregation is done with a majority vote. The aim of this class of methods is to avoid overfit.

Bagging Proposed by Breiman [15] with CART as predictors, bagging is the contraction of *bootstrap aggregating*. The main idea is to build, from the training sample, a number of bootstrap samples, then to aggregate the predictions. The generic bagging algorithm is presented in Algorithm 1. In our particular case, the chosen predictor is the classification tree of CART.

Algorithm 1 Bagging

Let X^0 be a set of inputs for which a forecast is wanted and $Z = (X^j, Y^j)_{j=1, \dots, N}$ a training sample.

For $b = 1, \dots, B$ **do**:

- Sample a bootstrap sample Z_b
- Estimate the predictor h_{Z_b} on this sample

End for

Compute the mean prediction $h_B(X^0) = \arg \max_j \# \{b | h_{Z_b}(X^0) = j\}$.

Boosting Proposed by Freund et Shapire [36], this type of algorithm is widely used with CART as predictors.

The principle is the sequential construction of models in which important weights are affected to misclassified individuals. The founding algorithm Adaboost (*Adaptive boosting*) is described in the case of a discrimination problem with two classes $\{-1, 1\}$. An initial bootstrap sample is sampled, where each individual has the same probability to appear. A classifier (predictor) is estimated, altogether with its classification error. A second bootstrap sample is generated, where misclassified individuals are more likely to appear. Another predictor is fitted and the algorithm continues. Each sample is generated according to the performance of the previous classifier. At the end, all

the classifiers are aggregated in function of their respective weights. A summary is presented in Algorithm 2.

Algorithm 2 Boosting

Let X^0 be a set of inputs for which a forecast is wanted and $Z = (X^j, Y^j)_{j=1, \dots, N}$ a training sample.

Initialize the weights $w_i = 1/N$; $i = 1, \dots, N$

For $m = 1, \dots, M$ **do**:

- Estimate classifier h_{Z_m} on the bootstrap sample weighted by w
- Compute the error rate:

$$err = \frac{\sum_{j=1}^N w_i 1_{\{h_{Z_m}(X^j) \neq Y^j\}}}{\sum_{j=1}^N w_i}$$

- Compute the logit $l_m = \log\left(\frac{1-err}{err}\right)$
- Compute the new weights $w_i := w_i \exp\left[-l_m 1_{\{h_{Z_m}(X^j) \neq Y^j\}}\right]$ $i = 1, \dots, N$

End for

Compute the mean estimation $h_B(X^0) = \text{sign}\left[\sum_{m=1}^M l_m 1_{\{h_{Z_m}(X^0) \neq Y^j\}}\right]$.

2.2.2.2 Random forests

The presented algorithm is RF-RI (*Random Forest - Random Input*) described by Breiman [16]. The main idea is to improve CART bagging with a step of random selection of inputs in the model. More specifically, a large number of trees are grown, each tree on a different bootstrap sample. At each node, m inputs among d are randomly selected, then the split is done. Section 1.3 of Genuer [39] presents a complete review for several versions of random forests. Algorithm 3 sums up the ideas.

Algorithm 3 Random Forests

Let X^0 be a set of inputs for which a forecast is wanted and $Z = (X^j, Y^j)_{j=1, \dots, N}$ a training sample.

For $b = 1, \dots, B$ **do**:

- Obtain a bootstrap sample Z_b
- Estimate a CART on this sample with variable randomisation:
 - at each node, randomly (uniform without replacement) pick m of the d inputs;
 - for each of the m variables, find the best split among the possible splits for the k -th variable;
 - among the m proposed splits, select the best one;
 - split the data using the selected best split;
 - repeat the previous steps until a maximal tree is growth.
- The final predictor is denoted h_{Z_b} .

End for

Compute the mean prediction $h_B(X^0) = \arg \max_j \# \{b | h_{Z_b}(X^0) = j\}$.

The default value for m in the classification context is $m = \sqrt{d}$. Notice that each tree is maximal and is not pruned. Some theoretical results on pure random forests (PRF) are available in Biau [12].

2.2.2.3 Structure stabilisation methods (Briand [19], section 4.4)

The presented stabilisation methods such as Bagging and Random Forests consist in the construction of a large number of classifiers on a randomized sample. These techniques improve the capacity of the predictors but the singular tree structure is lost. This singularity might be a requirement when the aim of the classification is the proposal of a decision tree. The techniques proposed hereafter aims at keeping the structure of the tree by stabilizing the nodes.

The method proposed by Ruey-Hsia [86] consists in including, for each node, logical structures. For instance, a division criterion might be " $2 \leq X_i$ and $X_k \geq 5$ ". The notion used to reach such a result is the existence of a division "almost as good" as the optimal. Briand [19] remarks that the existence of a large number of logical expressions might complicate the interpretation of the tree.

Choice is then set to use a method allowing a stabilisation of the nodes (division and variable associated) of the tree. The inspiration comes from Dannegger [29]. The main idea is to re-sample in a bootstrap fashion for each node. For each sample, the optimal division is searched. The variables most frequently selected are then used as a division variable for the treated node.

Briand proposed Dannegger's algorithm to build a maximal tree, then to prune the tree with a *reduced error pruning* method, Quinlan [81]. The couple tree growing/pruning is denoted REN method. An article by Briand *et al.* [20] proposes a similarity measure between trees - that might be of different structures. This similarity measure is used in Briand [19] to compare trees built with CART method or with REN method. It allows to assess the stability of the REN method to build classification trees.

2.2.3 Variable importance - Sensitivity analysis

2.2.3.1 Criteria definition

Tree-based classification methods are mostly used in the genomic domain, where the number of variables is much higher than the number of observations ($N \ll d$). Thereby, different importance measures have been considered by several authors. These measures are presented here, reminding that their aim is the selection of a few inputs among a large number of explanatory variables.

CART case A naive idea of variable ranking is that the variables most involved in the partition (and especially those which nodes are close from the root) are the most influential. A more refined idea has been proposed by Breiman *et al.* [17]. It is defined as the sum on the nodes of the heterogeneity reduction (for substitution divisions). An introduction on this index is presented in Ghattas [40]. It is also used altogether with the REN stabilisation method of Briand.

RFRI case When building a large number of trees, and randomizing each construction step, the unique structure described in the CART case is lost. Thereby, new sensitivity measures are proposed by Breiman [16].

- A first naive estimator of a variable's influence is the frequency of its apparition in the forest.
- A second estimator is said to be "local", it is based on the sum of the heterogeneity reduction (in the Gini index sense) on nodes where the variable is used. This criterion will be denoted GI in the following. The importance criterion V_{GI} is the sum of the heterogeneity decrease due to variable X_i , divided by the number of trees in the forest N_{trees} .

- Third measure is said to be "global" and is named MDA index (*Mean Decrease Accuracy*). It is based on a random permutation of the values of the considered variable. In a simplified way, if the variable is influential then the prediction error on the perturbed sample will be high. This prediction error will be smaller/null if the perturbation is done on a non-influential variable. More precisely, let us denote err_{oob} the "Out-of-Bag" error, the prediction error on the part of the sample (*OOB*) that has not been used to estimate the tree (the whole sample bereft of the bootstrap sample). The values of the i^{th} variable are permuted in the *OOB* sample; then the prediction error is computed on this sample. This error is denoted $err_{oob,i}$. The MDA index might be negative, and is defined as follow:

$$MDA(X_i) = \frac{1}{N_{trees}} \sum_{t=1}^{N_{trees}} (err_{oob,i}^t - err_{oob}^t)$$

2.2.3.2 Review of works on SA with CART/RFRI

In this part, a historical (from the oldest to the newest) review of the use of CART/RFRI for SA is presented. We tried to focus on the case $N \gg d$ or $N \simeq d$.

- Mishra *et al.* [69]. The topic of this article is SA. Four methods are presented, including one based upon CART classification. CART is used by classifying "extreme" events (10 and 90 percentiles of the output). This paper quotes the following one for the methodology and presents the same results.
- Mishra *et al.* [68]. This paper's topic is SA on a binary output (10 and 90 percentiles of a scalar continuous output). The studied model (nuclear waste repository field) presents 300 inputs. The authors use 60 datas. The sensitivity measure used is the most simple ("*The earliest splits contribute most to the reduction in deviance and are considered to be most important in the classification process*"). On the application case, it turns out that 5 variables are used to build the CART that classifies the output as "high" and "low". To the best of our knowledge, it is the first paper to perform SA on binary output.
- Frey *et al.* [37]. In this research report, the authors list SA methods then apply them on several cases where the output is a scalar continuous value (CART is used in a regression context). The used index is the reduction of deviance (sum of square of the mean departure) due to each node.
- Frey *et al.* [38]. This research report is a review on SA. With respect to CART, the recommended use is regression. For SA the authors' point of view is to consider the variables selected in the tree as influential; then to rank them by their proximity to the root. The previous report is quoted, advising to use the deviance reduction index.
- Pappenberger *et al.* [77]. To the best of our knowledge, this article is the first dealing with Random Forests (RF) to produce SA in the sense of the present work (it is noticeable that this paper quotes Sobol' and Saltelli). However, the use of RFRI is for regression, therefore the sensitivity measures are not the same as presented in Section 2.2.3. Two indices are presented, one based upon an information gain and another based upon permutation of input values (somehow close to the MDA index). An extension of this last measure is proposed for several variables, yet this measure is to be used with care due to an additive assumption. The point of view of the authors is that their method can be combined with Regional Sensitivity analysis, (Saltelli *et al.* [89], Hornberger *et al.* [48]). The first applicative example might be interpreted

as a failure function exceeding a threshold, thus presenting an interest for the present research. The SA part on RFRI consists in fitting a large number of regression trees and boxplotting the results. The authors show the interest of their method (SARS-RT) in comparison with rank regression SA. The ranking of the variables is the same for influential variables when using the two proposed indices. However, the ranking differs for the weakly influential variables.

- Strobl *et al.* [95]. This article deals with comparison of three sensitivity measures (Selection Frequency/GI/MDA, see Section 2.2.3) for RFRI. The framework is the one of $N \ll d$; and where the output is binary $\{0, 1\}$. The trees used are then classifiers. The main contribution of this article is to show the instability of variables ranking indices. These indices tend to show that multi-modal inputs are influential when they actually are not. The strong bias of GI measure is shown. The authors propose a tree building procedure called subsampling, building a tree on a sub sample without replacement of size $0.632N$ where N is the sample size. They show the good behaviour of their procedure in most test cases.
- Archer *et al.* [3]. This paper deals with variable ranking ("variable importance") in the genomic framework ($N \ll d$, classifier trees, a large number of correlated input variables). The authors show on simulations the similarity of the two tested sensitivity indices (GI/MDA) and their usefulness to identify influential variable (even in the correlated case).
- Pappenberger *et al.* [76]. This article is a review then an application of 5 SA methods on a flood model. There are no use of CART or RFRI, but the paper by Frey *et al.* [37] is quoted for the introduction of CART in SA.
- Briand [19]. The main idea of this PhD is the use of CART for SA. The main contribution is a procedure of tree stabilisation, presented in 2.2.2.3. An article by Briand *et al.* [20] dealing with a similarity measure between trees has also been produced. This measure can be used in a random forest to express a "median tree". A SA can then be performed on this tree.
- Genuer [39]. This PhD proposes a complete state of the art on the construction of random forests. It also studies the properties of the MDA sensitivity indices for automatic variable selection in the $N \ll d$ case. The aim is to select a few inputs to build a parsimonious model.
- Sauve *et al.* [91]. The aim of this theoretical article is more to select variables rather than to rank them by influence. Theoretical results on model selection are presented, in the regression and classification cases.

Conclusion This bibliography shows that sensitivity analysis can be performed on a binary output using CART/RFRI as classifiers. Further investigation will be done in 2.2.4. From the bibliography, Gini importance measures and MDA sensitivity indices seems promising. Additionally, the paper from Strobl *et al.* [95] brought up an important point: there is a possible bias with the Gini importance measure when dealing with inputs that vary in their spread. This behaviour will be tested in the experiments to come.

2.2.4 Applications

On the R environment, library `rpart` is used to build CART models. Library `randomForest`, based on Breiman's Fortran code, is used to deal with RFRI along this report.

2.2.4.1 Hyperplane 6410 Case

This numerical example is described in Appendix B.1. The following experiment is performed 100 times. A 10^5 points sample is generated; on which a forest of 500 trees is built. At each step of the tree construction, $m = \sqrt{d} = 2$ variables are randomly chosen. Results obtained with MDA and GI are boxplotted in Figure 2.2 respectively left and right.

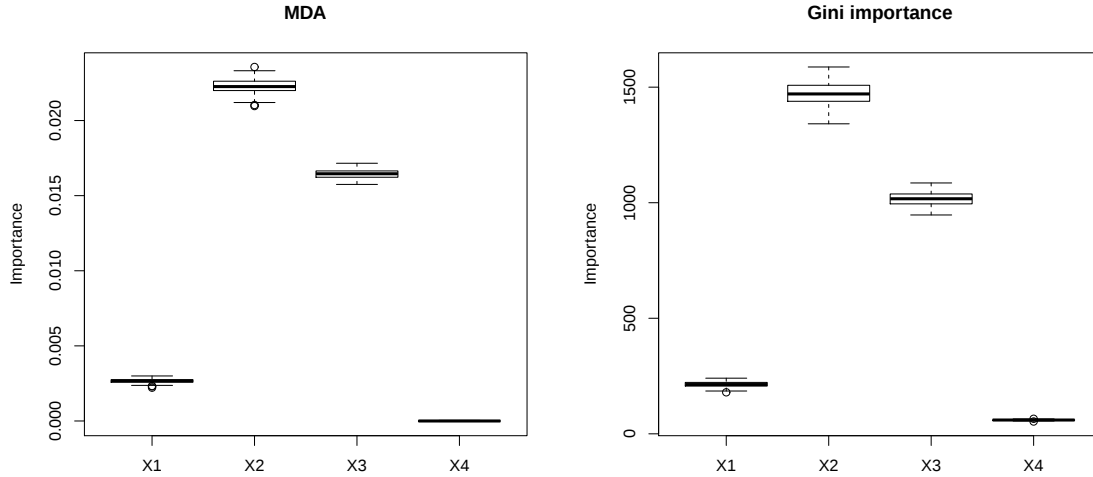


Figure 2.2: Boxplots of MDA indices (left) and GI indices (right) for the hyperplane 6410 test case

Both indices give the same variable ranking, identifying a strong influence for variable X_2 and X_3 . Variable X_1 is identified as weakly influential whereas variable X_4 is considered of very weak influence for GI indices and of null influence for MDA indices. This ranking is relevant given the coefficients of the variables.

2.2.4.2 Hyperplane 11111 Case

This numerical example is described in Appendix B.1. In term of SA, all the variables share the same influence. The following experiment is performed 100 times. A 10^5 points sample is generated; on which a forest of 500 trees is built. At each step of the tree construction, $m = 2$ variables are randomly chosen. Results obtained with MDA and GI are boxplotted in Figure 2.3 respectively left and right.

Both importance measures assess the same influence for all the variables. This was expected.

2.2.4.3 Hyperplane 15 variables test case

This numerical example is described in Appendix B.1. The following experiment is performed 100 times. A 10^5 points sample is generated; on which a forest of 500 trees is built. At each step of the tree construction, $m = 3$ variables are randomly chosen among the 15. Results obtained with MDA and GI are boxplotted respectively in Figures 2.4 and 2.5.

Both importance measures separate the influential variables (first 5), the weakly influential (6-10) and the non-influential (11-15). Once again, it is noticed that the GI measure does not allow to assess that a variable is "non-influential" but rather that a variable is less influential than the others, due to a non-null score. The explication of such a phenomenon might be the following. At

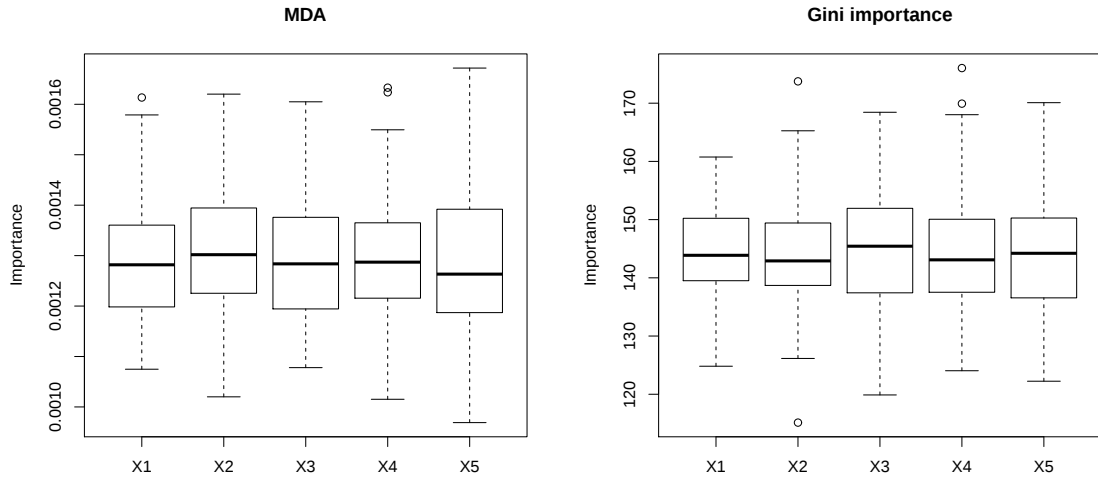


Figure 2.3: Boxplots of MDA indices (left) and GI indices (right) for the hyperplane 11111 test case

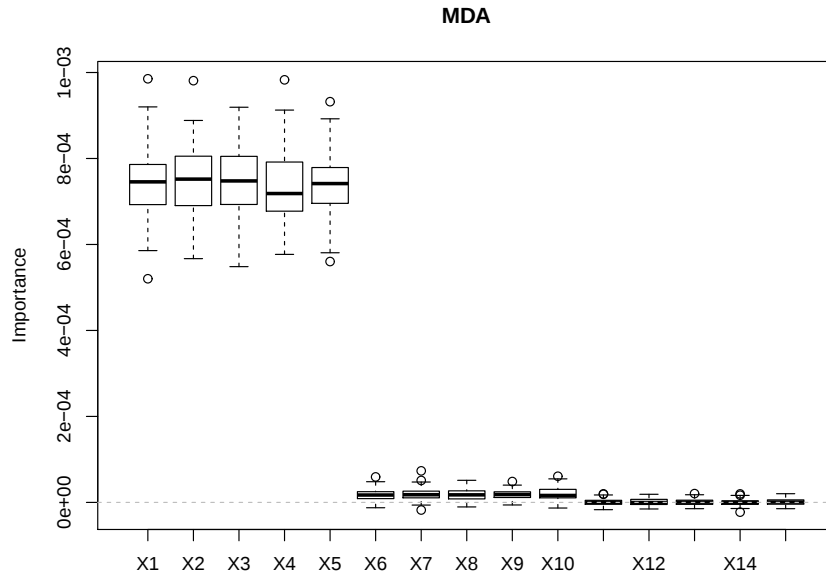


Figure 2.4: Boxplots of MDA indices for the hyperplane 15 variables test case

a node construction step, if the randomly chosen variables are only the non-influential ones, then the split will be done on one of these, thus reducing somehow the heterogeneity. This might explain the non-null GI measures for non-influential variables. However, MDA has a mean null score for non-influential variables, thus assessing their null impact on the failure probability.

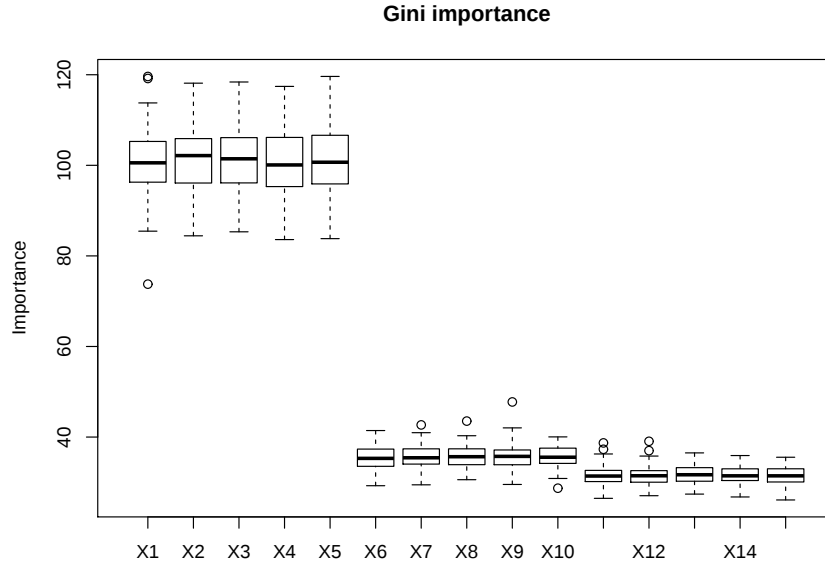


Figure 2.5: Boxplots of GI indices for the hyperplane 15 variables test case

2.2.4.4 Hyperplane with same importance and different spreads test case

This numerical example is described in Appendix B.1. The aim of such an example is to test the ability of both measures (MDA and GI) to give to each equally contributing variable the same importance despite their different spread. This test case is inspired by Strobl *et al.* [95] who have shown a strong bias for GI measure in case of multi modal or spread variables. The following experiment is performed 100 times. A 10^5 points sample is generated; on which a forest of 500 trees is built. At each step of the tree construction, $m = 2$ variables are randomly chosen. Results obtained with MDA and GI are boxplotted in Figure 2.6 respectively left and right.

It is noticeable that both measures show the same influence to all the variables, despite their different spreads. The boxplots do not present the bias of Strobl *et al.* [95]. Genuer [39] uses the MDA as a variable importance index over GI, due to the bias stressed by Strobl *et al.* [95]. However this "lack" of bias in our figures might come from the fact that these figures show an averaging of experience, thus an eventual bias might be neglected.

2.2.4.5 Tresholded Ishigami function

This numerical example is described in Appendix B.2. The parameters of the experiment are the following: 500 trees built on 10^5 points with $m = 2$ variables selected at each node construction step. Each experiment is reproduced 100 times. Results obtained with MDA and GI are boxplotted in Figure 2.7 respectively left and right.

According to the measures, there is no non-influential variable. The importance ranking differs with the measures. We recall that the problem raised with the GI measure is that one cannot assess that the less influential variable is non-influential. Our hypothesis on the different ranking is that binary trees do not fit efficiently separated failure surfaces. Figure B.1 is a plot of the shape of the failure surface for the Ishigami function: it seems difficult to fit a binary partition of the space for variables X_2 and X_3 .

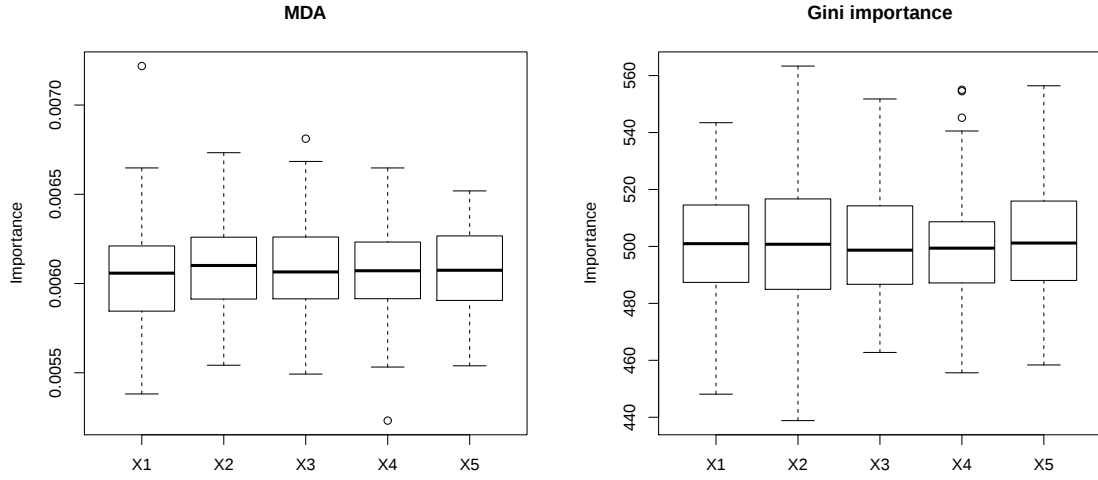


Figure 2.6: Boxplots of MDA indices (left) and GI indices (right) for the hyperplane different spreads test case

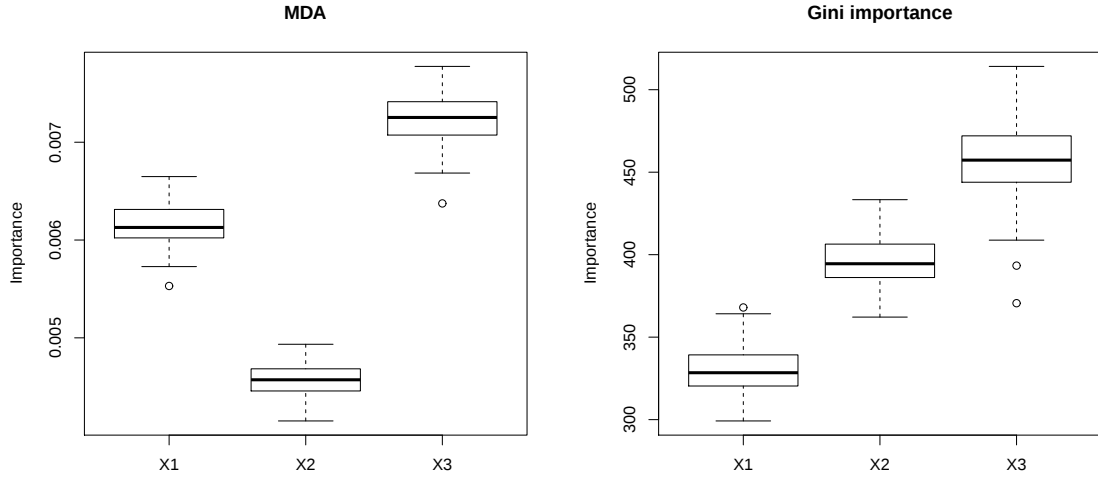


Figure 2.7: Boxplots of MDA indices (left) and GI indices (right) for the thresholded Ishigami test case

2.2.4.6 Flood Case

This example is described in Appendix B.3. The parameters of the experiment are the following: 500 trees built on 10^5 points with $m = 2$ variables selected at each node construction step. Each experiment is reproduced 100 times. Results obtained with MDA and GI are boxplotted in Figure 2.8 respectively left and right.

Variable ranking is the same on this test case. K_s is selected as the most influential variable, then comes Q . Z_v has a negligible influence while Z_m has a null influence (according to MDA indices).

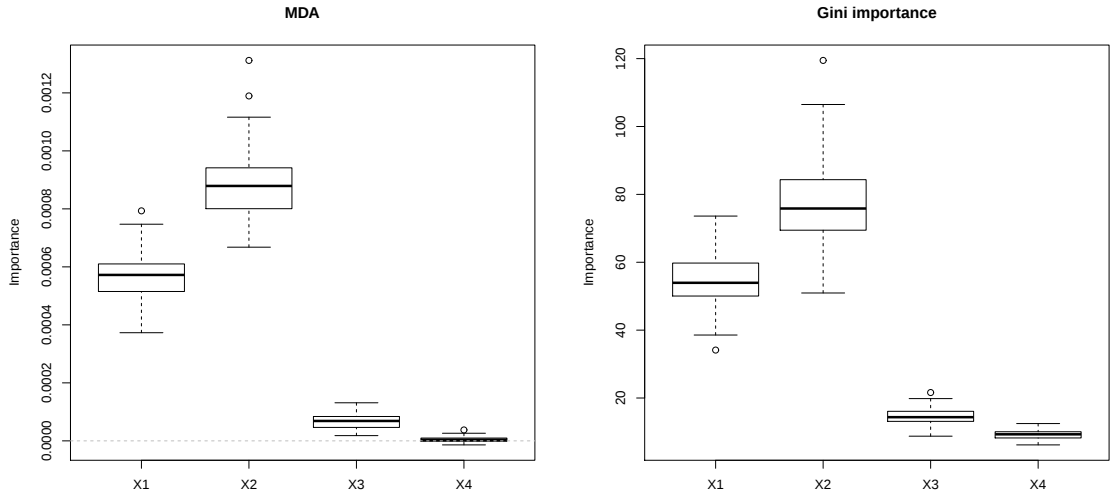


Figure 2.8: Boxplots of MDA indices (left) and GI indices (right) for the flood test case

2.2.5 Discussion

2.2.5.1 On the results of SA

Numerical experiments have shown the capacity for the proposed indices to rank the variables. This ranking is reproducible (boxplots with few covering on 100 repetitions). Except a complex case (thresholded Ishigami function), this ranking was the same for both studied measures.

However, GI measure can affect a non-null importance to a non-influential variable (as seen in Section 2.2.4.3). Even if on average, the same weight will be affected to all the non-influential variables, this numerical noise prevents to assess that a variable has a null influence. This drives us to prefer the MDA measure over the GI measure, since it allows the detection of non-influential variables.

2.2.5.2 On the model's quality

Problem noticing The study of fitted models (RFRI) shows that their quality is not satisfying. This might be a problem when drawing conclusions on SA with these models. More precisely, on a MC sample, the variable to be predicted presents two modalities in uneven quantities. For instance on the flood case, for a sample of 10^5 points there are 81 failure points whereas there are 99919 safe points. From this imbalance there is a tendency in getting "weak" predictors that make much more prediction error on the minority class. The confusion matrix (on the out-of-bag samples) of a forest of 500 trees is presented in Table 2.2.

		Observed		Class prediction error
		0	1	
Predicted	0	99912	7	7.01×10^{-5}
	1	27	54	3.33×10^{-1}

Table 2.2: Confusion matrix of the forest with default parameters

It is noticeable that the prediction error is around 5000 times higher for class 1 (failure) than for class 0 (safe mode). Given that the sensitivity measure chosen is an error averaging, it seems essential to improve the model's quality. The MDA ranking for this model is presented in Table 2.3.

	K_s	Q	Z_v	Z_m
MDA	6.28×10^{-4}	9.79×10^{-4}	5.22×10^{-5}	-1.96×10^{-6}

Table 2.3: MDA indices of the forest with default parameters

Class penalty A first idea to improve the models is to put a penalisation on the class so that the failure event is best predicted. This approach presents two drawbacks:

- making that choice turns the problem into the choice of the penalty;
- the model obtained might be a pessimistic one, predicting individuals of class 0 (safe mode) as being of class 1 (failure point).

A test affecting at each class weight proportionals to their frequency shows a weak improvement. The confusion matrix is presented in Table 2.4.

		Observed		Class prediction error
		0	1	
Predicted	0	99913	6	6.00×10^{-5}
	1	25	56	3.09×10^{-1}

Table 2.4: Confusion matrix of the forest with different weights

The MDA ranking for this model is presented in Table 2.5.

	K_s	Q	Z_v	Z_m
MDA	6.17×10^{-4}	9.74×10^{-4}	5.37×10^{-5}	3.64×10^{-6}

Table 2.5: MDA indices of the forest with different weights

The small modifications on the ranking and on the confusion matrix makes this solution inconclusive.

Increasing the number of trees Another solution is to increase the number of trees in the forest. A test is done on the same sample with 2000 trees (this value comes from Genuer [39]). The computing time is increased by a factor 10 on our machine. The confusion matrix and the MDA ranking are presented respectively in Tables 2.6 and 2.7.

		Observed		Class prediction error
		0	1	
Predicted	0	99914	5	5.00×10^{-5}
	1	26	55	3.21×10^{-1}

Table 2.6: Confusion matrix of the forest with 2000 trees

	K_s	Q	Z_v	Z_m
MDA	6.09×10^{-4}	9.71×10^{-4}	4.61×10^{-5}	4.49×10^{-6}

Table 2.7: MDA indices of the forest with 2000 trees

The confusion matrix does not present any improvement, despite the substantial increase of the computing time.

Increasing the sample size Another solution might be to increase the sample size. A test has been performed on a sample of size 5×10^5 for a forest of 500 trees. The computation failed due to the size of the sample. The solution is then inconclusive.

2.2.5.3 Importance sampling

To bypass the problem of the sample size, the use of importance sampling (see Section 1.2.1.3) is proposed. Therefore, the minority class will be artificially over-represented. For the flood case, the importance densities are the following:

- K_s follows a truncated Gumbel distribution with parameters 3000, 558 and a minimum 0;
- Q follow a truncated Gaussian distribution with parameters 10, 7.5 and a minimum 1;
- Densities of Z_v and Z_m are not modified.

Sampling 10^5 points according to these densities gives 49505 failure points (almost half of the sample). A forest of 500 trees is fitted on this sample. The confusion matrix is presented in Table 2.8.

		Observed		Class prediction error
		0	1	
Predicted	0	50001	494	9.98×10^{-3}
	1	498	49007	1.00×10^{-2}

Table 2.8: Confusion matrix of the forest built on an IS sample

Prediction error increases for class 0 (safe mode) with respect to Table 2.2. However prediction error decreases for class 1 (failure), this was wanted. Furthermore, the prediction errors for the two classes are of the same order of magnitude. The out-of-bag error on the whole model is around 1%.

MDA ranking on this model is presented in Table 2.9.

	K_s	Q	Z_v	Z_m
MDA	0.119	0.429	0.066	0.011

Table 2.9: MDA indices of the forest built on an IS sample

The ranking of the variables is the same, but the obtained values have a different order of magnitude. However, one cannot assess anymore that variable Z_m has a null influence.

To confirm these results, a forest of 1000 trees have been fitted. Results are similar and are not presented here.

However a question arises: do MDA indices computed on a sample that is not i.i.d. to the original densities have sense?

2.2.5.4 Using subset simulation

Another idea to solve the problem of unevenly represented classes without using importance sampling (that needs hypotheses on the importance densities) might be to use the results of a subset simulation. The sample would then have more failing points.

However, the MDA indices based on a coordinate permutation would not have sense anymore. Indeed, the individuals would not be i.i.d. with respect to the original densities, but block-wise i.i.d. to $f_{/D_k}$ where D_k are the subsets. One could then define an adapted measure of sensitivity:

$$MDA_S(X_i) = \frac{1}{N_{trees}} \sum_{t=1}^{N_{trees}} (err_{oob,i,S}^t - err_{oob}^t)$$

where the S stands for subset. The only difference here is in the way to compute $err_{oob,i,S}$. We propose the following: as the OOB sample is composed of individuals from different subsets $(D_1, D_2, \dots, D_K)$, perform the permutation of the i^{th} variable by subset (so that individuals coming from subset D_k are switched with individuals from the same subset). This error would be denoted $err_{oob,i,S}$.

These indices will be developed and tested on further works.

2.2.5.5 SA from the model selection point of view

Importance measures tested in this section have variable selection as primary objective. Their second objective is to fit parsimonious models (that do not use no more variables than necessary). The framework of such a procedure is generally the case $N \ll d$. Our studies is rather the case of discrimination of two classes unevenly present in a sample with $N \gg d$.

Nevertheless, this section has brought an interesting idea. This idea is to get a model collection built on the bagging principle, then to compute a sensitivity measure for each variable and to aggregate these measures to insure stability. It is definitely of interest and has to be explored in further works.

2.3 Using input cumulative distribution function departure as a measure of importance

In this section, a novel sensitivity measure is proposed. It is thought as a by-product of the subset sampling estimation technique (Section 1.2.3). The basic idea is to propose a sensitivity index for each variable at each step of the subset. The index is obtained as a departure in cumulative distribution function (c.d.f.) from the original. Subsection 2.3.1 introduces the idea and proposes some reminders. Subsection 2.3.2 makes a summary of all the distances analysed. The usual test cases are processed in Subsection 2.3.3. Finally, Subsection 2.3.4 sums up the ideas and concludes.

2.3.1 Introduction and reminders

As previously stated in the introduction, a sensitivity index for each variable at each step of the subset is proposed. The aim of such a proposition is to quantify step after step the influence of each variable on the failure probability. Let us give the informal definition: the sensitivity index is defined for the variable i and the subset step k as a departure between the empirical c.d.f. and the theoretical marginal c.d.f of the variable.

Considering M subset steps with $k = 1 \dots M$; denoting:

$$F_{n,i}^k = F_i(x|A_k), \quad (2.1)$$

the empirical c.d.f. of the i^{th} variable given that the subset A_k has been reached. Thus the proposed index writes:

$$\delta_i^{SS}(A_k) = d(F_{n,i}^k, F_i), \quad (2.2)$$

where F_i is the theoretical c.d.f. of the i^{th} variable, and d is a distance (defined further in Section 2.3.2).

Informally, an influential variable will have a strong departure in c.d.f. whereas a non-influential variable will have a weak departure in c.d.f., thus a weak index. Such a strategy is inspired by Monte-Carlo Filtering or Regionalised Sensitivity Analysis (RSA). However, it should be noted that several blocking points are identified:

- Information is neglected when working on the marginals. Moreover, when working with a particles cloud, the components are generally no longer independent. Thus the c.d.f. of the cloud is different from the product of the c.d.f. of the marginals. We decide to gloss over such problems for now.
- The choice of the distance measure will determine the importance ranking of the variables. It is therefore crucial to choose a distance adapted to the problem. The meaning of "influential" must then be set in advance (difference in the central tendency, difference in extremes...).

Choice has been set to work with empirical c.d.f. rather than with empirical densities for two reasons:

- Denoting $F_{n,i}$ an empirical c.d.f., Glivenko-Cantelli's theorem states that $\sup_x |F_{n,i}(x) - F_i(x)|$ converges almost surely to 0.
- More pragmatically, working with empirical densities (with a kernel smoothing) add an unnecessary processing.

2.3.2 Distances

We propose 3 distances coming from non-parametric statistics. These distances are used to define statistics of usual goodness-of-fit tests (Govindarajulu, [42]). Let us denote $F_{n,i}$ the empirical c.d.f. and F_i the c.d.f. to which it is compared (in our case, the theoretical original marginal c.d.f. of each variable).

2.3.2.1 Kolmogorov distance (L_∞ distance)

$$D_n = \sup_x |F_{n,i}(x) - F_i(x)|$$

The implementation of D_n is direct. D_n is the supremum of the departure between $F_{n,i}$ and F_i , it is thus the "worst case" distance.

2.3.2.2 Cramer-Von Mises distance (L_2 distance)

$$C_n = \int_{-\infty}^{+\infty} (F_{n,i}(x) - F_i(x))^2 dF_i(x)$$

The implementation of C_n can be done in two ways:

- C_n can be estimated using a numerical quadrature rule (such as Simpson's one);
- or denoting $U_j = F_i(X_j), j = 1, \dots, n$ and arranging this sample in order U_j^* then:

$$C_n = \frac{1}{n} \left[\sum_{j=1}^n \left(U_j^* - \frac{2j-1}{2n} \right)^2 + \frac{1}{12n} \right].$$

2.3.2.3 Anderson-Darling distance

$$A_n = \int_{-\infty}^{+\infty} \frac{(F_{n,i}(x) - F_i(x))^2}{F_i(x)(1 - F_i(x))} dF_i(x)$$

As C_n , A_n can be implemented in two ways:

- by quadrature ;
- or considering the U_j^* then:

$$A_n = \frac{1}{n} \left[-n + \frac{1}{n} \sum_{j=1}^n (2j-1-2n) \ln(1 - U_j^*) - (2j-1) \ln(U_j^*) \right].$$

Anderson-Darling distance is derived from the Cramer-Von Mises one but grants more weight to the extreme values.

2.3.3 Applications

2.3.3.1 Hyperplane 6410 test case

This numerical example is described in Appendix B.1.

Subset estimation First of all, the failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the true failure probability is $P_f = 0.014$. Note that in this case, the subset simulation method might not be the best adapted to estimate a "not so weak" failure probability. The parameters of the algorithm are the following:

- the proposal density is a Gaussian centred on the particle, with variance 1,
- $N = 10^4, \alpha = .75$.

The result with $15 \times N$ function calls is the exact result:

$$\hat{P} = 0.014$$

Plot of the c.d.f. For this first example, the c.d.f given that the third, the seventh and the fifteenth subset have been reached are plotted in Figure 2.9. One can see that whatever the distance used, the c.d.f. corresponding to the fifteenth subset is farther from the original one that the c.d.f. corresponding to the third subset (on this example).

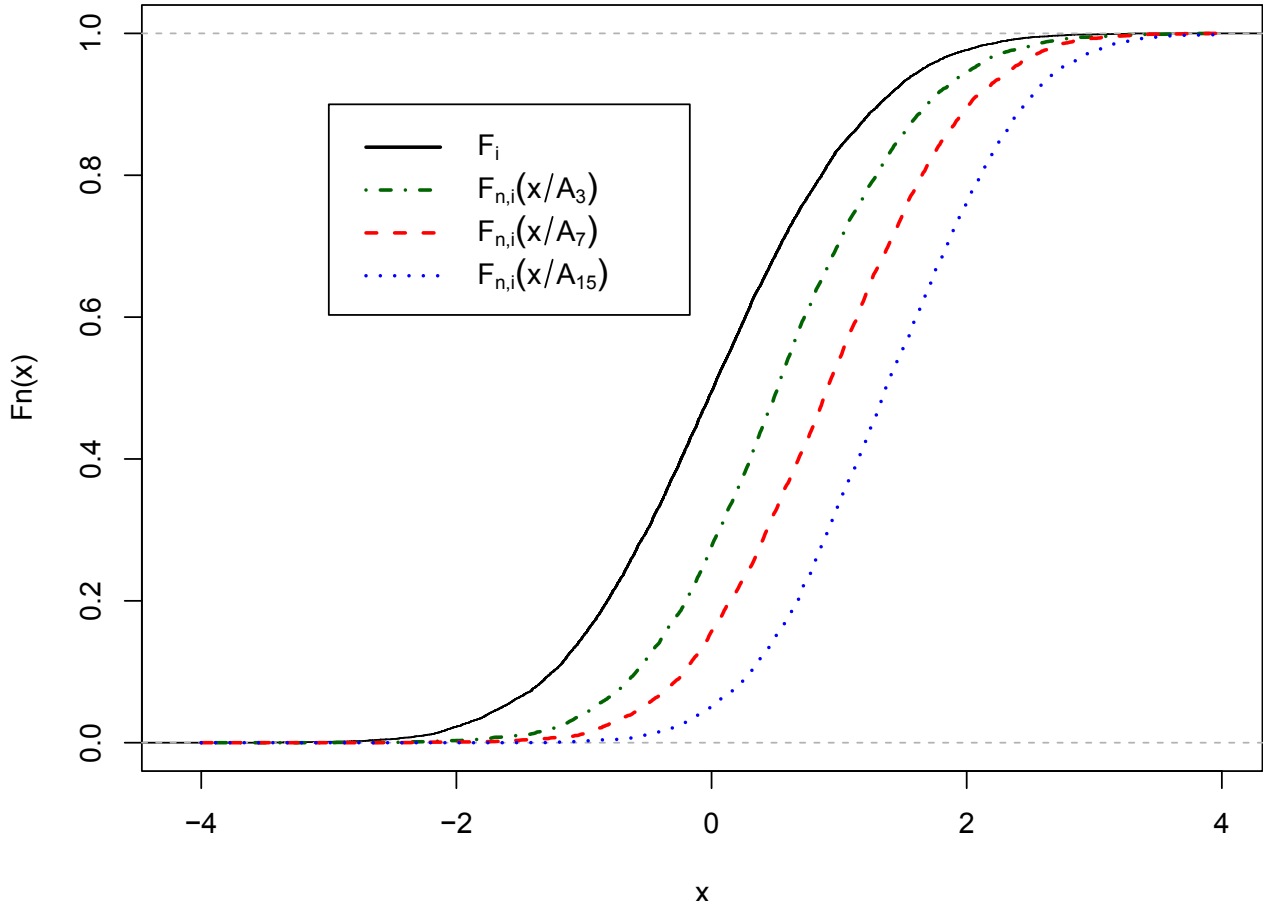


Figure 2.9: Several c.d.f.

Distance estimation The distance are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.10, 2.11 and 2.12. Variable X_1 is plotted in black, X_2 in blue, X_3 in green and X_4 in red.

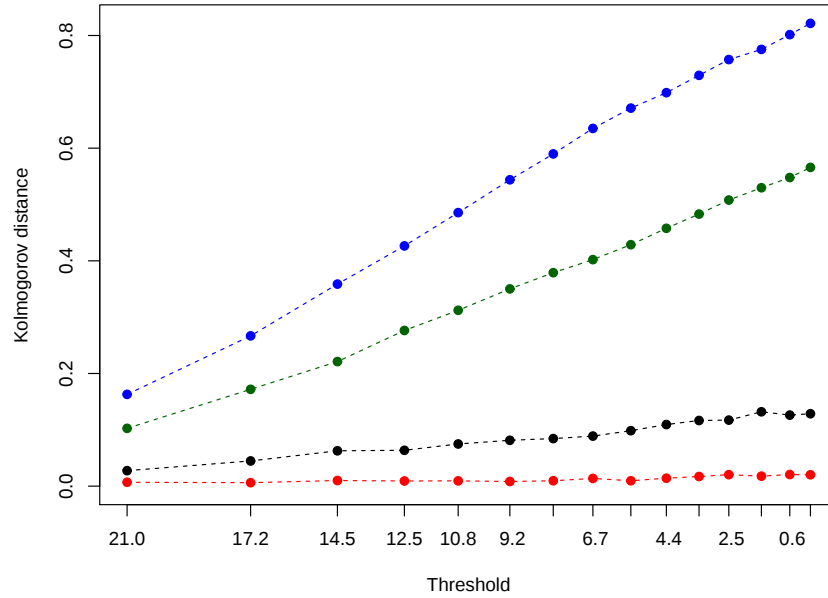


Figure 2.10: Hyperplane 6410 test case, Kolmogorov distance

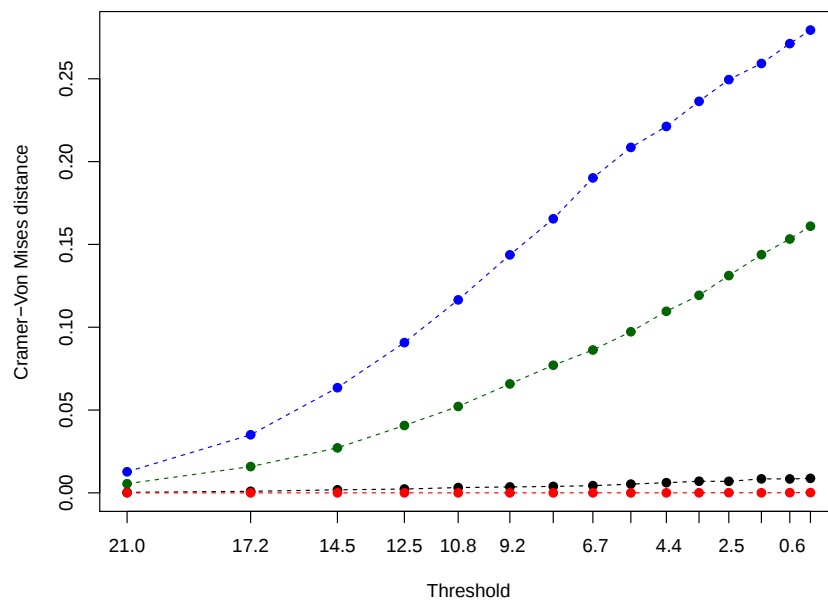


Figure 2.11: Hyperplane 6410 test case, Cramer-Von Mises distance

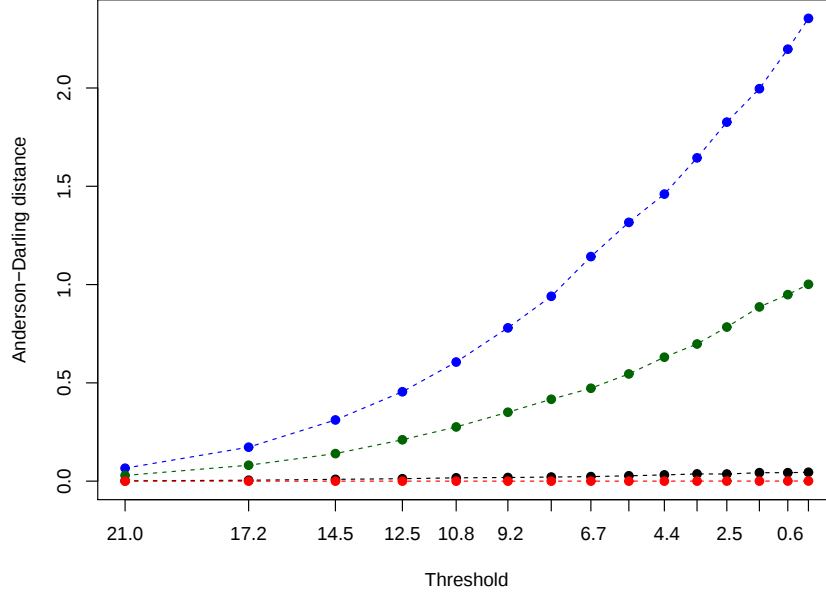


Figure 2.12: Hyperplane 6410 test case, Anderson-Darling distance

All the distances allow the following variable ranking: X_2, X_3, X_1 then X_4 . Notice that this is the same ranking than the one provided by the importance factors (see Table 3.3). All the distances tend to separate the variables in two groups. The Anderson-Darling distance seems to minimise the influence of the first variable (black).

2.3.3.2 Hyperplane 11111 test case

This numerical example is described in Appendix B.1. Recall that the aim of this test case is to assess the capability of the SA method to give the same importance to each input.

Subset estimation The failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the true failure probability is $P_f = 0.0036$. The algorithm's parameters are the following:

- the proposal density is a Gaussian centred on the particle, with variance 1,
- $N = 10^4$, $\alpha = .75$.

The result with $20 \times N$ function calls is the exact result:

$$\hat{P} = 0.0036$$

Distance estimation The distances are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.13, 2.14 and 2.15. A different color is used for each variable.

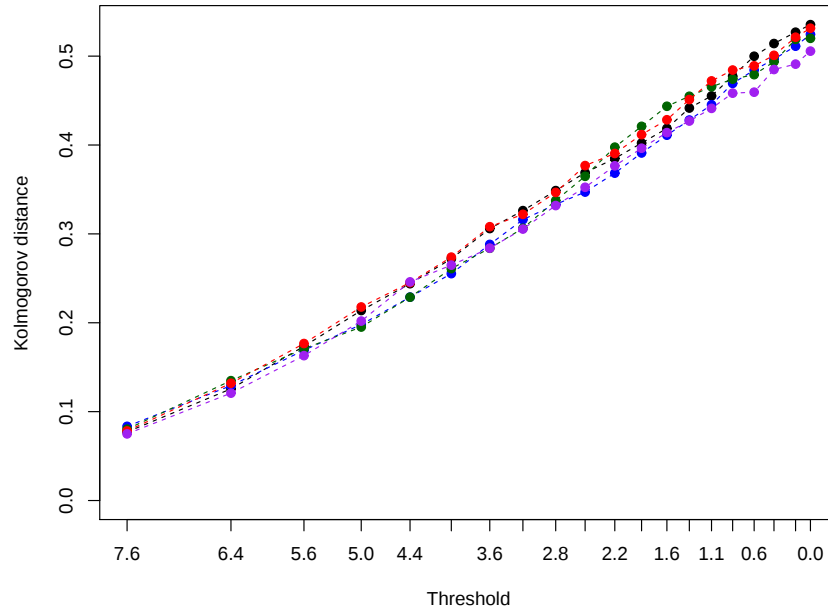


Figure 2.13: Hyperplane 11111 test case, Kolmogorov distance

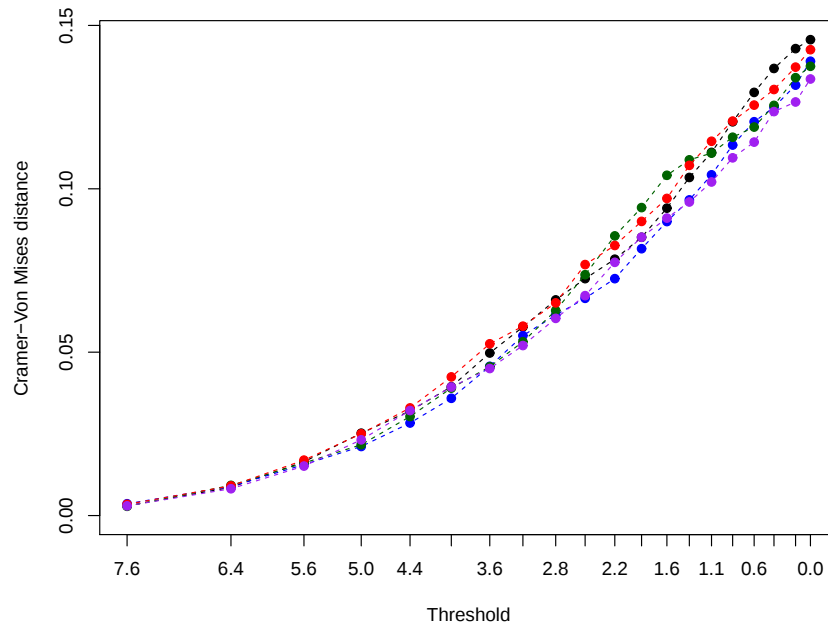


Figure 2.14: Hyperplane 11111 test case, Cramer-Von Mises distance

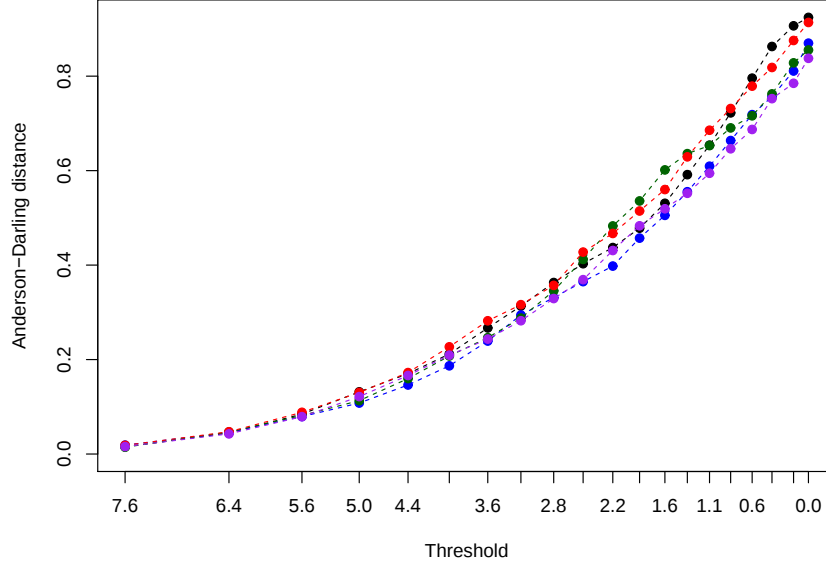


Figure 2.15: Hyperplane 11111 test case, Anderson-Darling distance

Every distance gives to the 5 variables the same importance. The distances growth with the threshold. So far, this SA method has proven that it can give the same influence to equally influential variables.

2.3.3.3 Hyperplane 15 variables test case

This numerical example is described in Appendix B.1. Recall that the aim of this test case is to class the inputs in 3 groups: influential, weakly-influential and non-influential.

Subset estimation The failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the true failure probability is $P_f = 0.00425$. The algorithm's parameters are the following:

- the proposal density is a Gaussian centred on the particle, with variance 1,
- $N = 10^4$, $\alpha = .75$.

The result with $19 \times N$ function calls is close from the exact result:

$$\hat{P} = 0.00454$$

Distance estimation The distances are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.16, 2.17 and 2.18. A different color is used for each variable. A different symbol (respectively a dot, a triangle and a square) is used for each group (respectively influential, weakly-influential and non-influential).

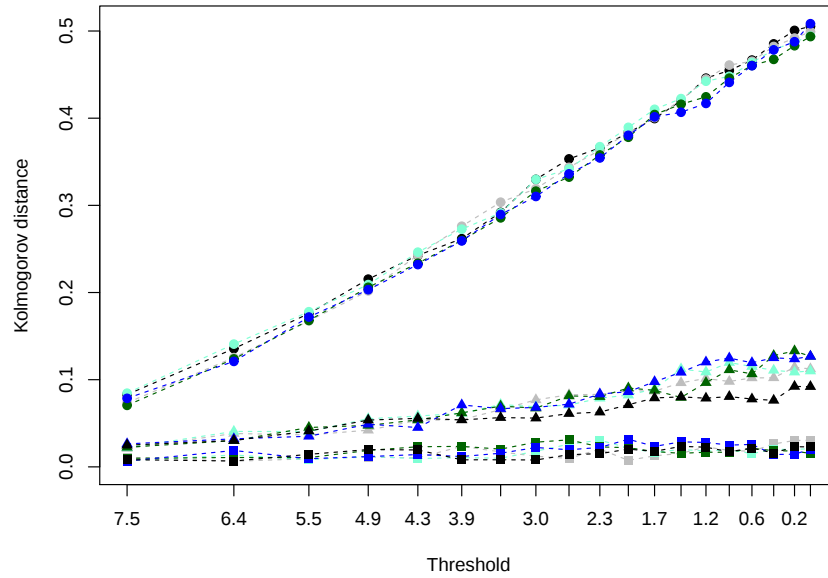


Figure 2.16: Hyperplane 15 variables test case, Kolmogorov distance

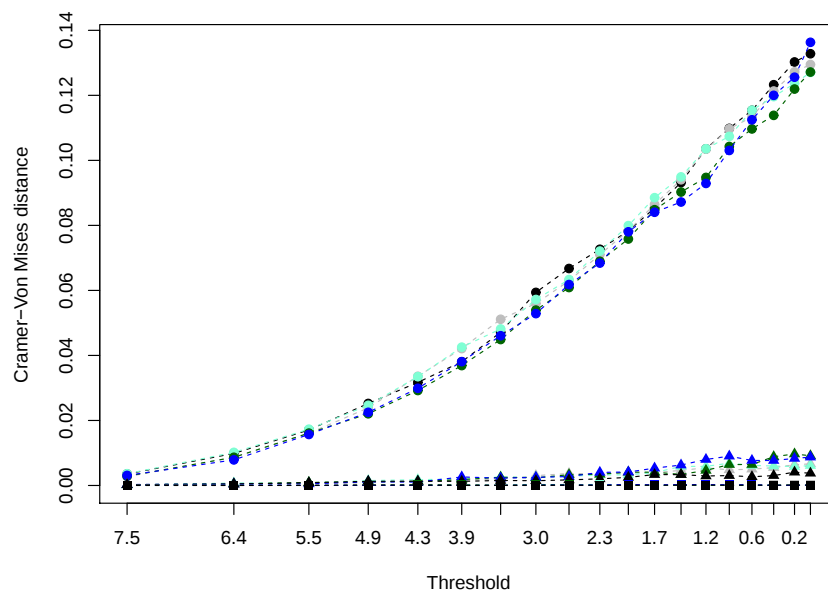


Figure 2.17: Hyperplane 15 variables test case, Cramer-Von Mises distance

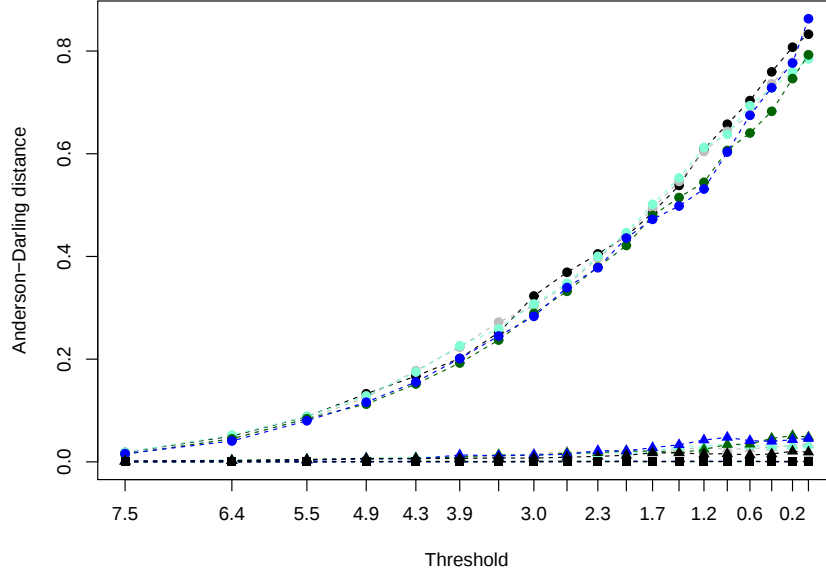


Figure 2.18: Hyperplane 15 variables test case, Anderson-Darling distance

All the distances growth with the threshold. Kolmogorov distance allows a separation of the inputs in 3 groups. On the other hand, both Cramer-Von Mises distance and Anderson-Darling separate the inputs in two groups: influential and non-influential.

2.3.3.4 Hyperplane different spread test case

This numerical example is described in Appendix B.1. Recall that the aim of this test is to assess the capability of the SA method to give to each equally contributing variable the same importance, despite their different spread.

Subset estimation The failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the true failure probability is $P_f = 0.0036$. The algorithm's parameters are the following:

- the proposal density is a Gaussian centred on the particle, with the same variance as the considered input,
- $N = 10^4$, $\alpha = .75$.

The result with $20 \times N$ function calls is close from the exact result:

$$\hat{P} = 0.0036$$

Distance estimation The distances are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.19, 2.20 and 2.21. A different color is used for every variable.

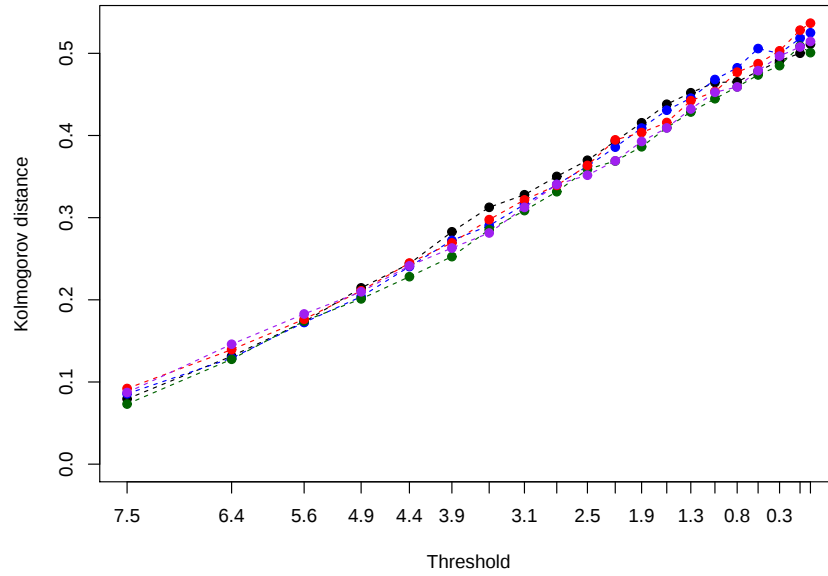


Figure 2.19: Hyperplane different spread test case, Kolmogorov distance

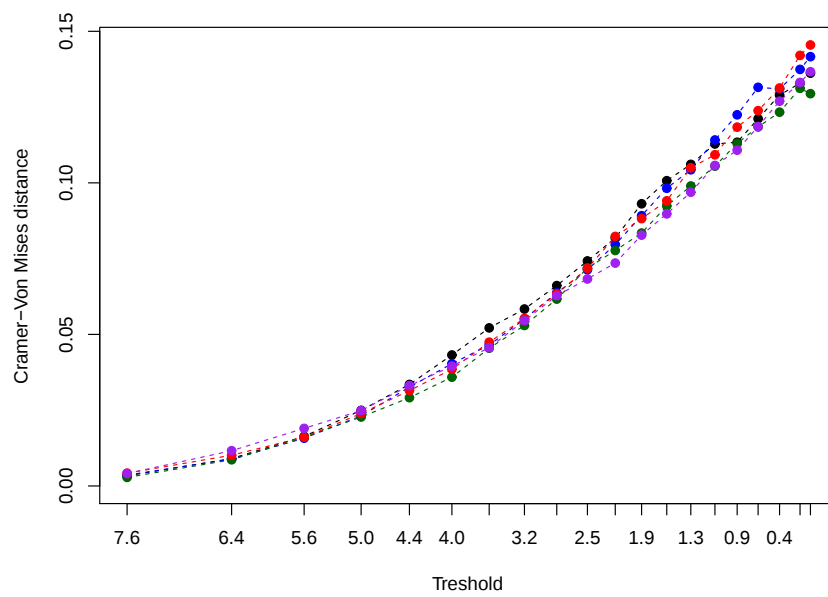


Figure 2.20: Hyperplane different spread test case, Cramer-Von Mises distance

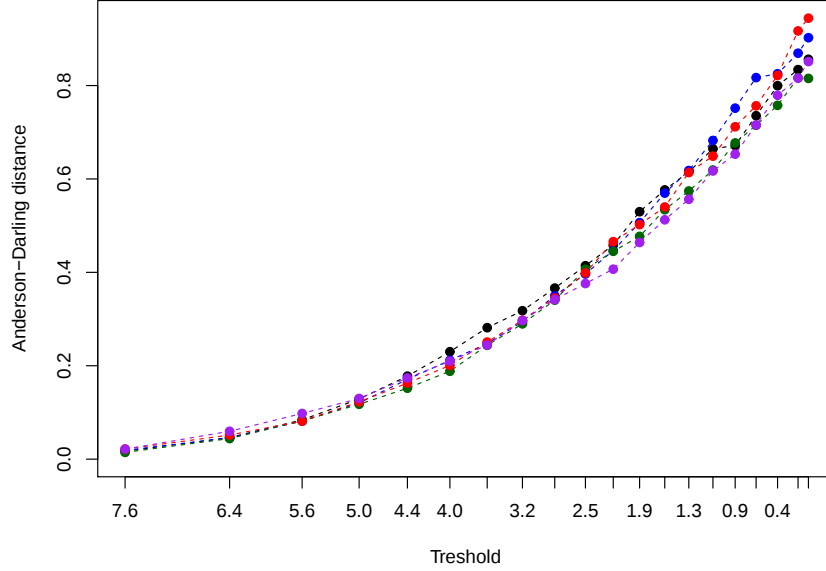


Figure 2.21: Hyperplane different spread test case, Anderson-Darling distance

Every distance growth with the threshold. All the distances pack the inputs variable together. So far, we can conclude that this SA method succeeds in giving to each equally contributing variable the same importance, despite their different spread.

2.3.3.5 Thresholded Ishigami test case

This more complex numerical example is described in Appendix B.2.

Subset estimation The failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the failure probability is roughly $P_f = 5.89 \times 10^{-3}$. The algorithm's parameters are the following:

- the proposal density is a truncated Gaussian centred on the particle, with variance 1, minimum and maximum respectively $-\pi$ and π ,
- $N = 10^4$, $\alpha = .75$.

The result with $18 \times N$ function calls is close from the exact result:

$$\hat{P} = 5.81 \times 10^{-3}$$

Distance estimation The distances are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.22, 2.23 and 2.24. A different color is used for every variable: X_1 is plotted in black, X_2 in blue and X_3 in red.

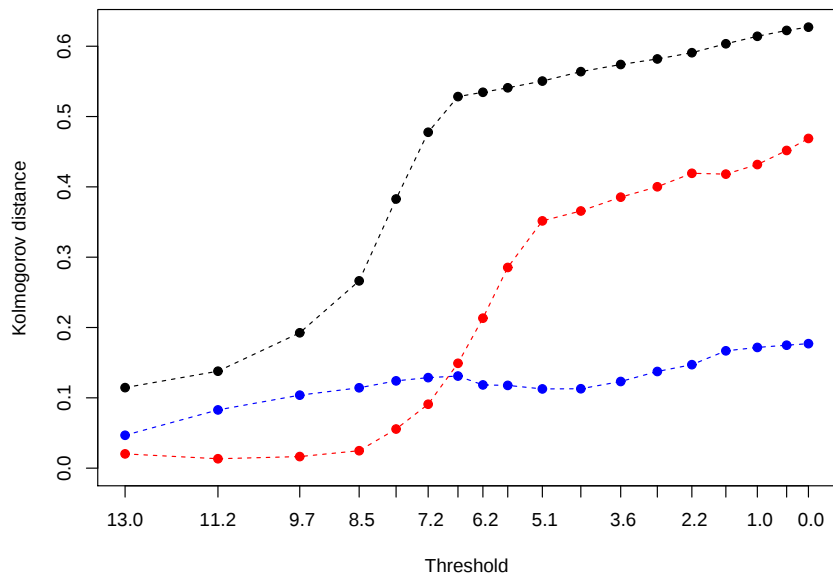


Figure 2.22: Thresholded Ishigami test case, Kolmogorov distance

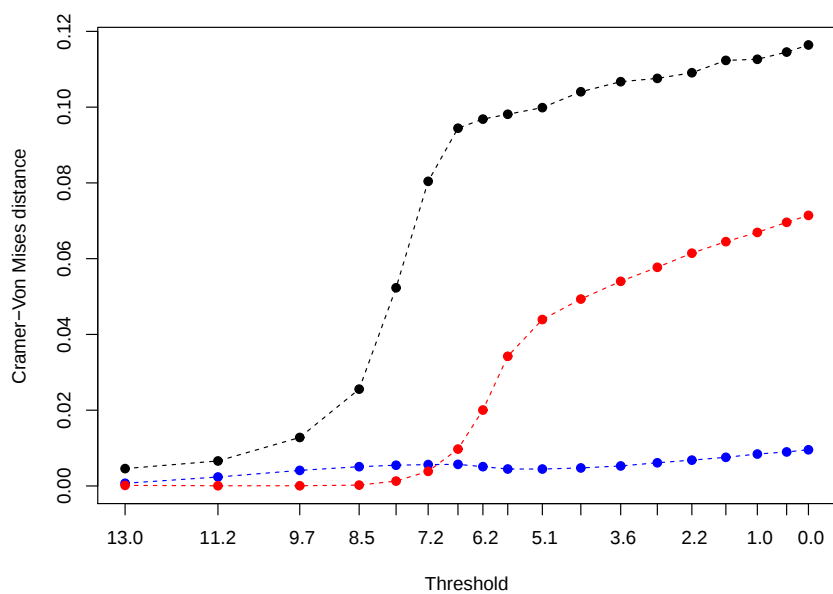


Figure 2.23: Thresholded Ishigami test case, Cramer-Von Mises distance

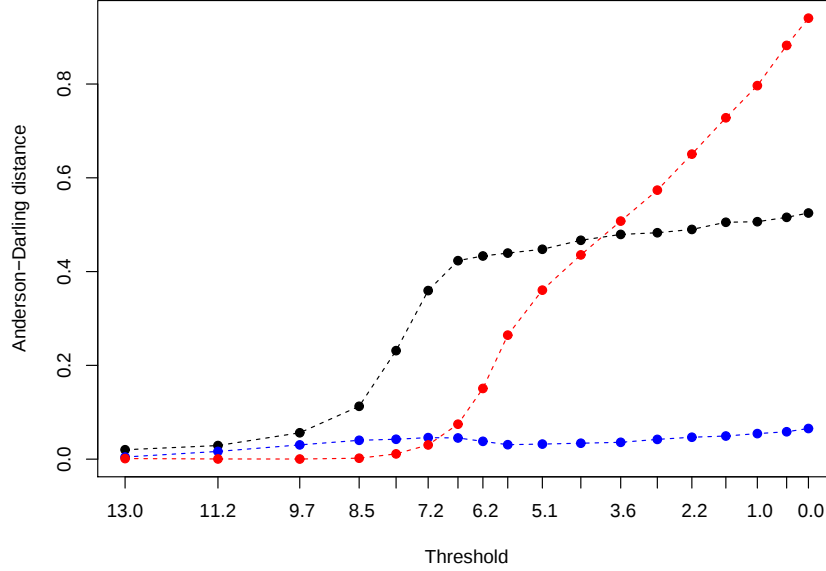


Figure 2.24: Thresholded Ishigami test case, Anderson-Darling distance

One can first comment that there is a non linearity in the growth of the distance for X_1 , for the three considered distances. Specifically, there is a raise in the growth between the threshold 8.5 and 6.7. For the three distances, there is a crossing of the values of the indices of X_2 and X_3 . Considering the Anderson-Darling distance, there is also a crossing between X_2 and X_1 .

On this test case, the 3 distances do not give equivalent results. Precisely, Kolmogorov and Cramer-Von Mises distances give the same final ranking (X_1, X_3, X_2); although the gap between variables X_1 and X_3 is larger with Kolmogorov distance. However, Anderson-Darling gives the ranking (X_3, X_1, X_2). We propose the following explanation: Anderson-Darling distance (being a re-weighting of Cramer-Von Mises distance) is said to grant more weight to the extremes. But in the final step of the subset, the third marginal of the sample of failure points consists in points distributed on the extrema (close of $-\pi$ and π). Notice that all the distances give variable X_2 as the less influential variables.

2.3.3.6 Flood test case

This numerical example emulating a real code is described in Appendix B.3.

Subset estimation The failure probability P_f is estimated using the adaptive subset simulation method (see Section 1.2.3). Recall that the failure probability is roughly $P_f = 7.88 \times 10^{-4}$. The algorithm's parameters are the following:

- the proposal density is always centred on the actual particle, and the densities are:
 - a truncated Gaussian with minimum 0 and standard deviation 10 for variable Q ;
 - a truncated Gaussian with minimum 1 and standard deviation 5 for variable K_s ;

- a truncated Gaussian with minimum 49, maximum 51 and standard deviation 1 for variable Z_v ;
- a truncated Gaussian with minimum 54, maximum 56 and standard deviation 1 for variable Z_m ;

- $N = 10^4$, $\alpha = .75$.

The result with $26 \times N$ function calls is close from the exact result:

$$\hat{P} = 7.07 \times 10^{-3}$$

Distance estimation The distances are estimated with the formulas given in 2.3.2. They are plotted in function of the threshold in Figures 2.25, 2.26 and 2.27. A different color is used for every variable: Q is plotted in black, K_s in blue, Z_v in green and Z_m in red.

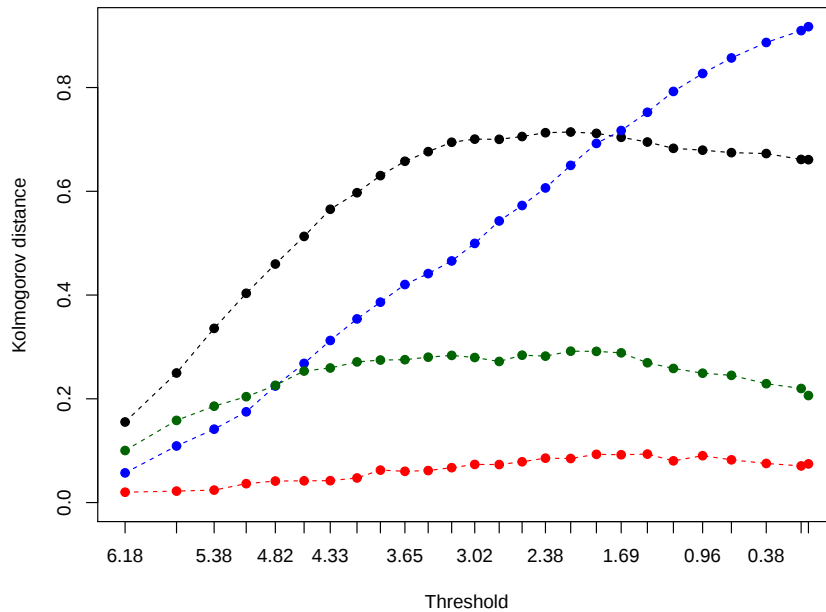


Figure 2.25: Flood test case, Kolmogorov distance

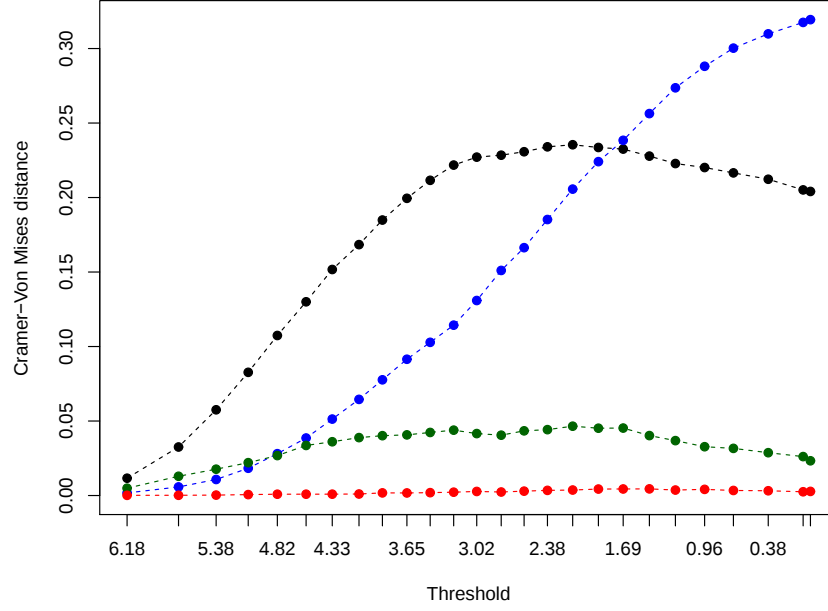


Figure 2.26: Flood test case, Cramer-Von Mises distance

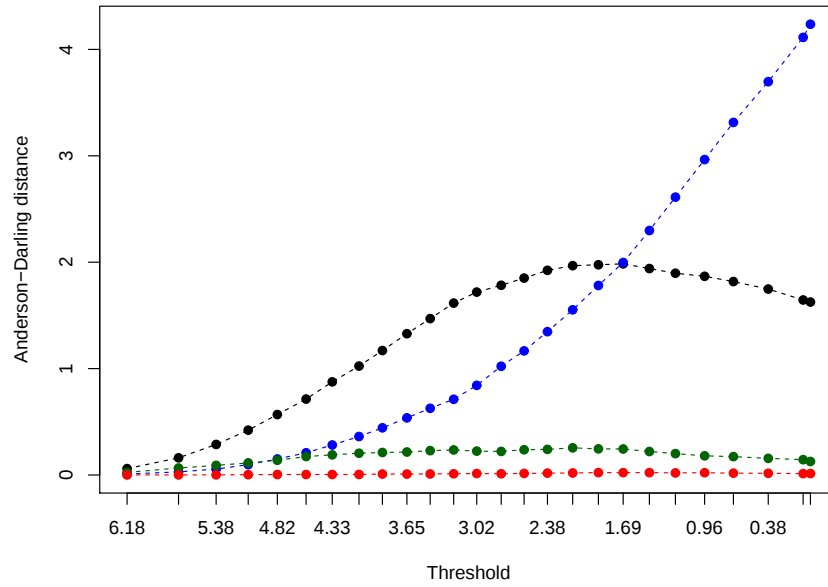


Figure 2.27: Flood test case, Anderson-Darling distance

On this test case, the 3 distances give equivalent results. The behaviour of variable Q is the same with the 3 distances: the distance between the original c.d.f. and the empirical one rises from the

beginning until the threshold reaches 2. Then the distance diminishes slowly. The behaviour is the same for variable Z_v although with much less amplitude. The distance for variable K_s grows with the subset. For variable Z_m , the distance stagnates around the minimal value. The final ranking is K_s, Q, Z_v, Z_m for the 3 distances, which is the one provided by the importance factors (see Table 3.16).

2.3.4 Conclusion

- The proposed SA technique allows an use of the subset simulation methods. In particular, we used adaptive levels algorithms.
- The computational time is negligible with respect to the computational time needed to obtain the failure sample.
- The three proposed distances bring complementary informations on the failure sample.
 - Kolmogorov distance is an L_∞ one. It expresses the maximal gap between the empirical c.d.f. of the failure sample and the original distribution. As far as we have noticed on the examples, it seems the more discriminant distance (see Figure 2.16 for instance).
 - Cramer-Von Mises distance is an L_2 one. The indices produced using this distance answer the question "what is the input which distribution varies most in central tendency when restricted to the failure domain?". The use of such a distance is then recommended if the aim of the SA is to fix the non-influential input variables to their central value.
 - Anderson-Darling distance grants more weight to the extreme values. The indices produced using this distance answer the question "what is the input which distribution varies most in the extremes when restricted to the failure domain?". The use of such a distance is recommended when the aim of the SA is to determine the relative influence of the boundaries or extremes of input distributions.
- So far, this SA method is recommended to get a similar information as the one provided by the Sobol' indices on the failure indicator (that is to say the detection of variables less influential than others).
- However, this method provides an interesting additional information: it shows how the threshold impacts each variable. This is interesting in the sense that, in some real cases, the threshold might not be fixed by the physics but by the regulation. A threshold given for a safety study might not be the same for another study. This method has shown (on the Ishigami test case) that the ranking might be different for several threshold (crossing of the curves between X_2 and X_3 for instance).

2.4 Synthesis

This chapter has presented two SA methods provide a variable ranking (objective 1, REM1, see Section 1.7). A first part was devoted to classification methods for SA, with a special attention paid to random forests. A second part was devoted to measuring the departure between the original and the empirical c.d.f. at several steps of a subset simulation method.

Table 2.10 is a short synthesis on the SA methods presented throughout this chapter.

Indice	Sensitivity type	Evaluation method	Pros/Cons
Gini indices	Global	• Random forests on a MC sample	+ By-product of the MC method – Can affect a non-null importance to a non-influential variable
MDA indices	Global	• Random forests on a MC sample	+ By-product of the MC method
Indices using the cdf departure $\delta_i^{SS}(A_k)$	Global	• Subset simulation technique	+By-product of a subset simulation technique –Information is neglected when working on the marginals.

Table 2.10: Synthesis on the presented SA methods

However this chapter provides some avenues for future research:

- An adapted reflection must be conducted on the pertinence of the random forests' sensitivity measures when using importance sampling.
- Still in the context of random forests, the MDA indices when using subset simulation must be implemented.
- The idea that consist in getting a model collection and aggregating their sensitivity measures to insure stability seems promising and is to be explored.
- When dealing with the second method proposed, a work including the copula theory might be conducted. In particular, the aim of this work could be to quantify the total departure of the particle cloud, and to assess which variable or interaction of variables contribute most to the failure event.

Chapter 3

Density Modification Based Reliability Sensitivity Indices

3.1 Introduction and overview

In most studies, sensitivity indices for failure probabilities are defined in strong correspondence with a given method of estimation (e.g. Lemaire [61], Munoz Zuniga *et al.* [73]). Their interpretation is consequently limited. In this chapter, it is proposed to define new generic sensitivity reliability indices. Our sensitivity index is based upon input density modification, and is adapted to failure probabilities. A methodology to estimate such indices is derived.

The proposed indices reflect the impact of the input density modification on the failure probability P_f . The indices are independent of the perturbation in the sense that the practitioner can set the perturbation adapted to his/her problem. Different modifications/perturbations will answer different problems.

For simplicity reasons, a classical Monte Carlo framework is considered in the following, although the estimation process will be extended to the use of subset and importance sampling methods. The sensitivity index can be computed using the sole set of simulations that has already been used to estimate the failure probability P_f , thus limiting the number of calls to the numerical model, as specified in the constraints of the CWNR case (page 24)

The outline of this chapter is the following: first, the indices and their theoretical properties are presented in Section 3.2, altogether with the estimation methodology. Second, Section 3.3 deals with several perturbation methodologies. These perturbations can be classified into two main families: Kullback-Leibler minimization methods and parameter perturbations methods. The behaviour of the indices is examined in Section 3.4 through numerical simulations in various complexity settings (see Appendix B). Comparisons with two reference sensitivity analysis methods (FORM's importance factors and Sobol' indices, see Section 1.3) highlight the relevance of the new indices in most situations. In Section 3.5, it is proposed to improve the DMBRSI estimation with importance sampling and with subset simulation. The main advantages and remaining issues are finally discussed in the last section of the chapter, that introduces avenues for future research.

This chapter is the extended version of the paper [63].

3.2 The indices: definition, properties and estimation

3.2.1 Definition

Given a unidimensional input variable X_i with pdf f_i , let us call $X_{i\delta} \sim f_{i\delta}$ the corresponding perturbed random input. This perturbed input takes the place of the real random input X_i , in a sense of modelling error : what if the correct input were $X_{i\delta}$ instead of X_i ? More about this replacement is proposed thereafter, see Section 3.3.1.1. Recall that we consider that $(X_1, \dots, X_d)$ are mutually independent.

The perturbed failure probability becomes:

$$P_{i\delta} = \int \mathbf{1}_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(\mathbf{x}) d\mathbf{x} \quad (3.1)$$

where x_i is the i^{th} component of the vector $\mathbf{x}$. Independently of the mechanism chosen for the perturbation (see next section for proposals), a good sensitivity index $S_{i\delta}$ should have intuitive features that make it appealing to reliability engineers and decision-makers. We argue that the following definition can fulfill these requirements.

Definition 3.2.1 *Define the Density Modification Based Reliability Sensitivity Indices (DMBRSI) as the quantity $S_{i\delta}$:*

$$S_{i\delta} = \left[\frac{P_{i\delta}}{P_f} - 1 \right] \mathbf{1}_{\{P_{i\delta} \geq P_f\}} + \left[1 - \frac{P_f}{P_{i\delta}} \right] \mathbf{1}_{\{P_{i\delta} < P_f\}} = \frac{P_{i\delta} - P_f}{P_f \cdot \mathbf{1}_{\{P_{i\delta} \geq P_f\}} + P_{i\delta} \cdot \mathbf{1}_{\{P_{i\delta} < P_f\}}}.$$

3.2.2 Properties

- Firstly, $S_{i\delta} = 0$ if $P_{i\delta} = P_f$, as expected if X_i is a non-influential variable or if δ expresses a negligible perturbation.
- Secondly, the sign of $S_{i\delta}$ indicates how the perturbation impacts the failure probability qualitatively. It highlights the situations when $P_{i\delta} > P_f$ i.e. if the remaining (*epistemic*) uncertainty on the modelling $X_i \sim f_i$ can increase the failure risk. In this case, the uncertainty on the concerned variable should be more accurately analysed. Conversely, if $P_{i\delta} < P_f$, P_f can be interpreted as a conservative assessment of the failure probability, with respect to variations of X_i . In such a case, deeper modelling studies on X_i appear less essential.
- Thirdly, given its sign, the absolute value of $S_{i\delta}$ has simple interpretation and provides a level of the conservatism or non-conservatism induced by the perturbation. A value of $\alpha > 0$ for the index means that $P_{i\delta} = (1 + \alpha)P_f$. If $S_{i\delta} = -\alpha < 0$ then $P_{i\delta} = (1/(1 + |\alpha|))P_f$.

3.2.3 Estimation

The postulated ability of $S_{i\delta}$ to enlighten the sensitivity of P to input perturbations must be tested in concrete cases (see Section 3.4), when an estimator $\hat{P}_N$ of P_f can be computed using an already available design of N numerical experiments. In the following, N is assumed to be large enough such that statistical estimation stands within the framework of asymptotic theory. Besides, a standard Monte Carlo design of experiments is assumed for simplicity (see Section 1.2.1). This allows to write:

$$\hat{P}_N = \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}}$$

where the $\mathbf{x}^1, \dots, \mathbf{x}^N$ are independent realisations of X . The strong Law of Large Numbers (LLN) and the Central Limit Theorem (CLT) ensure that for almost all realisations $\hat{P}_N \xrightarrow[N \rightarrow \infty]{} P_f$ and

$$\sqrt{\frac{N}{P_f(1-P_f)}}(\hat{P}_N - P_f) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, 1). \quad (3.2)$$

The Monte Carlo framework allows $P_{i\delta}$ to be consistently estimated without new calls to G , through a "reverse" importance sampling mechanism:

$$\hat{P}_{i\delta N} = \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)}. \quad (3.3)$$

This property holds in the more general case when P is originally estimated by importance sampling rather than simple Monte Carlo, which is more appealing when G is time-consuming, Beckman and McKey, Hesterberg [8, 45]. This generalization is discussed further in the text (Section 3.5). The following lemma ensures the asymptotic behaviour of such an estimator.

Lemma 3.2.1 *Assume the usual conditions*

$$(i) \text{ } \text{Supp}(f_{i\delta}) \subseteq \text{Supp}(f_i),$$

$$(ii) \int_{\text{Supp}(f_i)} \frac{f_{i\delta}^2(x)}{f_i(x)} dx < \infty,$$

then $\hat{P}_{i\delta N} \xrightarrow[N \rightarrow \infty]{} P_{i\delta}$ and $\sqrt{N}\sigma_{i\delta N}^{-1}(\hat{P}_{i\delta N} - P_{i\delta}) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, 1)$. The exact expression of $\sigma_{i\delta N}^{-1}$ is given in Appendix D.1, equation (D.1). It can be consistently estimated by

$$\hat{\sigma}_{i\delta N}^2 = \frac{1}{N} \sum_{n=1}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \left(\frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)} \right)^2 - \hat{P}_{i\delta N}^2.$$

The proof of this Lemma is given in Appendix D.1.

We stress that Equation 3.3 is valid as long as the assumptions of Lemma 3.2.1 are respected. This means that whatever the perturbation chosen, the estimation of $\hat{P}_{i\delta N}$ does not require new function calls.

The asymptotic properties of any estimator of $S_{i\delta}$ will depend on the correlation between $\hat{P}_N$ and $\hat{P}_{i\delta N}$. The next proposition summarizes the features of the joint asymptotic distribution of both estimators.

Proposition 3.2.1 *Under assumptions (i) and (ii) of Lemma 3.2.1,*

$$\sqrt{N} \left[\begin{pmatrix} \hat{P}_N \\ \hat{P}_{i\delta N} \end{pmatrix} - \begin{pmatrix} P_f \\ P_{i\delta} \end{pmatrix} \right] \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}_2(0, \Sigma_{i\delta})$$

where $\Sigma_{i\delta}$ is given in Appendix D.1, Equation (D.2) and can be consistently estimated by

$$\hat{\Sigma}_{i\delta} = \begin{pmatrix} \hat{P}_N(1 - \hat{P}_N) & \hat{P}_{i\delta N}(1 - \hat{P}_N) \\ \hat{P}_{i\delta N}(1 - \hat{P}_N) & \hat{\sigma}_{i\delta N}^2 \end{pmatrix}.$$

The proof of this Proposition is given in Appendix D.1.

Given $(\hat{P}_N, \hat{P}_{i\delta N})$, the plugging estimator for $S_{i\delta}$ is:

$$\hat{S}_{i\delta N} = \left[\frac{\hat{P}_{i\delta N}}{\hat{P}_N} - 1 \right] \mathbf{1}_{\{\hat{P}_{i\delta N} \geq \hat{P}_N\}} + \left[1 - \frac{\hat{P}_N}{\hat{P}_{i\delta N}} \right] \mathbf{1}_{\{\hat{P}_{i\delta N} < \hat{P}_N\}}. \quad (3.4)$$

In corollary of Proposition 3.2.1, applying the continuous-mapping theorem to the function $s(x, y) = \left[\frac{y}{x} - 1 \right] \mathbf{1}_{\{y \geq x\}} + \left[1 - \frac{x}{y} \right] \mathbf{1}_{\{y < x\}}$, $\hat{S}_{i\delta N}$ converges almost surely to $S_{i\delta}$. The following CLT results from Theorem 3.1 in Van der Vaart [98].

Proposition 3.2.2 *Assume that assumptions (i) and (ii) of Lemma 3.2.1 hold and further that $P \neq P_{i\delta}$, we have*

$$\sqrt{N} \left[\hat{S}_{i\delta N} - S_{i\delta} \right] \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, d_s^T \Sigma d_s) \quad (3.5)$$

with $d_s = \left(\frac{\partial s}{\partial x}(P_f, P_{i\delta}), \frac{\partial s}{\partial y}(P_f, P_{i\delta}) \right)^T$ for $x \neq y$, and

$$\begin{aligned} \frac{\partial s}{\partial x}(x, y) &= -y \mathbf{1}_{\{y \geq x\}}/x^2 - \frac{1}{y} \mathbf{1}_{\{y < x\}}, \\ \frac{\partial s}{\partial y}(x, y) &= \frac{1}{x} \mathbf{1}_{\{y \geq x\}} + x \mathbf{1}_{\{y < x\}}/y^2. \end{aligned}$$

This holds when $P_f = P_{i\delta}$. Indeed, one has for $x^* \neq 0$:

$$\lim_{\substack{y \geq x \\ \begin{pmatrix} x \\ y \end{pmatrix} \rightarrow \begin{pmatrix} x^* \\ y^* \end{pmatrix}}} \nabla s(x, y) = \lim_{\substack{y < x \\ \begin{pmatrix} x \\ y \end{pmatrix} \rightarrow \begin{pmatrix} x^* \\ y^* \end{pmatrix}}} \nabla s(x, y) = \left(-\frac{1}{x^*}, \frac{1}{x^*} \right)^T.$$

3.2.4 Framework

Figure 3.1 summarises the use of DMBRSI.

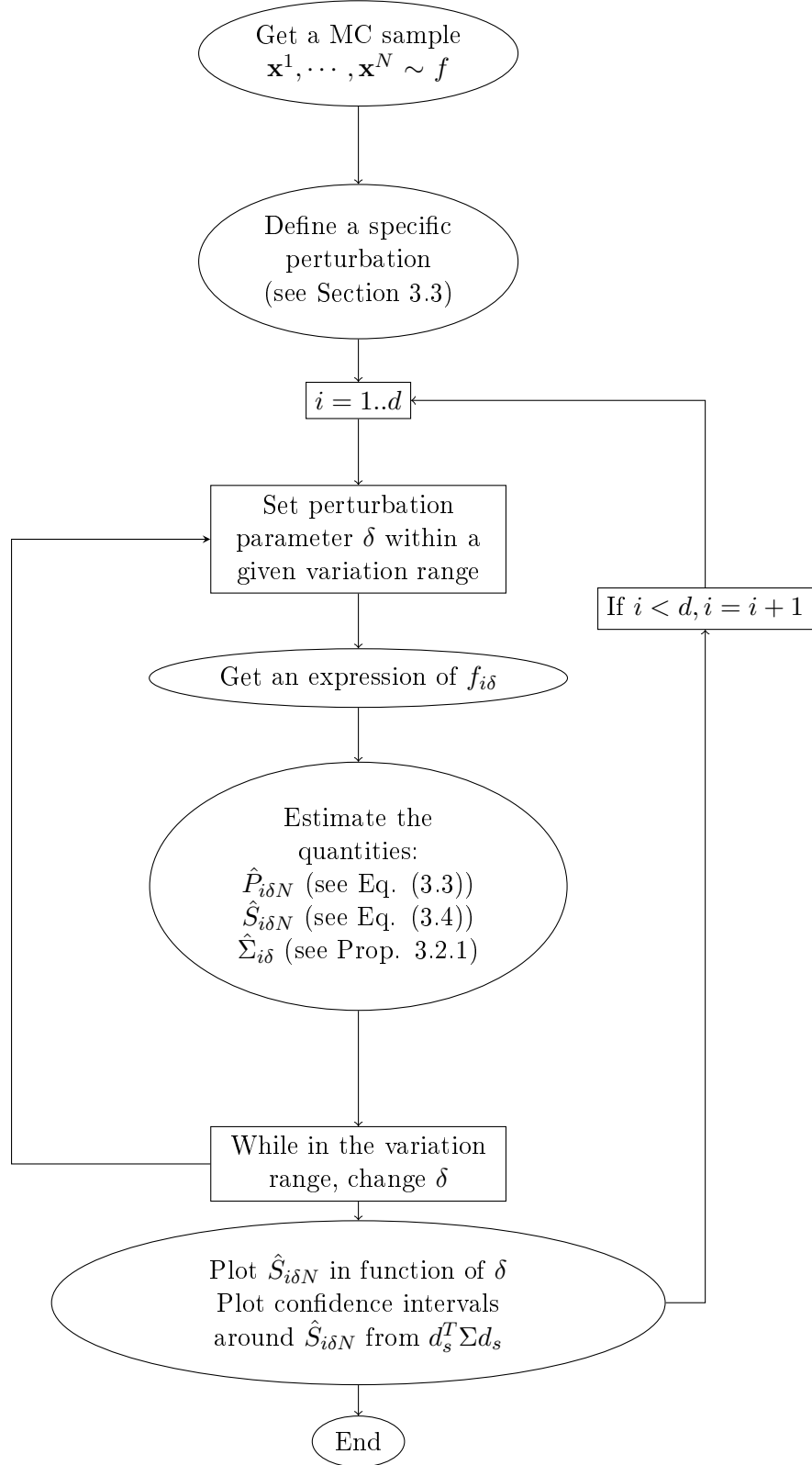


Figure 3.1: General DMBRSI framework

The notion of perturbation is discussed in the next section. Which perturbation to choose according to the objective is also discussed as well as recommendations on the variation range of δ .

3.3 Methodologies of input perturbation

This section proposes several perturbation methodologies. However the DMBRSI and its estimation techniques remain valid for any perturbation, as long as the support constraints (Lemma 3.2.1) are respected. Here, two main families of method are presented. The first one determines the perturbed density minimizing the Kullback-Leibler divergence under some constraints given by the practitioner. Several constraints are proposed, each one dealing with a different SA objective. The second method is to be used when the practitioner wants to test the sensitivity of P_f to the parameters of the distributions. Both subsections will be introduced by toy-examples.

This section illustrates the DMBRSI's capacity to deal with several SA objectives. The practitioner is invited to propose new perturbation methodologies that would answer his questions. Recommendations of perturbation regarding the objectives are itemized at the end of the section.

3.3.1 Kullback-Leibler minimization

The DMBRSI requires to define a perturbation for each input. In general, and especially in preliminary reliability studies, there is no prior rule allowing to elicit a specialized perturbation for each input variable. Thus a simple perturbation methodology is exposed -denoted KLM for Kullback-Leibler minimization- allowing the practitioner to answer the questions itemized in Section 1.7 of the present thesis.

3.3.1.1 First example

Let us assume we have an input X_i distributed according to f_i . This random input models for instance a physical uncertain quantity. The distribution f_i is known, altogether with its parameters. This modelling was done by physic expert, engineers, practitioners, statistical analyst from field data ... Moments of X_i are also known given they exist.

We would like to fairly perturb this input to represent "the lack of certitude" on some quantity. This quantity might be, as a simple example, the first moment. Let us assume the input X_i is distributed according to a Gaussian, $\mathcal{N}(0, 1)$. What if the expectation of X_i was badly modelled? What if the data used to calibrate f_i were wrong?

We will thus suppose the existence of *another* random variable $X_{i\delta}$ (distributed according to $f_{i\delta}$), close from X_i in some sense, and we will process it through the model, as if input X_i was replaced by the perturbed input $X_{i\delta}$. δ represents here the perturbation, its amplitude for instance.

Thus the example is an expectation perturbation. What if the mean of the perturbed input were 2? New data can lead to such a situation. So we want the new input to have:

$$\mathbb{E}[X_{i\delta}] = 2, \tag{3.6}$$

obviously

$$\int f_{i\delta}(x)dx = 1 \tag{3.7}$$

and $X_{i\delta}$ must be close in some sense to X_i . Notice that Equation (3.6) rewrites

$$\int x f_{i\delta}(x)dx = 2. \tag{3.8}$$

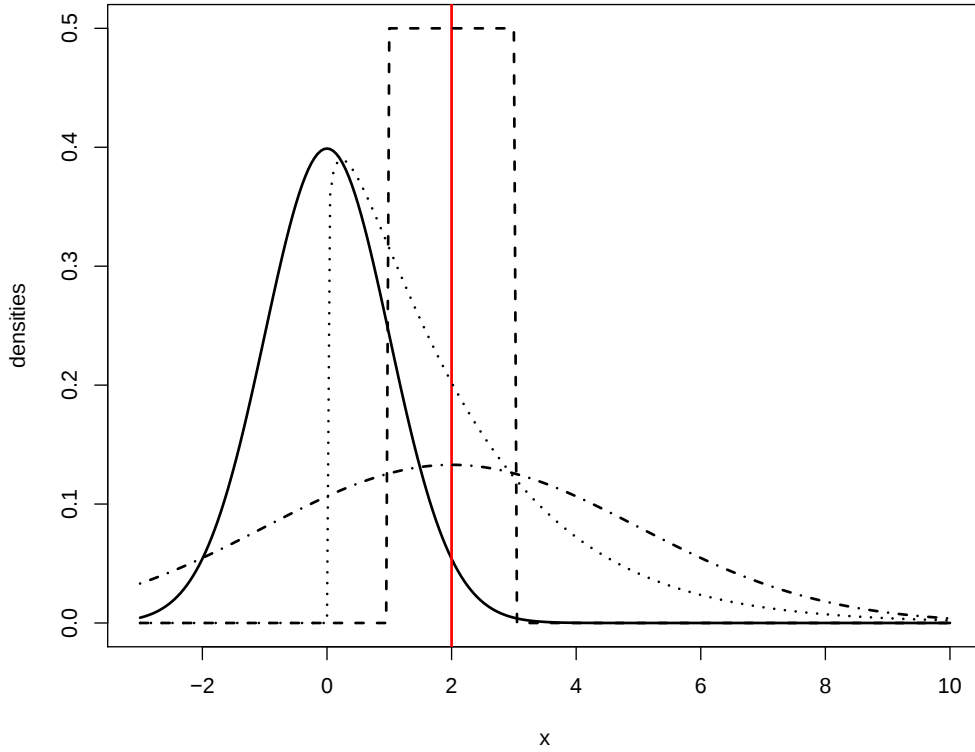


Figure 3.2: The original density of mean 0 (full line) and several candidates densities of mean 2

Several candidates for $f_{i\delta}$ exist. In Figure 3.2 are plotted some choices, altogether with the original density. Some candidates are "closer" to f_i than others in some sense not yet defined. Let us now focus on the needs. We would like to take $f_{i\delta}$ as the density, among all the densities satisfying the constraints (in our example, constraint (3.6)), that is the minimum argument of a departure D between densities.

$$f_{i\delta} = \underset{f_{\text{mod}} | \text{constraints holds}}{\operatorname{argmin}} D(f_{\text{mod}}, f_i) \quad (3.9)$$

Distance quantifying the departure between two densities are numerous (Cha [24]). Information-theoretical arguments (Cover and Thomas [25]) led to choose the Kullback-Leibler divergence (KLD) between $f_{i\delta}$ and f_i as a measure of the discrepancy to minimize under constraints (definition of KLD is reminded in 3.10). This comes at "adding" as few information as possible on $f_{i\delta}$ other than the constraints.

By simple calculus, it may be shown that the density minimizing the KLD from f_i and satisfying constraints 3.6 is a Gaussian, of mean 2 and of the same variance as f_i . The computation of the indices expressed in Section 3.2 can now be done as $f_{i\delta}$ is provided.

Next subsection formalises this example.

3.3.1.2 Kullback-Leibler minimization

Here, a perturbed input density $f_{i\delta}$ is defined as the closest distribution to the original f_i in the entropic sense and under some constraints of perturbation.

Later (see Sections 3.3.1.3, 3.3.1.4), specific perturbations corresponding to a mean shift, a variance shift and a quantile shift will be presented.

Recall that between two pdf p and q we have:

$$KL(p, q) = \int_{-\infty}^{+\infty} p(y) \log \frac{p(y)}{q(y)} dy \text{ if } \log \frac{p(y)}{q(y)} \in L^1(p(y)dy). \quad (3.10)$$

Let $i = 1, \dots, d$, the constraints are expressed as follows in function of the modified density f_{mod} :

$$\int g_k(x_i) f_{\text{mod}}(x_i) dx_i = \delta_{k,i} \quad (k = 1 \dots K). \quad (3.11)$$

Here, for $k = 1, \dots, K$, g_k are given functions and $\delta_{k,i}$ are given real. These quantities will lead to a perturbation of the original density. The modified density $f_{i\delta}$ considered in our work is:

$$f_{i\delta} = \underset{f_{\text{mod}} | (3.11) \text{ holds}}{\operatorname{argmin}} KL(f_{\text{mod}}, f_i) \quad (3.12)$$

and the result takes an explicit form (Csiszar, [26]) given in the following proposition.

Proposition 3.3.1 *Let us define, for $\lambda = (\lambda_1, \dots, \lambda_K)^T \in \mathbb{R}^K$,*

$$\psi_i(\lambda) = \log \int f_i(x) \exp \left[\sum_{k=1}^K \lambda_k g_k(x) \right] dx, \quad (3.13)$$

where the last integral can be finite or infinite (in this last case $\psi_i(\lambda) = +\infty$). Further, set $\operatorname{Dom} \psi_i = \{\lambda \in \mathbb{R}^K | \psi_i(\lambda) < +\infty\}$. Assume that there exists at least one pdf f_{mod} satisfying (3.11) and that $\operatorname{Dom} \psi_i$ is an open set. Then, there exists a unique λ^ such that the solution of the minimisation problem (3.12) is*

$$f_{i\delta}(x_i) = f_i(x_i) \exp \left[\sum_{k=1}^K \lambda_k^* g_k(x_i) - \psi_i(\lambda^*) \right]. \quad (3.14)$$

The theoretical technique to compute λ is provided in Appendix D.2.

3.3.1.3 Moments shifting

Mean shifting The first moment is often used to parametrize a distribution. Thus the first perturbation presented here is a mean shift, that is expressed with a single constraint:

$$\int x_i f_{\text{mod}}(x_i) dx_i = \delta_i. \quad (3.15)$$

In terms of SA, this perturbation should be used when the user wants to understand the sensitivity of the inputs to a mean shift - that is to say "what if the mean of input X_i were δ_i instead of $\mathbb{E}[X_i]$ ". Notice that for most distributions, this amounts to testing the sensitivity to the central tendency.

Proposition 3.3.2 *Considering constraint (3.15), under the assumptions of Proposition 3.3.1, the expression of the optimal perturbed density is*

$$f_{i\delta_i}(x_i) = \exp(\lambda^* x_i - \psi_i(\lambda^*)) f_i(x_i) \quad (3.16)$$

where λ^ is such that Equation (3.15) holds.*

Notice that Equation (3.13) becomes

$$\psi_i(\lambda) = \log \int f_i(x_i) \exp(\lambda x_i) dx_i = \log (M_{X_i}(\lambda)) \quad (3.17)$$

where $M_{X_i}(u)$ is the moment generating function (m.g.f.) of the i -th input. With this notation, λ^* is such that:

$$\int x_i \exp(\lambda^* x_i - \log(M_{X_i}(\lambda^*))) f_i(x_i) dx_i = \delta_i ,$$

which leads to:

$$\int x_i \exp(\lambda^* x_i) f_i(x_i) dx_i = \delta_i M_{X_i}(\lambda^*) .$$

This can be simplified to:

$$\frac{M'_{X_i}(\lambda^*)}{M_{X_i}(\lambda^*)} = \delta_i . \quad (3.18)$$

This equation is easy to solve when the expression of the mgf of the input X_i and of its derivative is known.

Variance shifting In some cases, the expectation of an input may not be the main source of uncertainty. One might be interested in perturbing its second moment. This case may be treated considering a couple of constraints. The perturbation presented is a variance shift, therefore the set of constraints is:

$$\begin{cases} \int x_i f_{\text{mod}}(x_i) dx_i = \mathbb{E}[X_i] , \\ \int x_i^2 f_{\text{mod}}(x_i) dx_i = V_{\text{per},i} + \mathbb{E}[X_i]^2 . \end{cases} \quad (3.19)$$

The perturbed distribution has the same expectation $\mathbb{E}[X_i]$ as the original one and a perturbed variance $V_{\text{per},i} = \text{Var}[X_i] \pm \delta_i$. In terms of SA, for most distributions, this amounts to testing the sensitivity to the tails of the distribution, keeping the central tendency untouched.

Proposition 3.3.3 *Under the assumptions of Proposition 3.3.1, for constraint (3.19), the expression of the optimal perturbed density is:*

$$f_{i\delta_i}(x_i) = \exp(\lambda_1^* x + \lambda_2^* x^2 - \psi_i(\boldsymbol{\lambda}^*)) f_i(x_i)$$

where λ_1^* and λ_2^* are so that equation (3.19) holds.

Perturbation of Natural Exponential Family In general, when perturbing the input densities with the KLM method, the shape is not conserved. However in the specific case of Natural Exponential Family (NEF), the following proposition can be derived.

Proposition 3.3.4 *Assume that the original random variable X_i belongs to the NEF, i.e. its pdf can be written as:*

$$f_{i,\theta}(x_i) = b(x_i) \exp[x_i \theta - \eta(\theta)]$$

where θ is a parameter from a parametric space Θ , $b(\cdot)$ is a function that depends only of x_i and

$$\eta(\theta) = \log \int b(x) \exp[x_i \theta] dx_i$$

is the cumulant distribution function. Considering the assumptions of Proposition 3.3.1, the optimal pdfs proposed respectively in Proposition 3.3.2 and Proposition 3.3.3 are also distributed according to a NEF.

The proof comes from Theorem 3.1 in Csiszar [26]. The details of computation are given for a mean shift and a variance shift in Appendix D.3.

Some shapes As an example, the two kinds of perturbations previously presented are provided for two families of inputs (Gaussian and Uniform) in Figure 3.3. The perturbations are respectively a mean and variance increasing. It is noticeable (as proven in Proposition 3.3.4) that the shape is conserved for the Gaussian distribution when shifting the mean or the variance. On the other hand, when increasing its mean, the Uniform distribution is packed down on the right-hand boundary of its support. When increasing its variance, the density is packed down on both boundaries of its support.

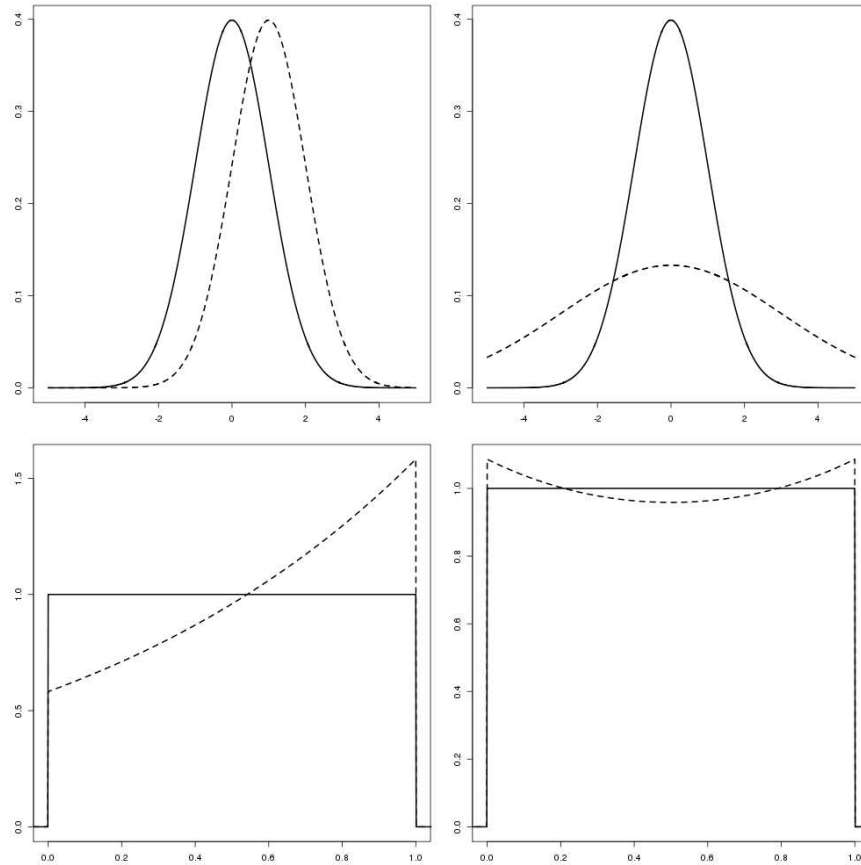


Figure 3.3: Mean shifting (left) and variance shifting (right) for Gaussian (upper) and Uniform (lower) distributions. The original distribution is plotted in solid line, the perturbed one is plotted in dashed line.

Some limitations, notion of equivalent perturbation In this paragraph, we focus on a mean shift but the same problems arise for a variance shift. What if two inputs do not have the same mean and we want to assess the impact of their mean shift on P_f ? How to conduct an equivalent perturbation on both inputs? Let us imagine an example in which an input has mean 0 and another has mean 100. If a perturbation is conducted on each variable separately, the interpretation is complicated as the ranges of variation will be separated. It is thus complicated or impossible to assess the impact of an equivalent perturbation. Conversely, it is impossible in this case to make a "relative mean shift" as one of the input has mean 0. The following solution is proposed for the mean perturbation: shift the mean relatively to the standard deviation, hence including the spread of the various inputs in their respective perturbation. So for any input, the original distribution

is perturbed so that its mean is the original's one plus δ times its standard deviation and the perturbation is conducted on δ (for instance ranging from -1 to 1). This solution is applied in the flood case (Section 3.4.7) where the inputs are not distributed according to the same density. However this solution might not be effective in every case, for instance when inputs do not have defined moments. This consideration led us to another kind of perturbation that we though is more equivalent: quantile shifting (see Section 3.3.1.4). Moreover in the following of this thesis, the perturbation will be conducted on the parameters of the input densities (see Section 3.3.2) but this falls outside of the KLM framework.

3.3.1.4 Quantile shifting

Based on the practitioner's experience, it has been noticed that the values of the input leading to the failure event seldom lies around the central tendency, but more in the extreme quantiles. From this point, another way to perturb the densities is proposed, keeping the KLM framework. Compared to the first two moment perturbations previously presented, we argue that this one seems more suitable to deal with inputs that are not identically distributed (see previous paragraph for a discussion on equivalent perturbations).

First example Let us first recall the definition of a quantile.

Definition 3.3.1 *For a given random variable X of probability density function f and of cumulative distribution function F , the α -quantile is the value q_α so that:*

$$\mathbb{P}(X < q_\alpha) = F(q_\alpha) = \int_{-\infty}^{q_\alpha} f(x)dx = \alpha \quad (3.20)$$

Then consider a random variable, modelling for instance an unknown physical phenomena value, defined as a standard Gaussian. Its 5% quantile or 5th percentile is $q_{5\%} = -1.64$.

As far as we noticed, in most cases, the values of the input leading to the failure event comes from the tails of the input distributions. What if these tails were badly modelled? Therefore a perturbation based on the quantiles is proposed.

In this first toy example, the aim is to increase the weight of the left tail. That is to say that the value $q_{5\%}$ is wished to become for the modified density, for instance the 7% quantile. This can be written:

$$\int 1_{]-\infty; q_{5\%}]}(x) f_{\text{mod}}(x) dx = 7\% \quad (3.21)$$

In Figure 3.4 are plotted the regular (black) and the perturbed (blue) densities. The shaded areas worth respectively $\int_{-\infty}^{q_{0.05}} f(x)dx = 0.05$ in grey and $\int_{-\infty}^{q_{0.05}} f_\delta(x)dx = 0.07$ in blue. One can remark that there is no longer a conservation of the shape with such a perturbation, since f_δ is not Gaussian. Additionally, the density is no longer continuous.

In a similar way, one could decide to perturb the densities in such a way that the tail is less weighted, meaning that the extreme values become less frequent. For instance, it can be written:

$$\int 1_{]-\infty; q_{5\%}]}(x) f_{\text{mod}}(x) dx = 3\% \quad (3.22)$$

meaning that the 5% quantile becomes the 3% quantile. The regular and the perturbed densities are pictured in Figure 3.4. A discontinuity at $q_{5\%}$ is present.

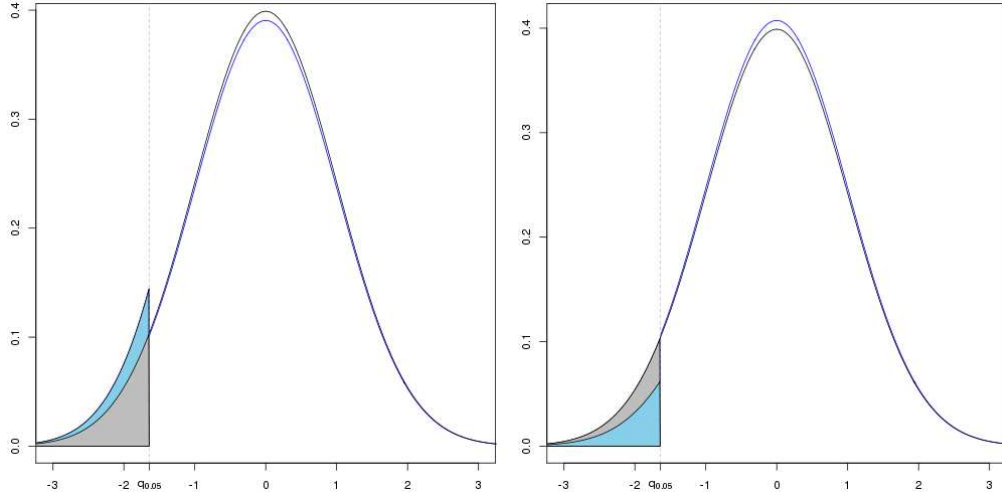


Figure 3.4: Standard Gaussian and perturbed density: quantile increase (left) and quantile decrease (right)

Methodology of input perturbation Let us denote by q_r the reference quantile, e.g. the value such that:

$$\int_{-\infty}^{q_r} f(x)dx = r, \quad 0 < r < 1 \quad (3.23)$$

The constraint is:

$$\int_{-\infty}^{q_r} f_{\text{mod}}(x)dx = \delta, \quad (3.24)$$

meaning that f_{mod} is the density such that its δ -quantile is q_r . Equivalently, the constraint can be written in the general fashion defined in Section 3.3.1.2, Equation 3.11:

$$\int 1_{]-\infty; q_r]}(x) f_{\text{mod}}(x)dx = \delta \quad (3.25)$$

Proposition 3.3.5 *Under the assumptions of Proposition 3.3.1, and under the constraint 3.25, the expression of the corresponding perturbed density is:*

$$f_{\delta}(x) = f(x) \exp [\lambda^* 1_{]-\infty; q_r]}(x) - \psi(\lambda^*)] \quad (3.26)$$

with

$$\psi(\lambda) = \log \left(\int f(x) \exp [\lambda^* 1_{]-\infty; q_r]}(x)] dx \right) \quad (3.27)$$

and λ^* is a real number such that (3.25) holds.

Some shapes In Figure 3.5 are displayed the original (solid black) and perturbed (dashed blue) pdf for the following families: Uniform, Triangle and Truncated Gumbel. The parameters used for these variables are the ones from the flood case (Appendix B.3.). In each case, the perturbation is:

$$\int_{-\infty}^{q_{0.05}} f_{\text{mod}}(x)dx = 0.07, \quad (3.28)$$

that is to say increasing the weight of the left-hand tail from 5% to 7%.

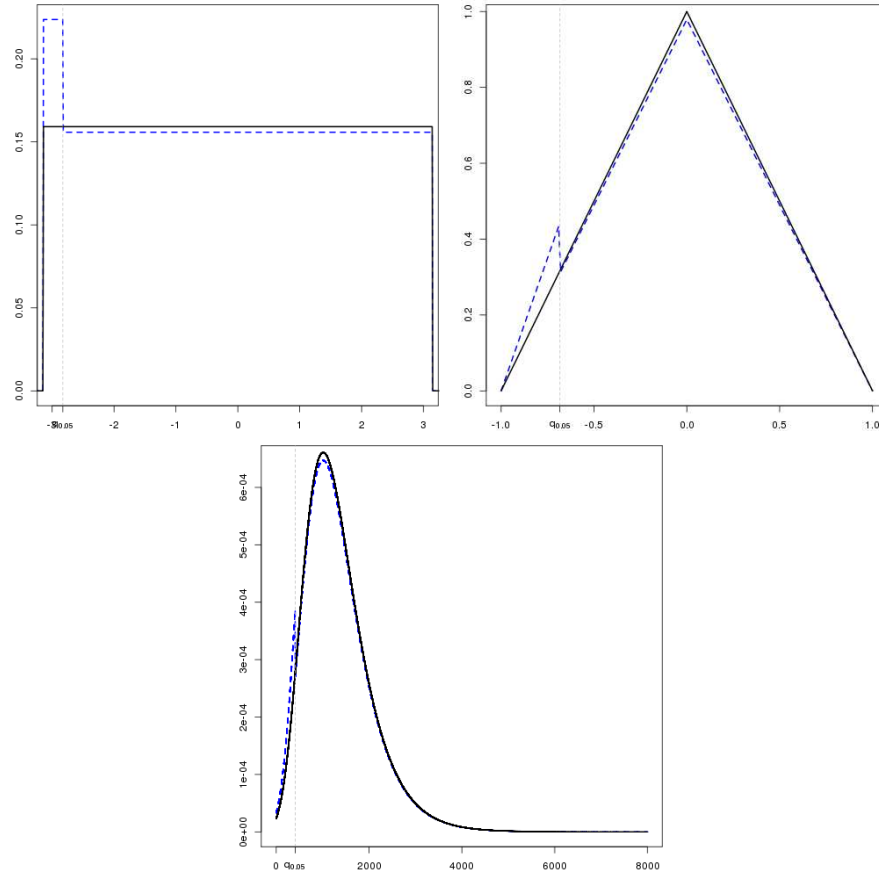


Figure 3.5: Uniform, Triangle and Truncated Gumbel pdf: quantile increase

3.3.2 Parameters perturbation

3.3.2.1 First example

Problem Assume that we have an input distribution, characterized by its parameters which are data-driven. The question of interest is "*how does a parametrisation error affects the failure probability ?*". To do so, the use of the DMBRSI is proposed - although the moments perturbations might not answer the question. Specifically, a perturbation based on the parameters is proposed. The indices are then plotted in function of the departure in a given divergence (Hellinger, Definition 3.3.3). Let us first illustrate the idea on a first example.

The input distributions and the model For the sake of clarity the Weibull distribution expression (Rinne [83]) is reminded here:

Definition 3.3.2 A random variable X has a three-parameters Weibull distribution if its pdf, defined on $\mathbb{R}^+$ is:

$$f(x|a, b, c) = \frac{c}{b} \left(\frac{x-a}{b} \right)^{c-1} \exp \left[- \left(\frac{x-a}{b} \right)^c \right]$$

where parameter a , defined on $\mathbb{R}$ in the same unit as x , is called the origin. It is a location parameter. The second parameter b is defined on $\mathbb{R}^+$ in the same unit as x and is called the scale parameter. The third parameter c bears no dimension, is defined on $\mathbb{R}^+$ and is called the shape parameter.

The expectation of such a random variable writes:

$$\mathbb{E}[W_{a,b,c}] = a + b\Gamma\left(1 + \frac{1}{c}\right),$$

and the variance is:

$$\text{Var}[W_{a,b,c}] = b^2 \left(\Gamma\left(1 + \frac{2}{c}\right) - \Gamma\left(1 + \frac{1}{c}\right)^2 \right),$$

where Γ is the Gamma function. In the following it will be stated that $a = 0$ and this location parameter will be omitted.

For this first example, an input is distributed according to a Weibull distribution and another input is distributed according to a standard Gaussian. Assume that the failure model is:

$$G(\mathbf{X}) = G(X_1, X_2) = \frac{1}{2}X_1 + \frac{1}{10}X_2 + 1.5$$

where $X_1 \sim \mathcal{N}(\mu, \sigma)$ and $X_2 \sim W(b, c)$ with $\mu = 0$, $\sigma = 1$, $b = 1.5$ and $c = \pi$. The failure probability is roughly $\hat{P} = 4.8 \times 10^{-3}$.

Use of DMBRSI for sensitivity to the parameters Let us assume that the practitioner is interested in testing the sensitivity of its model to the parameters of the distributions. When dealing with the Gaussian input, a perturbation of the 2 first centred moments is equivalent to a perturbation of the parameters (see Section 3.3.1.3). On the other hand, perturbing the moments of a Weibull distribution is far from perturbing its parameters, as proven by the expressions of such moments. The interpretation of the indices (see the graphs in Section 3.4) might be hard for the practitioner.

Therefore a new representation of the indices is proposed, in which the parameters of the input distributions are perturbed. For instance a parameter perturbation is presented in Figure 3.6, where 3 Weibull pdfs are plotted: the original pdf with parameters $(1.5, \pi)$ and two modified pdf where each parameter varies.

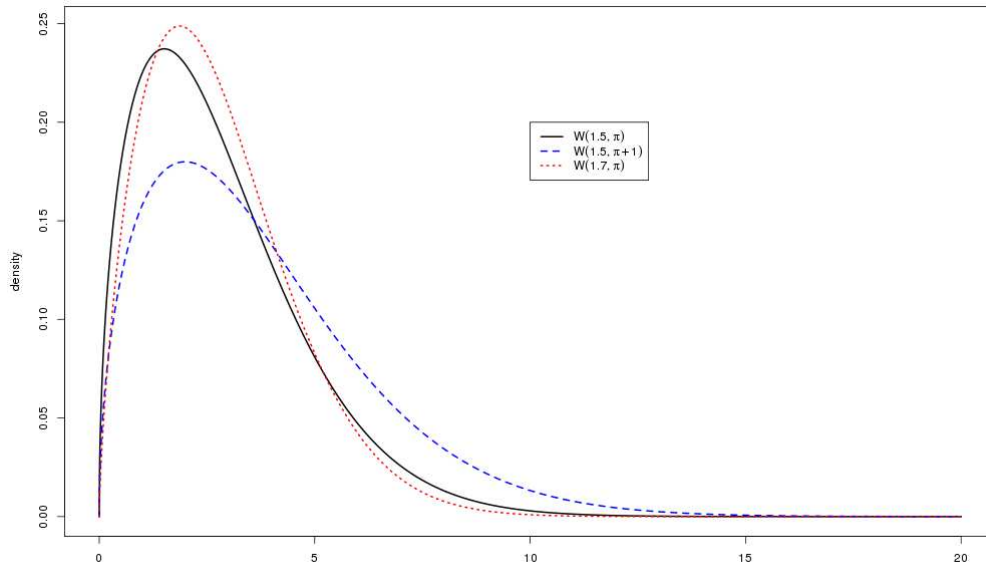


Figure 3.6: Original and perturbed Weibulls pdfs

This graph shows that each parameter variation produces different effects on several parts of the support. Precisely, increasing the scale parameter (dotted red curve) decreases the weight of the right-hand tail whereas increasing the shape parameter (dashed blue curve) increases the weight of the tail. The effect is reversed on the weight of the mode.

Given the input distributions, it can be inferred that increasing the mean μ will diminish the failure probability, increasing the variance σ^2 will increase the failure probability. It can also be stated that increasing the scale parameter b will concentrate the samples in the mode, thus increasing the failure probability whereas increasing the shape parameter c will increase the weight of the tail, thus diminish the failure probability. We are interested in the following: assuming that the true value of the parameters might not be the ones given, which of those 4 parameters causes the most uncertainty on the failure probability?

The use the DMBRSI is proposed, and it is suggested to plot them in function of the departure in density caused by the perturbation of the parameter.

Measure of the departure caused by parameters perturbation Distance quantifying the departure between two densities are numerous (Cha [24]), we propose the use the square of the Hellinger distance, which is defined as follows.

Definition 3.3.3 *The Hellinger Distance $H(P, Q)$ between two probability measures is the L_2 -distance between the square roots of the corresponding pdfs (Pollard [80]).*

$$H^2(P, Q) = \int \left(\sqrt{p(x)} - \sqrt{q(x)} \right)^2 dx = 2 - 2 \int \sqrt{p(x)q(x)} dx. \quad (3.29)$$

The Hellinger distance satisfies the inequality:

$$0 \leq H(P, Q) \leq \sqrt{2}. \quad (3.30)$$

The reasons for using the Hellinger distance over Kullback-Liebler divergence are:

- it is numerically practicable to estimate (the integral might be estimated by Simpson's rule);
- it is bounded;
- it is a distance thus symmetrical.

As the practitioner might not be familiar with the use of the Hellinger distance, tables eliciting the relationship between a parameter perturbation and the occasioned departure will be provided. For instance, when referring to Figure 3.6, the Hellinger distance between the original density and the one obtained when increasing the scale parameter (dotted red curve) is 0.0072. Conversely, the Hellinger distance between the original density and the one obtained when increasing the shape parameter (dashed blue curve) is 0.0422.

Dealing with the example When dealing with the example, the parameters are perturbed and the indices are plotted in function of the departure caused by the perturbation in Figure 3.7. We must stress that these are actually two graphs concatenated, in a sense that we plot the DMBSRI in function of the (square of the) Hellinger distance - yet for each parameters there are two perturbations that correspond to a given departure: the one corresponding to an increase, the other to a decrease. On Figure 3.7, the indices corresponding to an increase of the parameters appear on the right side of the graph, and the indices corresponding to a decrease of the parameters are plotted on the left

side. Confidence intervals are available thanks to asymptotic formulae provided in Section 3.2.3; yet they are not plotted here since it is an illustrative example.

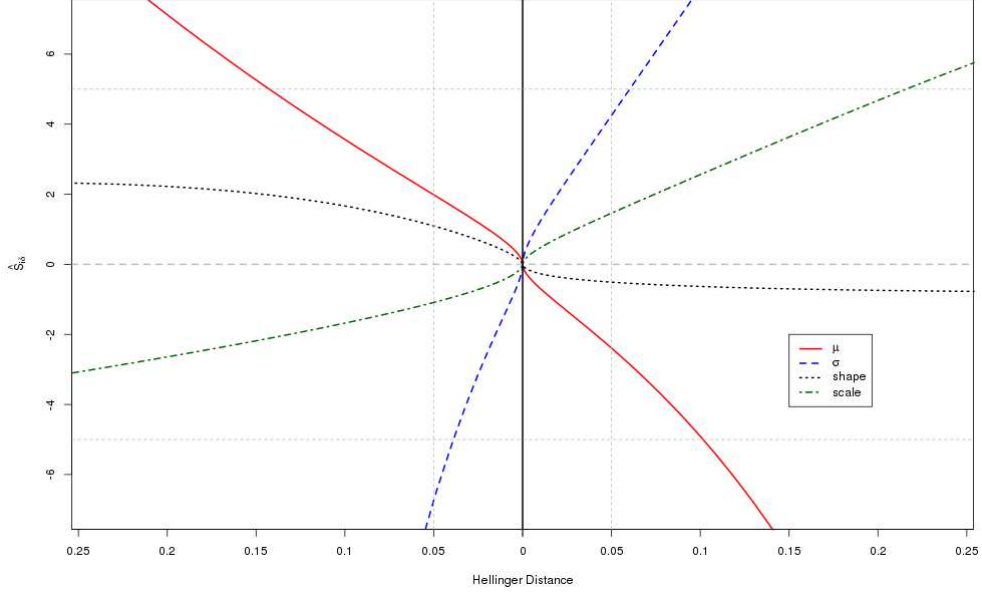


Figure 3.7: DMBRSI with parameters perturbations

Altogether with the Figure, Table 3.1 is provided: it expresses the departure in terms of parameters variation. The aim of such a table is to help the practitioner with quantifying the departure in terms of parameters perturbation. Note that Table 3.1 only focuses on parameters increasing (right-hand part of Figure 3.7). In the numerical examples of Section 3.4, both parameters increasing and decreasing will be dealt with.

	$X_1 \sim \mathcal{N}(\mu = 0, \sigma = 1)$		$X_2 \sim W(b = 1.5, c = \pi)$	
	$\mu \sigma = 1$	$\sigma \mu = 0$	$b c = \pi$	$c b = 1.5$
$H^2(X_i, X_{i\delta}) = 0$	0	1	1.5	π
$H^2(X_i, X_{i\delta}) = 0.05$	0.450	1.378	2.102	$\pi + 1.104$
$H^2(X_i, X_{i\delta}) = 0.1$	0.641	1.585	2.440	$\pi + 1.691$
$H^2(X_i, X_{i\delta}) = 0.15$	0.790	1.773	2.753	$\pi + 2.213$
$H^2(X_i, X_{i\delta}) = 0.2$	0.918	1.958	3.064	$\pi + 2.715$

Table 3.1: Hellinger distance in function of the parameter perturbation

The indices in Figure 3.7 show some central symmetry. This graph states that a variation in σ has the largest effect on the failure probability. Then comes μ , then the scale parameter b and finally the shape parameter c .

Conclusion, notion of equivalence This first example shows how the DMBRSI can be used to assess the influence of each input distributions' parameter on the failure probability.

We also argue that the perturbation is "equivalent" in the sense evoked in the last paragraph of Section 3.3.1.3. Indeed, when perturbing two parameters for instance expressed in different units or

different orders of magnitude, the Hellinger distance allows to quantify "equivalently" the amplitude of the departure produced by the parameter shift.

3.3.2.2 Methodology of input perturbation

In this subsection, we formalize what has been done in the previous first example.

Let us suppose that the i -th variable X_i of the input vector is distributed according to f_i . The i -th input has p_i parameters: it is parametrized by the vector $\Theta_i = (\theta_{i,1}, \dots, \theta_{i,p_i})$. The perturbation will be on the j -th parameter, and will be of the following form:

$$\theta_{i,j,\delta} = \theta_{i,j} + \delta_{i,j} \tag{3.31}$$

where $\delta_{i,j}$ is a given real such that $\Theta_{i\delta} = (\theta_{i,1}, \dots, \theta_{i,j} + \delta_{i,j}, \dots, \theta_{i,p_i})$ is still a parametrization vector for the input f_i (for instance a variance parameter cannot become negative). Vector $\Theta_{i\delta}$ parametrizes the modified pdf $f_{i\delta}$. It must be noticed as well that the support of the perturbed pdf $f_{i\delta}$ must lie within the support of f_i (for estimation purposes, see conditions of Lemma 3.2.1).

The framework given in Figure 3.1 is modified in Figure 3.8 to consider the parameters perturbations.

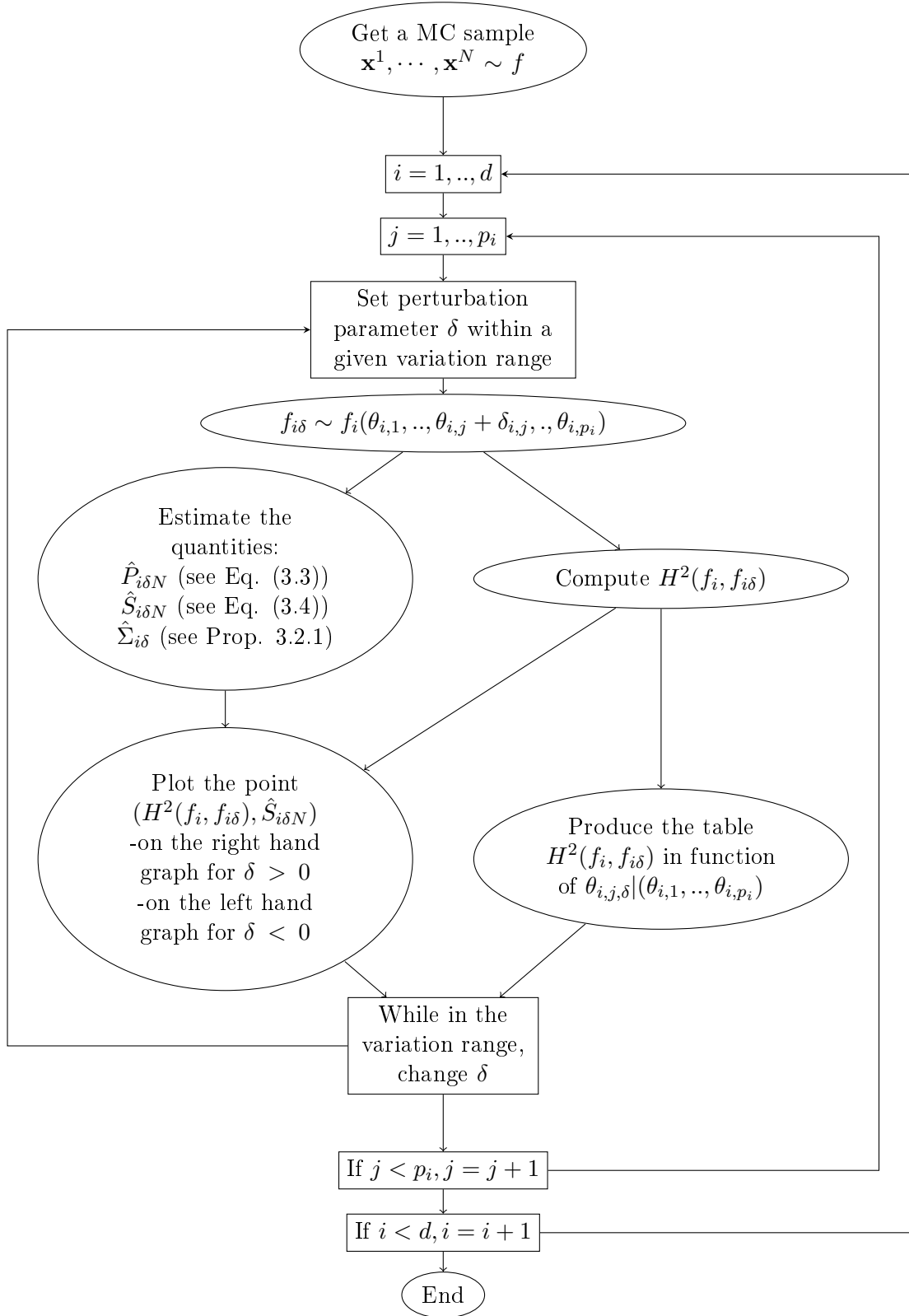


Figure 3.8: Specific DMBRSI framework for parameters perturbations

3.3.3 Choice of the perturbation given the objectives

3.3.3.1 Types of perturbations and variation ranges

The types of perturbations presented in this section are reminded and summarized here. Some recommendations are given on the range of the perturbations.

- Mean shifting (Eq. 3.15): if the inputs are identically distributed, then the perturbation is straightforward (standard mean shift for all the variables). The range of the perturbation must be chosen so that the confidence intervals of the indices are not too spread (and if possible separated). If the inputs are not identically distributed, the perturbation proposed in the last paragraph of Section 3.3.1.3 is the following: the original distribution is perturbed so that its mean is the original's one plus δ times its standard deviation and the perturbation is conducted on δ . For the moment, a range proposed for δ is from -1 to 1 .
- Variance shifting (Eq. 3.19): we argue that this perturbation is only to be used if the inputs are identically distributed. The new variances must be chosen so that the confidence intervals of the indices are not too spread.
- Quantile shifting (Eq. 3.25): the following strategy is proposed. First, fix a reference quantile (namely q_{ref}), then perturb this quantile for all the inputs. For the beginning of the study, we propose to perturb the 1st, 2nd and 3rd quartiles altogether with the 5th and 95th percentiles. Other quantiles might be perturbed in the following of the study if necessary.
- Parameters shifting (Eq. 3.31): this perturbation allows to deal with inputs that are not identically distributed. Here, the strategy is to perturb all the parameters of the input distributions. The range of the perturbation is driven by the square of the Hellinger distance between the original and the perturbed distribution. A perturbation so that this distance is $H^2 = .1$ seems enough to us (given our numerical tests).

3.3.3.2 Relationship between objectives and perturbations

In this paragraph are reminded the different objectives presented in Section 1.7. We propose the adapted perturbations for any given objective.

- REM1 (absolute ranking when the inputs are set): in this case we propose to perform the three KLM perturbations (mean shift, variance shift and quantile shift). For each perturbation, an input ranking can be produced.
- REM2 (quantify the sensitivity to the family or shape): in this case, we propose to perform only a quantile perturbation, as the quantiles allow to define a distribution.
- REM3 (assess the sensitivity to the parameters): in this specific case, we propose to use the parameters perturbation. This meets perfectly the objective.
- Objective 1 (variable ranking, assess which input "most needs better determination"). In this case, we propose the three KLM perturbations.
- Objective 2 (model simplification). This case is not treated in the manuscript but we can propose the following solution. A specific perturbation can be created, in which the perturbed input is a narrow distribution within the support of the original input (e.g. an input is set to a reference value and this reference value is moved along the support). The impact on the failure probability can be deduced from the indices thus meeting the objective.

- Objective 3 (model understanding). As the objective is to determine which particular values of some inputs leads to some behaviour of the output, we propose to perform the three KLM perturbations. Each perturbation provides supplementary knowledge on which part of the support of the input leads to the failure event.
- Objective 4 (calibration sensitivity). In this case we propose to perform the 4 perturbations type. The perturbations respectively allows to test the sensitivity to the moments, the tails and the parameters of the inputs.

Table 3.2 summarises the main ideas developed in this subsection.

	REM1	REM2	REM3	Obj. 1	Obj. 2	Obj. 3	Obj. 4
Mean shifting	×			×		×	×
Variance shifting	×			×		×	×
Quantile shifting	×	×		×		×	×
Parameters shifting			×				×
Specific					×		

Table 3.2: Type of perturbation recommended given the objective or the motivation

In addition with Table 3.2, we stress that the reference methods (FORM's Importance factors and Sobol' indices) only fulfill REM1 and Objective 1 (variable ranking).

3.4 Numerical experiments

3.4.1 Testing methodology

In this section, the proposed indices are tested on the numerical cases defined in Appendix B. A comparison with two references method (FORM's Importance factors and Sobol' indices) is provided. Importance factors and Sobol' indices are computed using the methodologies given in Lemaire [61] and Saltelli [87], respectively. The R packages *mistral* and *sensitivity* have been used. The Sobol' indices are computed using two initial samples of size 10^6 , resulting into $N = 10^6 \times (d + 2)$ function calls (Saltelli *et al.* [88]). The results of the Sobol' indices analysis were already provided in Section 1.4.

3.4.2 Hyperplane 6410 test case

This first test case was defined in Appendix B.1. Remind that all variables are independent standard Gaussian. Also recall that variable X_2 is most influential, then comes variable X_3 . X_1 has a small influence and X_4 has no influence at all. Finally remind that the failure probability is $P_f = 0.014$.

3.4.2.1 Importance factors

In this ideal hyperplane failure surface case, FORM provides an approximated value $\hat{P}_{FORM} = 0.01398$, which is as expected (Lemaire [61]) close to the exact value. 39 model calls have been required. The importance factors, given in Table 3.3, provide an accurate variable ranking for the failure function.

Variable	X_1	X_2	X_3	X_4
Importance factor	0.018	0.679	0.302	0

Table 3.3: Importance factors for hyperplane 6410 function

3.4.2.2 Sobol' indices

We reproduce here table 1.4 and the resulting conclusions.

Index	S_1	S_2	S_3	S_4	S_{T1}	S_{T2}	S_{T3}	S_{T4}
Estimation	0.002	0.254	0.054	0	0.200	0.940	0.720	0

Table 3.4: Estimated Sobol' indices for the hyperplane 6410 case

The total indices assess that X_2 is extremely influential, and that X_3 is highly influential. X_1 has a moderate influence and X_4 has a null influence. This last point is interesting: it shows that this SA method can detect the non-influential variables.

3.4.2.3 DMBRSI

The method presented throughout this chapter is applied on the first hyperplane function. As explained in section 3.3, several ways to perturb the input distributions exist. A mean shifting, a variance shifting, a quantile shifting and a parameters perturbation will be performed. We follow the methodology displayed in Figures 3.1 and 3.8. We stress that all the indices are estimated with the same MC sample. The MC estimation gives $\hat{P} = 0.01446$ with 10^5 function calls.

Mean shifting For the mean shifting (see Eq. (3.15)), the domain variation for δ ranges from -1 to 1 with 40 points, reminding that $\delta = 0$ cannot be considered as a perturbation since it is the expectation of the original density. The results of the estimation of the indices $\widehat{S}_{i\delta}$ are plotted in Figure 3.9, altogether with 95% symmetrical confidence intervals (CI).

The indices $\widehat{S}_{i\delta}$ behave in a monotonic way given the importance of the perturbation. The slope at the origin is directly related to the value of a_i . For influential variables (X_2 and X_3), the increasing or the decreasing is faster than linear, whereas the curve seems linear for the slightly influential variable (X_1). Modifying the mean with a positive amplitude slightly rises the failure probability for X_1 , highly decreases it for X_2 and increases it for X_3 . The effects are reversed with similar amplitude for negative δ . It can be seen that X_4 has no impact on the failure probability for any perturbation. Those results are consistent with the expression of the failure function. One can see that the CI associated to all variables are fairly well separated, except for the small absolute value of δ .

Variance shifting For the variance shifting (see Eq. (3.19)), the variation domain for V_{per} ranges from $1/20$ to 3 with 28 points, where $V_{\text{per}} = 1$ is not a perturbation. The estimated indices are plotted in Figure 3.10. The 95% symmetrical CI are plotted around the indices, using the presented asymptotic formulas in Section 3.2.

Increasing the variance of inputs X_2 and X_3 increases the failure probability, whereas it decreases when decreasing the variance. Modifying the variance of X_1 and X_4 have no effect on the failure probability. The increasing of the indices is linear for X_2 and X_3 , and the decreasing of the indices is faster than linear, especially for X_2 . Considering the CI, one can see that they are well separated

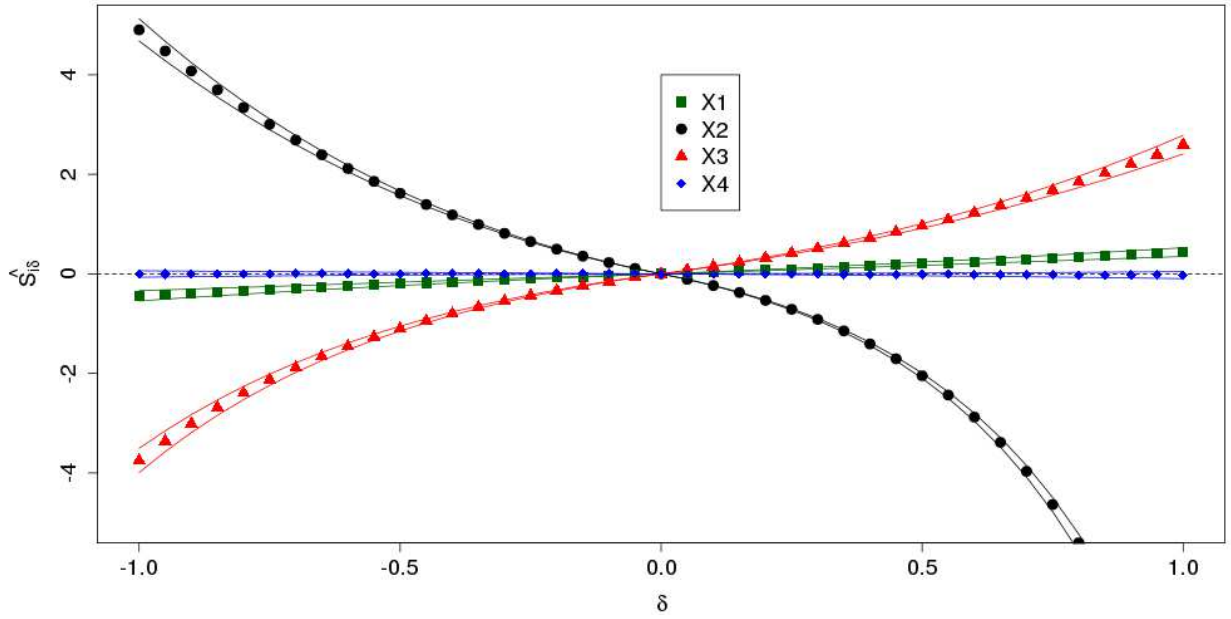


Figure 3.9: Estimated indices $\widehat{S}_{i\delta}$ for the 6410 hyperplane function with a mean shifting

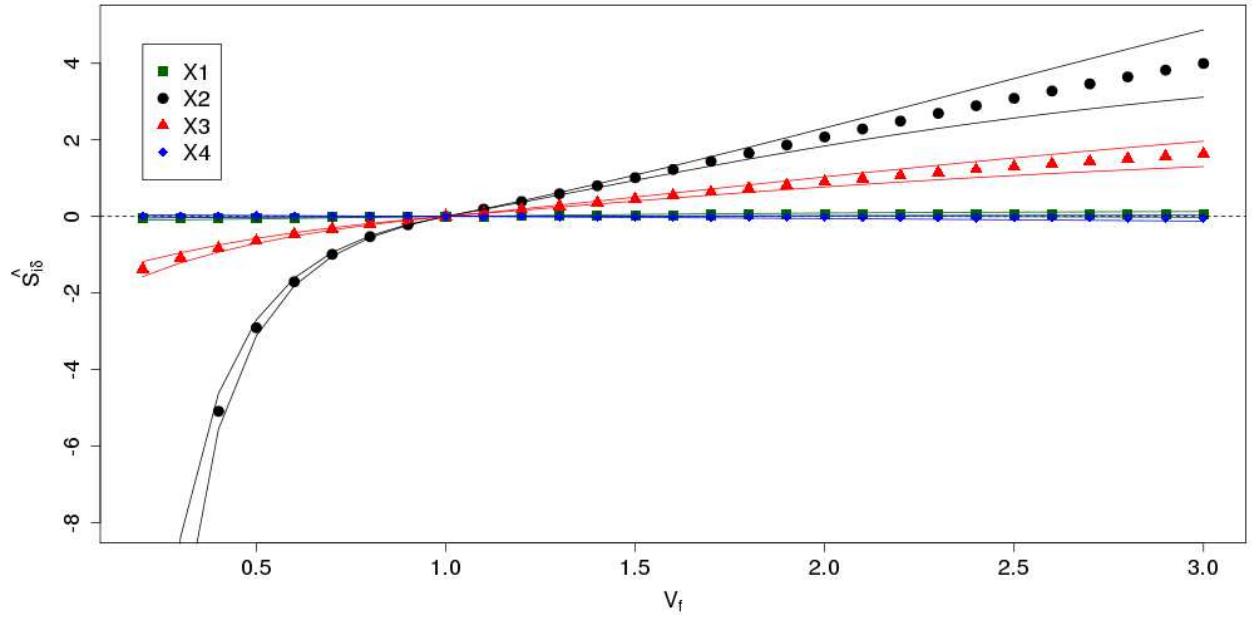


Figure 3.10: Estimated indices $\widehat{S}_{i,V_f}$ for hyperplane function with a variance shifting

for variables X_2 and X_3 , assessing the relative importance of these variables. On the other hand, the CI associated to X_1 and X_4 are not separated and contain 0. Influence of X_1 and X_4 cannot thus be separated - but is estimated as null for both variables.

Quantile shifting

We first perturb the 5th percentile. The tail is perturbed in such ways that it weights between 1%

and 10%. The results are displayed in Figure 3.11.

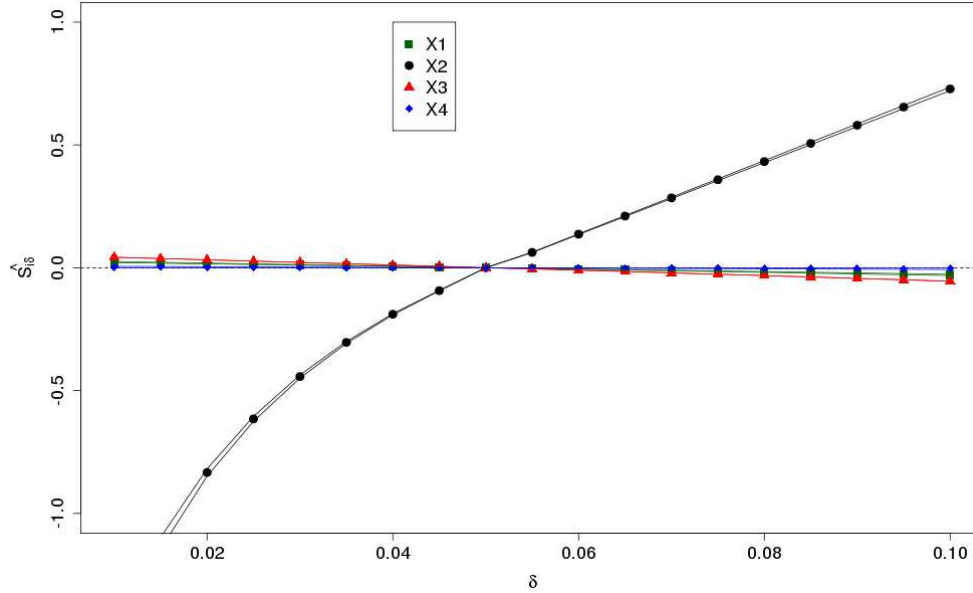


Figure 3.11: 5th percentile perturbation on the hyperplane 6410 test case

Concerning the left-hand tail, this figure shows the dominant role of variable X_2 . Effects of variables X_1 and X_3 are small whereas the indices associated to X_4 are null, assessing the non-influence of the last variable - at least when perturbing the left-hand tail.

The first quartile or 25th percentile is then perturbed. The weight of the tail under the 25th percentile (meaning the left-hand tail) of the input varies between 10% and 40%. The result of the numerical experiments are displayed in Figure 3.12.

This plot shows that an increase of the 1st quartile leads to an increase of the failure probability for variable X_2 whereas it leads to a decrease for variables X_3 and X_1 in order of influence. A quantile perturbation on variable X_4 has no effect on the failure probability. On the other hand, when decreasing the weight of the 1st quartile, the failure probability increases for variable X_3 and X_1 , and decreases for variable X_2 .

We then perturb the second quartile or median. The density is perturbed so that the left-hand tail weight varies between 25% and 75%. The results are displayed in Figure 3.13.

This last graph shows the relative importance of X_3 and X_2 . X_1 behaves as X_3 , only with a smaller effect. This is relevant given the expression of the model.

Let us now perturb the third quartile or 75th percentile. The weight of the pdf under the 75th percentile of the standard Gaussian varies between 60% and 90% - which is the same as perturbing the weight of the right-hand tail between 10% and 40%. The result of the numerical experiments are displayed in Figure 3.14.

This shows that the most influential variable when perturbing the 3rd quartile is variable X_3 , then comes variable X_2 , then variable X_1 . Perturbing variable X_4 has no effect on the failure probability, as expected. We proceed as before and perturb a more extreme quantile, namely the 95th percentile. It varies between 90% and 99%. The results are displayed in Figure 3.15.

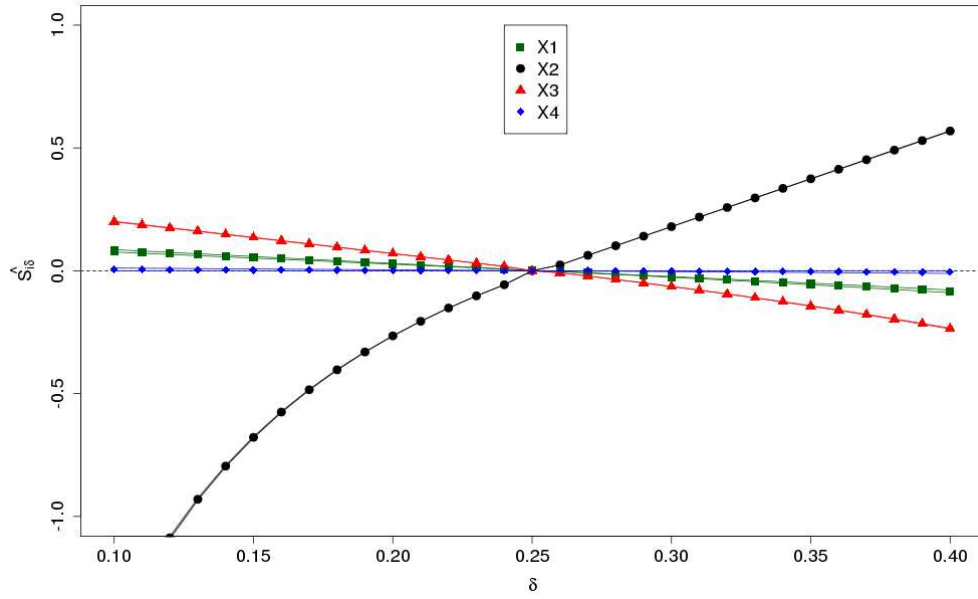


Figure 3.12: 1st quartile perturbation on the hyperplane 6410 test case

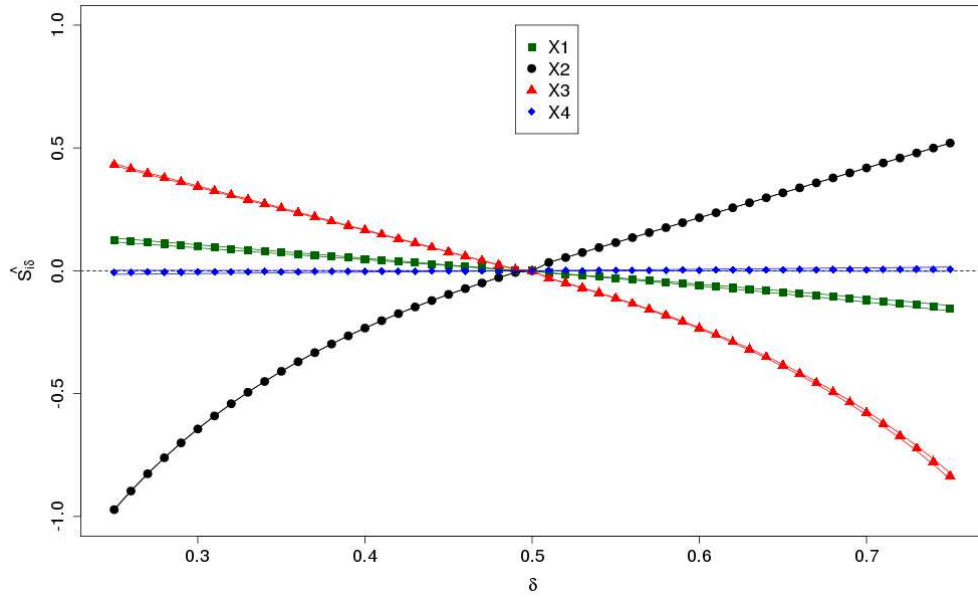
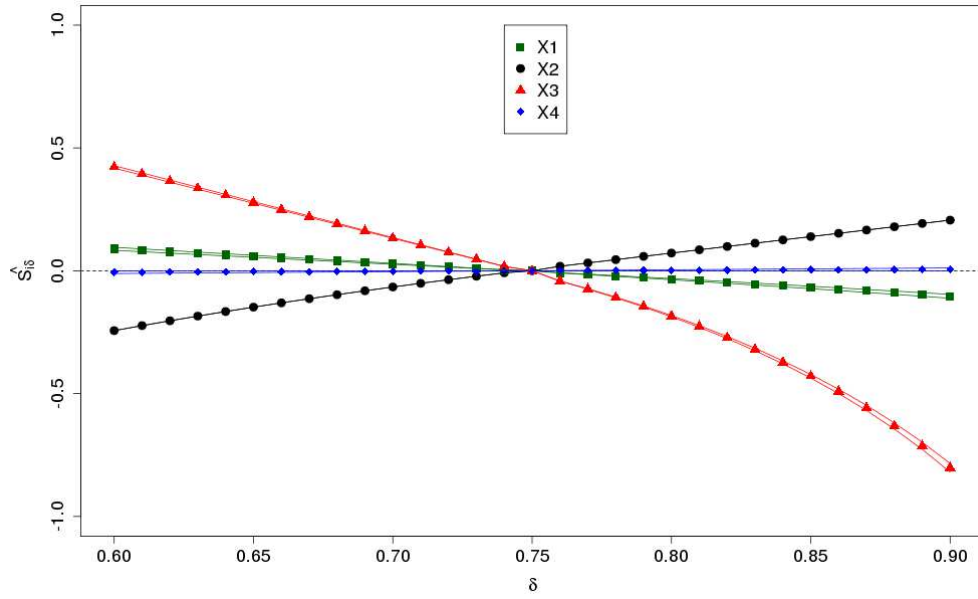
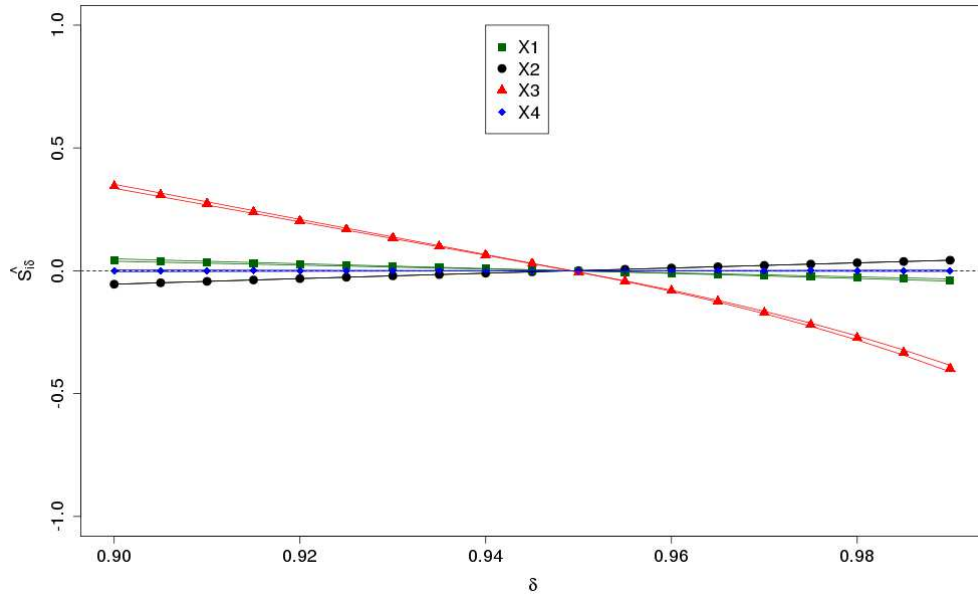


Figure 3.13: Median perturbation on the hyperplane 6410 test case

This shows the main influence of variable X_3 when dealing with perturbations of the right-hand tail.

As a conclusion on this monotonic test case, it can be said that the input values leading to the failure event are mostly the extremes values of the left-hand tail for variable X_2 and the extremes values of the right-hand tail for variable X_3 .

Parameters perturbation The methodology presented in subsection 3.3.2 is tested here. There are 8 parameters governing this model: the means and standard deviations of each of 4 variables.

Figure 3.14: 3rd quartile perturbation on the hyperplane 6410 test caseFigure 3.15: 95th percentile perturbation on the hyperplane 6410 test case

Based on the same 10^5 MC sample, Figure 3.16 can be plotted.

This figure has to be interpreted altogether with table 3.5. Recall that all the inputs follow standard Gaussian.

Interpreting both Figure 3.16 and table 3.5 lead us to conclude the following. The most influential parameter with respect to the failure probability is the standard deviation of X_2 . Increasing this quantity so that the H^2 distance between the original and the perturbed density is 0.05 triples the failure probability. On the other side of the graph, diminishing the variance of X_2 strongly diminishes the failure probability with respect to the other parameters. Then, the other influential

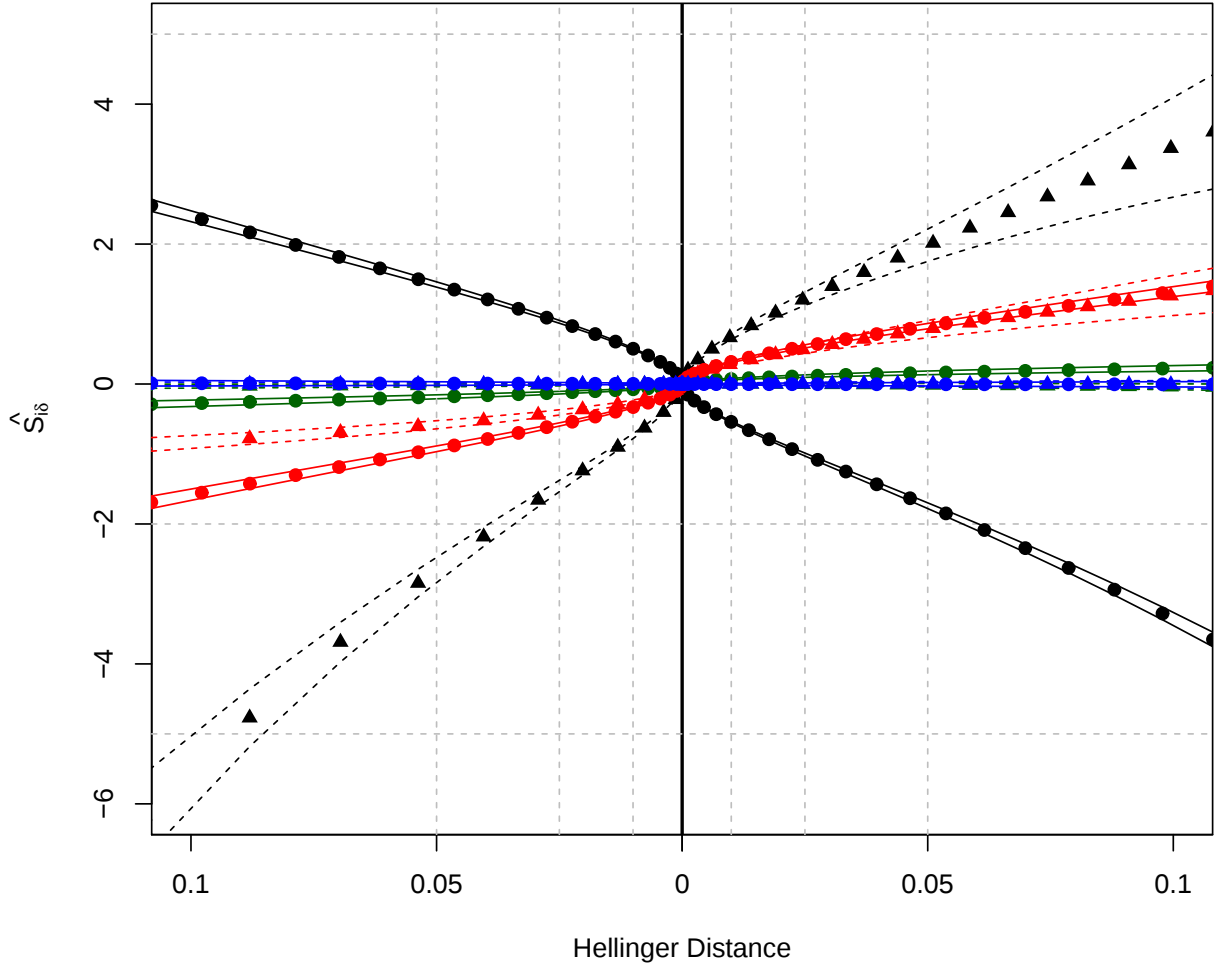


Figure 3.16: Parameters perturbation on the hyperplane 6410 test case. Dots are for means, triangle for the standard deviations. Green corresponds to X_1 , black to X_2 , red to X_3 and blue to X_4 .

	$X_i \sim \mathcal{N}(\mu = 0, \sigma = 1)$	
	$\mu \sigma = 1$	$\sigma \mu = 0$
$H^2(X_i, X_{i\delta}) = 0$	0	1
$H^2(X_i, X_{i\delta}) = 0.01$	0.200/-0.200	1.152/0.868
$H^2(X_i, X_{i\delta}) = 0.025$	0.317/-0.317	1.252/0.798
$H^2(X_i, X_{i\delta}) = 0.05$	0.450/-0.450	1.378/0.725
$H^2(X_i, X_{i\delta}) = 0.1$	0.641/-0.641	1.585/0.631

Table 3.5: Hellinger distance in function of the parameter perturbation. The first value is an increase of the parameter (right hand of the graph) whereas the second is a decrease of the parameter (left hand of the graph). Both perturbation lead to the same H^2 departure.

parameter is the mean of X_2 . It is slightly less important than the standard deviation of X_2 yet it is much more influential than others parameters. When increasing the standard deviation and (not at the same time) the mean of X_3 , it affects positively the failure probability. The estimated indices are confounded, but the CI are slightly larger for the standard deviations. When decreasing these

last two parameters, the failure probability decreases. Yet in this case, the mean is more influential than the standard deviation. This is an interesting result. When dealing with the parameters of X_1 , it must be noticed that the estimated indices for the standard deviations lie around 0 and are confounded with the one for X_4 . However the indices for the mean are slightly positive and increasing when increasing this mean while they are slightly negative and decreasing when diminishing this parameter. The indices associated to X_4 , both mean and standard deviation are null, thus assessing the non-influence of this last variable.

Conclusion and discussion The DMBRSI has brought the following conclusions:

- When shifting the mean (that is to say the central tendency in this case), the most influential variable is X_2 , followed by X_3 . X_1 is slightly influential while X_4 is not influential at all.
- When shifting the variance, variable X_2 is more influential than variable X_3 . Variables X_1 and X_4 have no impact when shifting the variance that is to say when we are interesting in the tails behaviour.
- The many graphs associated with several quantiles shifts lead to the conclusion that the influential regions leading to the failure event are the extreme left-hand tail values for variable X_2 and the extreme right-hand tail values for variable X_3 .
- When shifting the parameters, it lead to the conclusion that the most influential parameters are the standard deviation of X_2 , the mean of X_2 , then the mean of X_3 followed by the standard deviation of X_3 . Others parameters have a small to null influence.

These results are consistent with each other. We argue that all these information are much richer than the ones provided by importance factors and by Sobol' indices. Indeed, the information is provided about regions of the input space leading to failure event; or on parameters whose variation will provide a broad change on the failure probability. This is, in our opinion, more of interest to the practitioner than a "simple" variable ranking.

3.4.3 Hyperplane 11111 test case

This second test case was defined in Appendix B.1. Remind that all variables are independent standard Gaussian. Also recall that all variables have the same influence. Finally remind that the failure probability is $P_f = 0.0036$.

3.4.3.1 Importance factors

In this ideal hyperplane failure surface case, FORM provides an approximated value $\hat{P}_{FORM} = 0.0036$, which is as expected (Lemaire [61]) close to the exact value. 33 model calls have been required. The importance factors, given in Table 3.6, provide an exact variable ranking for the failure function. They assess that all variables have the same importance. That was the sought after result.

Variable	X_1	X_2	X_3	X_4	X_5
Importance factor	0.2	0.2	0.2	0.2	0.2

Table 3.6: Importance factors for hyperplane 11111 function

3.4.3.2 Sobol' indices

We reproduce here Table 1.5 and the resulting conclusions.

On Table 3.7 the estimated Sobol indices with 2 samples of size 10^6 , using the Saltelli 02 method. The total number of function evaluations is 7×10^6 .

Index	S_1	S_2	S_3	S_4	S_5	S_{T1}	S_{T2}	S_{T3}	S_{T4}	S_{T5}
Estimation	0.015	0.013	0.014	0.009	0.015	0.677	0.673	0.695	0.674	0.685

Table 3.7: Estimated Sobol' indices for the hyperplane 11111 case

The weak first order indices (less than 2% of the variance explained) and the high total indices assess that all the variables are influential in interaction with the others. All the total indices are approximatively the same showing that this SA method can give the same importance to each equally contributing input.

3.4.3.3 DMBRSI

As in the previous example, all the types of perturbations proposed in section 3.3 will be tested on this second numerical case. The methodology displayed in Figures 3.1 and 3.8 is used. We again stress that all the indices are estimated with the same MC sample. The MC estimation gives $\hat{P} = 0.00353$ with 10^5 function calls, which is a good order of magnitude.

Mean shifting The mean of all the variables is shifted (one variable at a time), see Eq. (3.15). The domain variation for δ ranges from -1 to 1 with 40 points, reminding that $\delta = 0$ cannot be considered as a perturbation since it is the expectation of the original density. The result is plotted in Figure 3.17, with a different color and different sign for each variable. 95% confidence intervals are plotted.

For small values (of absolute value smaller than 0.5) of new mean, the estimated indices are similar for all the variables. When the values of the new mean get higher (in absolute value), some numerical noise spreads the indices. However, the confidence intervals are not disconnected. We conclude from this graph that, when dealing with the central tendency, all the variables involved in the code have the same influence on the failure probability.

Variance shifting The variance of all the variables is now shifted (still one variable at a time), see Eq. (3.19). The domain variation for V_f (the perturbed variance) ranges from 0.2 to 3 with 71 points, reminding that $V_f = 1$ is not a perturbation. The result is plotted in Figure 3.18, with a different color and different sign for each variable. 95% confidence intervals are plotted.

For small values of perturbation (variance ranging from 0.5 to 1.5), the indices are confounded. When increasing the strength of the perturbation, one can see that the indices get disjointed. However the confidence intervals are not disconnected, thus one can infer that the values of the indices are roughly the same (they are theoretically the same in this model). An interesting fact is that all confidence intervals do not have the same width. A conclusion from this graph is that, when dealing with the tails, all the variables involved in the code have the same influence on the failure probability.

Quantile shifting As previously, we perturbed the 1st, 2nd and 3rd quartiles altogether with the 5th and 95th percentiles. As all the graphs have a similar shape, only one (for the median) is displayed in Figure 3.19.

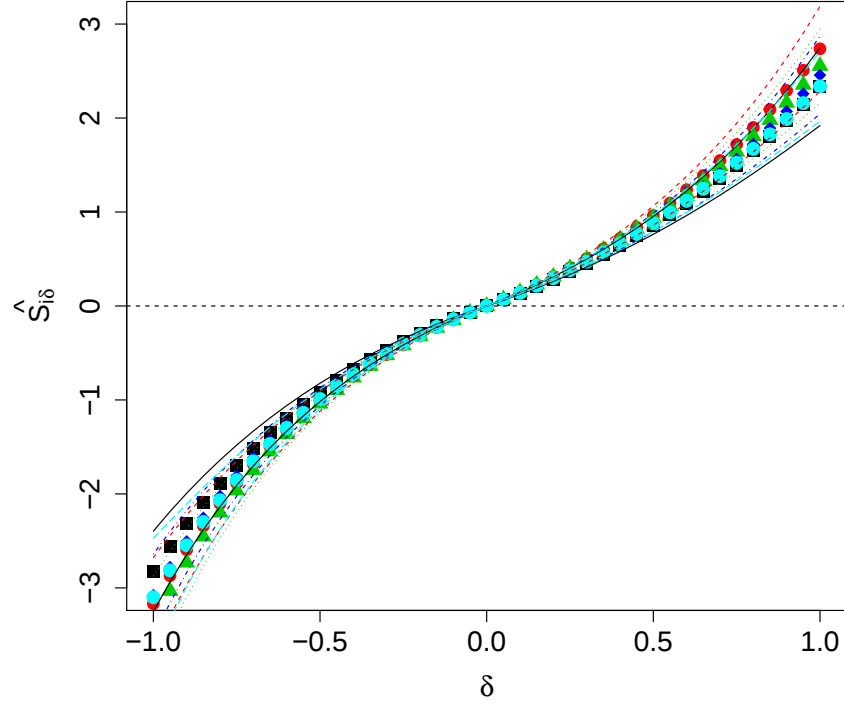


Figure 3.17: Estimated indices $\hat{S}_{i\delta}$ for the 11111 hyperplane function with a mean shifting

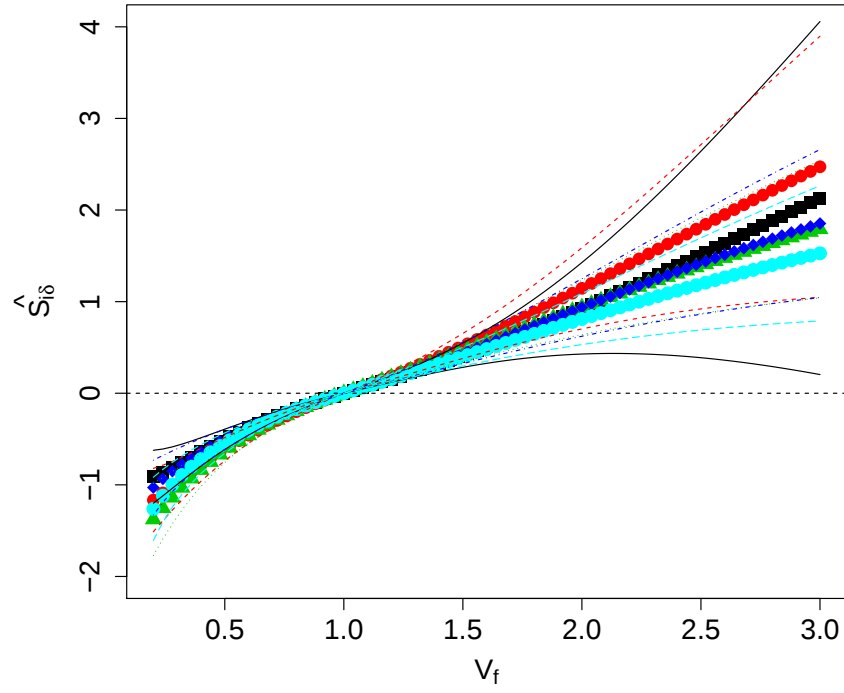


Figure 3.18: Estimated indices $\hat{S}_{i\delta}$ for the 11111 hyperplane function with a variance shifting

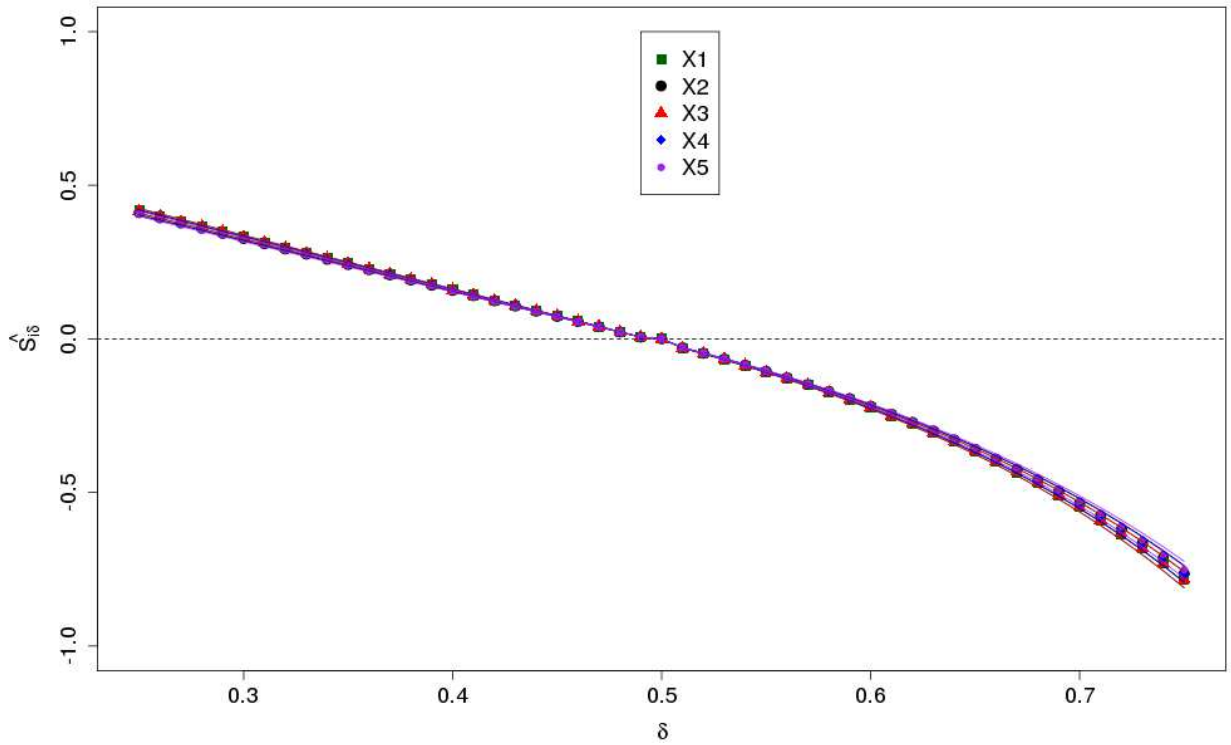


Figure 3.19: Median perturbation on the hyperplane 11111 test case

This graph shows that all the variables have an equivalent behaviour when their quantiles are perturbed.

Parameters shifting 10 parameters drive the model: a variance and a standard deviation for each Gaussian input. Each of these parameters is perturbed and the estimated indices are plotted in function of the Hellinger distance in Figure 3.20, as explained in Figure 3.8. 95% confidence intervals are provided as well.

This figure leads to several comments and needs to be interpreted with table 3.5. Increasing any parameter leads to an increase of the failure probability whereas diminishing any parameter leads to a reduction of the failure probability. When increasing the parameters, indices are badly separated. A closer look shows that the indices associated to the means (dots) are packed down to (slightly) lower values than the indices associated to the standard deviations (triangles), which are more dispersed. The confidence intervals (solid lines for the means, dashed lines for the standard deviations) are smaller for the means than for the standard deviations. On the other side of the graph, when reducing the parameters, an "equivalent" (in the H^2 sense) reduction of the mean has more impact (on the reduction of the failure probability) than a reduction of the standard deviations. The confidence intervals are well separated. In all cases, there is no way to distinguish the effects of several variables, which was expected in this model.

Conclusion and discussion When shifting the mean, for small perturbations, all the variables are ranked with the same importance. This goes the same for a variance shift and a quantile shift. Similarly, a parameter perturbation does not allow to say that a variable is more influential than

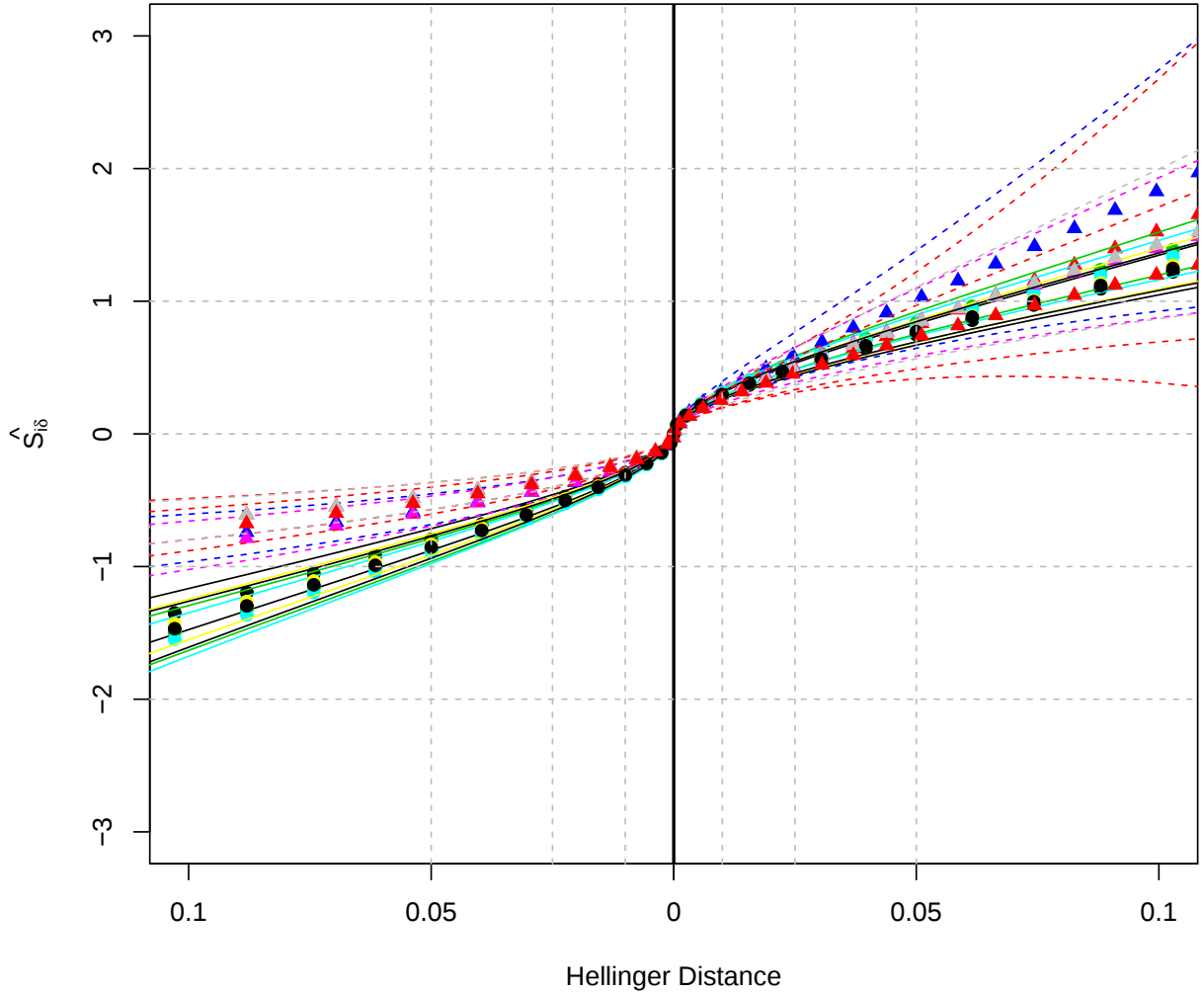


Figure 3.20: Parameters perturbation on the hyperplane 11111 test case. Dots are for means, triangle for the standard deviations. A different color is used for each variable.

another (however the parameters of a given variable does not have the same influence on the failure probability).

If the objective was a pure variable ranking, then small variations of moments and quantile are adapted - at least on this case it has shown the ability to affect roughly the same indices to equivalently influential variables.

If the objective of the SA is to know which parameters impact the most the failure probability (and a realistic objective would be "where to reduce the uncertainty in order to reduce the failure probability"), we stress here that the parameters shift has allowed to conclude that for this case the means of the variables have more influence than their standard deviations.

3.4.4 Hyperplane with 15 variables test case

This third test case was defined in Appendix B.1. Remind that all variables are independent standard Gaussian. Also recall that the aim of this example is to test the ability of the proposed method

Variable	X_1 to X_5	X_6 to X_{10}	X_{11} to X_{15}
Importance factor	0.192	7.69×10^{-3}	0

Table 3.8: Importance factors for the hyperplane 15 variables

to discriminate the variables in three classes: influential, weakly-influential, non-influential. Finally remind that the failure probability is $P_f = 0.00425$.

3.4.4.1 Importance factors

In this ideal hyperplane failure surface case, FORM provides an exact value $\hat{P}_{FORM} = 0.00425$, as expected. 31 model calls have been required. The importance factors, given in Table 3.8, provide an exact variable ranking for the failure function. They give to each group of variable different values of influence. The ranking is correct, namely the influential variables are detected as such, the weakly-influential variables have a very small importance factor and the non-influential variables have importance factors of 0. That was the sought after result.

3.4.4.2 Sobol' indices

We reproduce here Table 1.6 and the resulting conclusions.

On Table 3.9 are presented the estimated Sobol' indices with 2 samples of size 10^6 , using the Saltelli [87] method. The total number of function evaluations is 17×10^6 .

Index	S_1 to S_5	S_6 to S_{10}	S_{11} to S_{15}
Estimation	0.014 to 0.018	0.001 to 0.002	0
Index	S_{T1} to S_{T5}	S_{T6} to S_{T10}	S_{T11} to S_{T15}
Estimation	0.655 to 0.673	0.141 to 0.150	0

Table 3.9: Estimated Sobol' indices for the hyperplane with 15 variables case

The first order indices are all weak, yet separated in three groups. The total indices give a good separation between the influential, weakly influential and non influential variables. The Sobol' indices SA method is able to deal with problems of medium dimension; however it has an heavy computational cost in this case.

3.4.4.3 DMBRSI

As in the previous example, all the types of perturbations proposed in section 3.3 will be tested on this third numerical case. The methodology displayed in Figures 3.1 and 3.8 is used. We stress again that all the indices are estimated with the same MC sample. The MC estimation gives $\hat{P} = 0.0042$ with 10^5 function calls, which is close from the real result.

Mean shifting The mean of all the variables is shifted (one variable at a time), see Equation (3.15). The domain variation for δ ranges from -1 to 1 with 40 points, reminding that $\delta = 0$ cannot be considered as a perturbation since it is the expectation of the original density. The result is plotted in Figure 3.21, with a different color for each variable and different sign for each group variable. 95% confidence intervals are plotted.

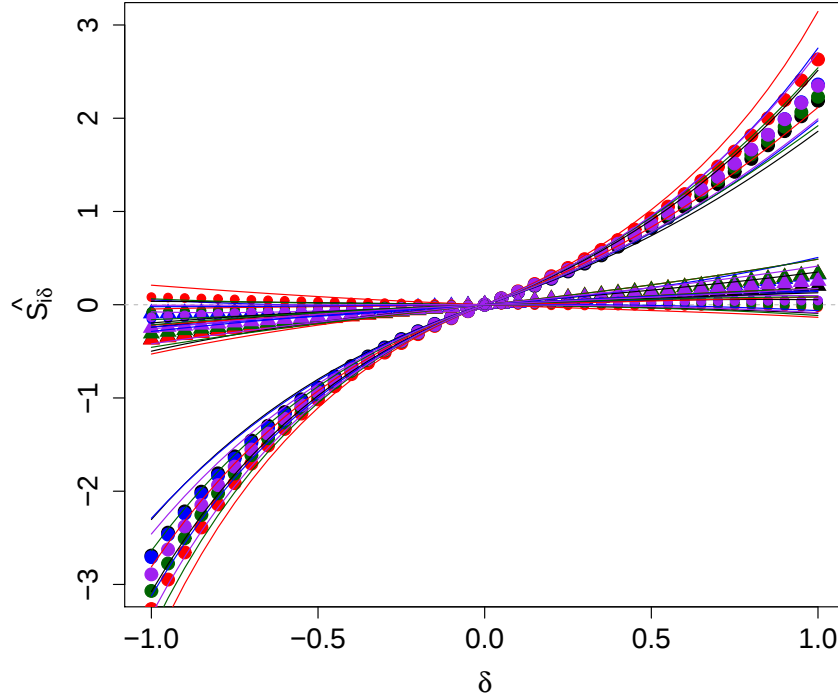


Figure 3.21: Estimated indices $\hat{S}_{i\delta}$ for the 15 variables hyperplane function with a mean shifting

For the influential variables (big dots), increasing the mean increases the failure probability whereas decreasing this parameter decreases the failure probability. However distinguish the effects of the weakly-influential variables (triangles) from the effects of the non-influential variables (small dots) is not possible due to the covering of the confidence intervals. So far, DMBRSI does not allow to separate the effects of the two last groups of variables. However, another test with a MC size of 10^6 draws (graphs non provided here) allows a good separation of the weakly and non-influential variables.

Variance shifting The variance of all the variables is now shifted (still one variable at a time), see Equation (3.19). The domain variation for V_f (the perturbed variance) ranges from 0.2 to 3 with 71 points, reminding that $V_f = 1$ is not a perturbation. The result is plotted in Figure 3.22, with a different color and different sign for each variable. 95% symmetrical confidence intervals are plotted.

The influential variables (big dots) are well separated from the others. As expected for these variables, increasing (respectively decreasing) the variance increases (respectively decreases) the failure probability. However, the effects for the weakly-influential (triangles) and non-influential (small dots) variables, the effects are hardly separable (see the confidence intervals). As well as previously, DMBRSI does not allow to separate the effects of the two last groups of variables (weakly and non-influential).

However, increasing the sample size of a factor 10 (graph not provided here) still does not allow to separate the effects of the last two groups of variable. This might be due to the relative null-influence of a variance shift in the last 10 variables.

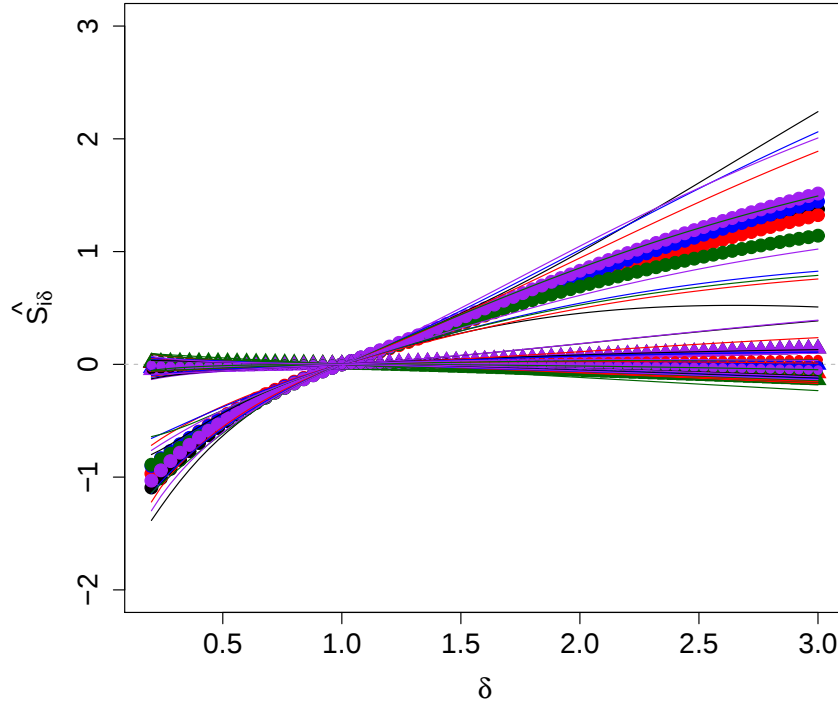


Figure 3.22: Estimated indices $\hat{S}_{i\delta}$ for the 15 variables hyperplane function with a variance shifting

Quantile shifting As in the previous numerical experiments, the 1st, 2nd and the 3rd quartiles altogether with the 5th and 95th percentiles were perturbed. All the graphs are similar, only the left scale (the value of the sensitivity indices) varies, thus only one (relative to the median perturbation) is displayed in Figure 3.23.

This graph somehow allows the ranking in influential, weakly-influential and non-influential variables. This graph shows that the method allows a separation of the 15 variables into 3 groups of influence: medium, small and null influence although the separation between the two last groups is not straightforward.

The 10 first variables (2 first groups of 5 variables) have an equivalent behaviour when their quantiles are perturbed: increasing the weight of the left-hand tail increases the failure probability whereas it decreases this probability when increasing the weight of the right-hand tail. The indices associated to the last 5 variables have confidence interval values that include 0.

Increasing the sample size by a factor 10 allows to obtain a graph that accurately separates the diverse groups of variables (the graph is not provided here as it is the same as Figure 3.23).

With this type of perturbation, the DMBRSI allows to separate the variables by group of influence.

Parameters perturbation The model is driven by 30 parameters: a variance and a standard deviation for each Gaussian input. Each of these parameters is perturbed and the estimated indices are plotted in function of the Hellinger distance in Figure 3.24, as explained in Figure 3.8. 95% confidence intervals are provided as well. As the graph gets too complicated for an adequate representation, only one variable per group is plotted.

Table 3.5 is needed as well to interpret this graph. From the graph with all the indices plotted

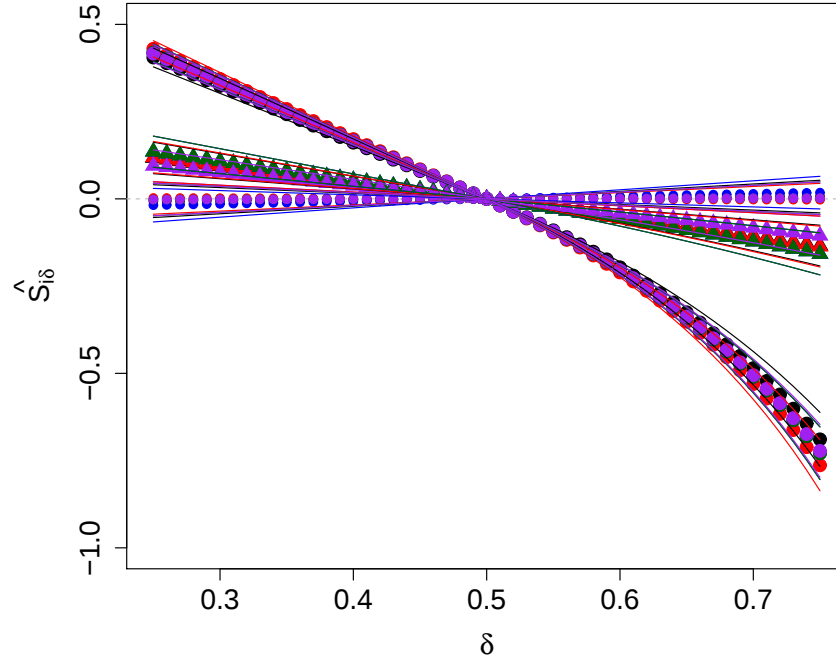


Figure 3.23: Median perturbation on the hyperplane with 15 variables test case

(not showed here) and from Figure 3.24, one can infer the following. The parameters related to the first variable - related to the first influence group - (black, dots for the mean, triangle for the standard deviation) are the most influential, with a bigger influence from the mean when decreasing the parameter. When increasing the parameters, the effects of the standard deviation and of the mean are not discernible. The confidence interval for the standard deviations (dashed lines) is quite wider than the one associated with the mean. However the indices associated with the means and variance of the other groups of variables are too noisy and cannot be interpreted.

Conclusion and discussion DMBRSI is not adapted to this medium dimension case. Indeed, only the quantile perturbation is able to distinguish the weakly from the non-influential variables. The parameter perturbation method especially leads to representation problem, with 30 curves to plot plus the confidence intervals. This leads to the conclusion that DMBRSI should not be used as a screening method.

3.4.5 Hyperplane with same importance and different spreads test case

This fourth test case was defined in Appendix B.1. Remind that all variables are independent Gaussian with mean 0 and increasing standard deviation. Also recall that the aim of this example is to give to equivalently influential variables that are not distributed similarly the same importance. Finally remind that the failure probability is $P_f = 0.0036$.

3.4.5.1 Importance factors

In this ideal hyperplane failure surface case, FORM provides an approximated value $\hat{P}_{FORM} = 0.0036$, which is as expected (Lemaire [61]) close to the exact value. 33 model calls have been required. The importance factors, given in Table 3.10, provide an exact variable ranking for the

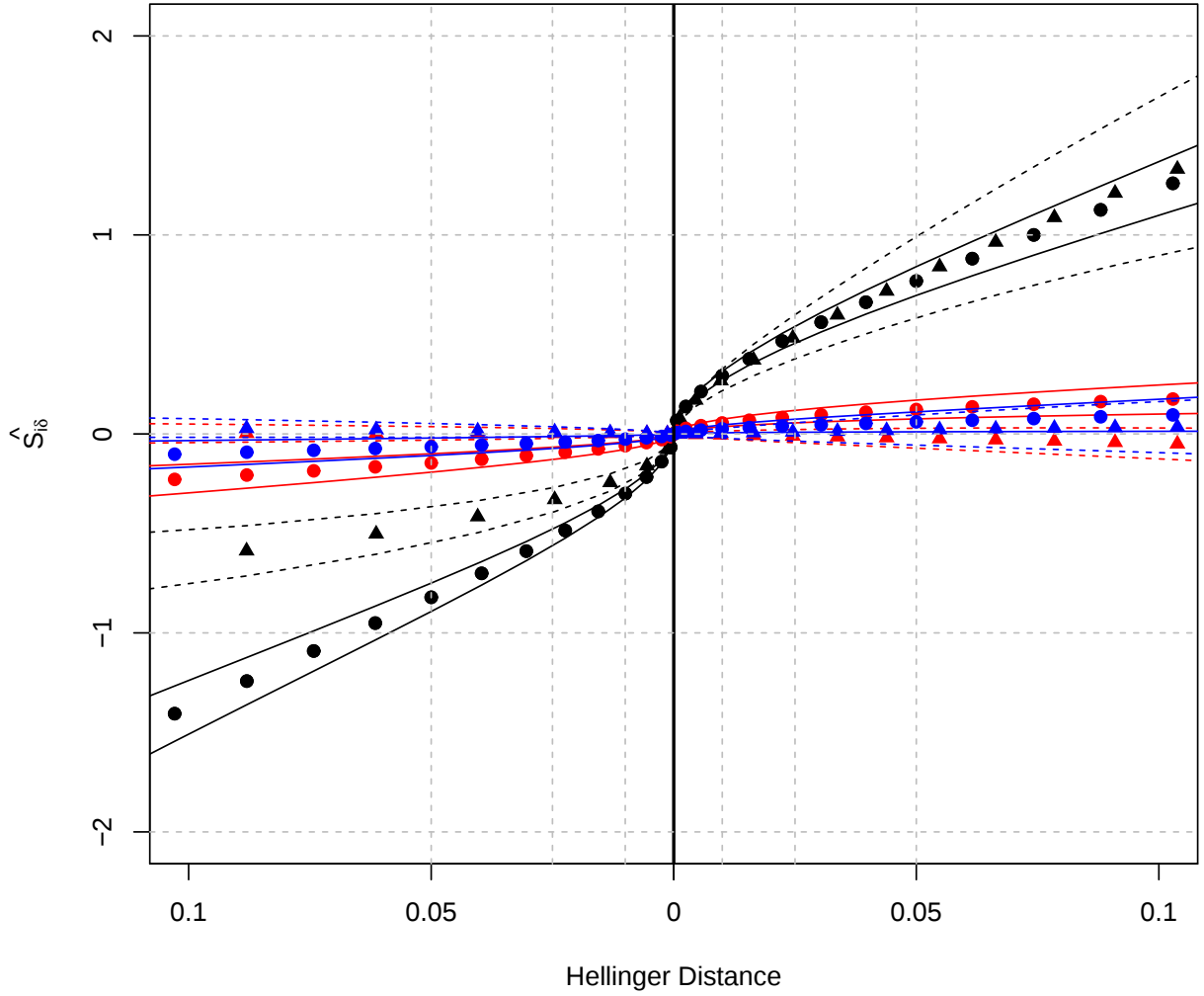


Figure 3.24: Parameters perturbation on the 15 variables hyperplane test case. Dots are for means, triangle for the standard deviations. Black is for the first group of influence, red is for the second and blue for the third.

failure function. They assess that all variables have the same importance. That was the expected result.

Variable	X_1	X_2	X_3	X_4	X_5
Importance factor	0.2	0.2	0.2	0.2	0.2

Table 3.10: Importance factors for hyperplane with different spreads function

3.4.5.2 Sobol' indices

We reproduce here Table 1.7 and the resulting conclusions.

On Table 3.11 are presented the estimated Sobol' Indices. The computation was done with 2 samples of size 10^6 , using the Saltelli [87] method. The total number of function evaluations is 7×10^6 .

Index	S_1	S_2	S_3	S_4	S_5	S_{T1}	S_{T2}	S_{T3}	S_{T4}	S_{T5}
Estimation	0.027	0.028	0.025	0.025	0.028	0.611	0.622	0.618	0.618	0.624

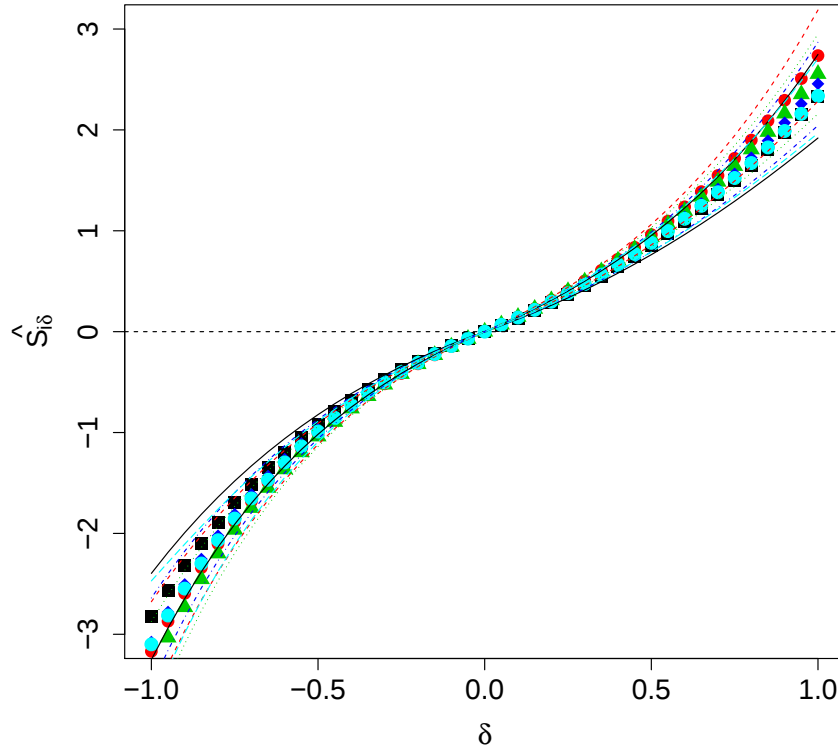
Table 3.11: Estimated Sobol' indices for the hyperplane with different spreads case

The weak first order indices (less than 3% of the variance explained) and the high total indices assess that all variables are influential in interaction with the others, and that no variable is influential on its own. All the total indices are approximatively equal showing that this SA method gives to each equally contributing variable the same importance, despite their different spread.

3.4.5.3 DMBRSI

One can notice that the different inputs follow various distributions (unlike the other examples), thus the question of "equivalent" perturbation arises. Due to this non-similarity of the distributions, only a (modified) mean shift, a quantile shift and a parameter shift will be applied on this test case. It has been discussed further in Section 3.3.1.3.

Mean shifting As stressed in Section 3.3.1.3 the choice has been made to shift the mean relatively to the standard deviation, hence including the spread of the various inputs in their respective perturbation. So for any input, the original distribution is perturbed so that its mean is the original one plus δ times its standard deviation, δ ranging from -1 to 1 with 40 points. The results of the numerical experiment are displayed in Figure 3.25.

Figure 3.25: Estimated indices $\hat{S}_{i\delta}$ for the hyperplane with different spreads case with a mean shifting

The indices have similar values for similar perturbations, thus assessing the equal impact of the variables. However this information was obtained with a fine tuning of the perturbations.

Quantile shifting As in the previous numerical experiments, the 1st, 2nd and the 3rd quartiles altogether with the 5th and 95th percentiles were perturbed. As the graphs behave in a similar way, only one is displayed in Figure 3.26.

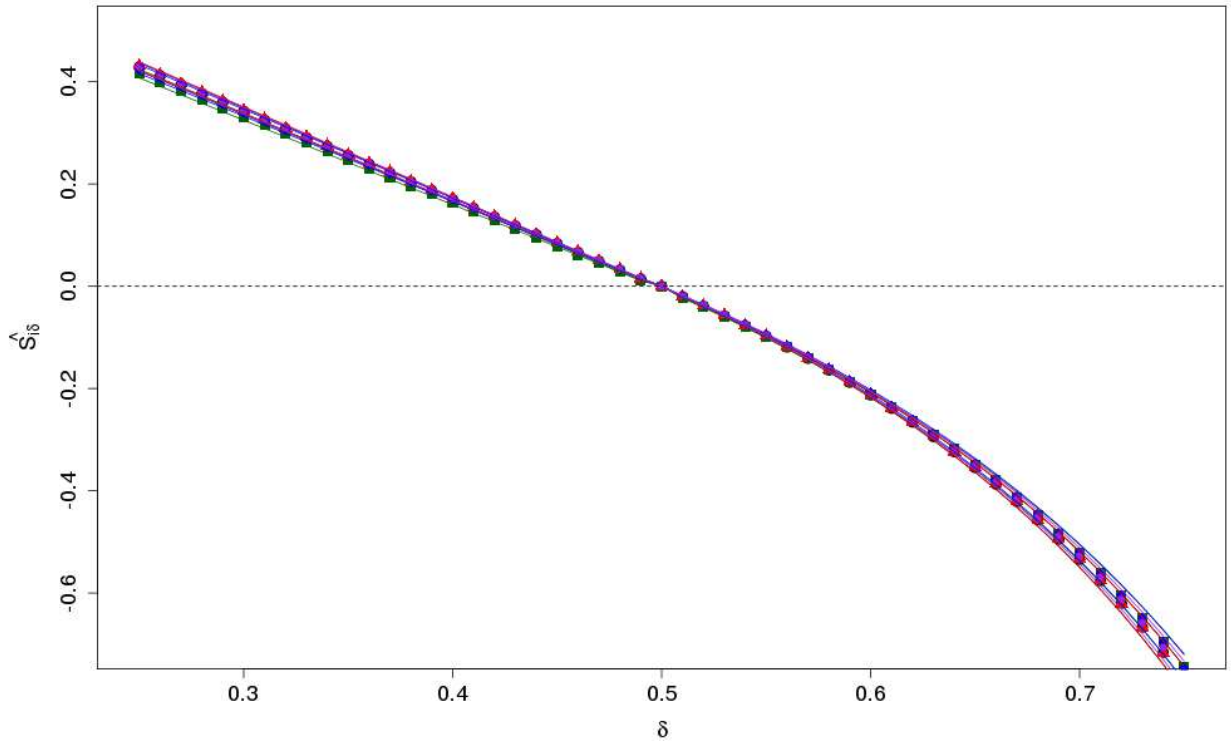


Figure 3.26: Estimated indices $\hat{S}_{i\delta}$ for the hyperplane with different spreads case with a median shifting

The perturbation of the 2nd quantile affects all the variables in the same way, despite their different distributions. This shows that the quantile perturbation method gives to each equally contributing variable the same importance.

Additionally, we can conclude the following on the application of the quantile perturbation on monotonic cases (3.4.2 to 3.4.5):

- the graphs for the median perturbation are similar to the ones relative to a mean perturbation.
- when a left-hand quantile α_1 (if $\alpha_1 < 50\%$) is influent (meaning a perturbation of $\delta\%$ of this quantile produces an index superior to a threshold t) then $\alpha_2 < \alpha_1$ has more influence. In the case of a right-hand quantile (if $\alpha_1 > 50\%$) then $\alpha_2 > \alpha_1$ has more influence.

Parameters perturbation The model is driven by 10 parameters: a variance and a standard deviation for each Gaussian input. Each of these parameters is perturbed and the estimated indices are plotted in function of the Hellinger distance in Figure 3.27 as explained in Figure 3.8. 95% confidence intervals are provided as well. As the graph gets too complicated for an adequate representation, only three variables are plotted: X_1 (black), X_3 (red) and X_5 (blue). As usual, the

indices associated with the means are plotted as dots and the indices associated with the standard deviations are plotted as triangles.

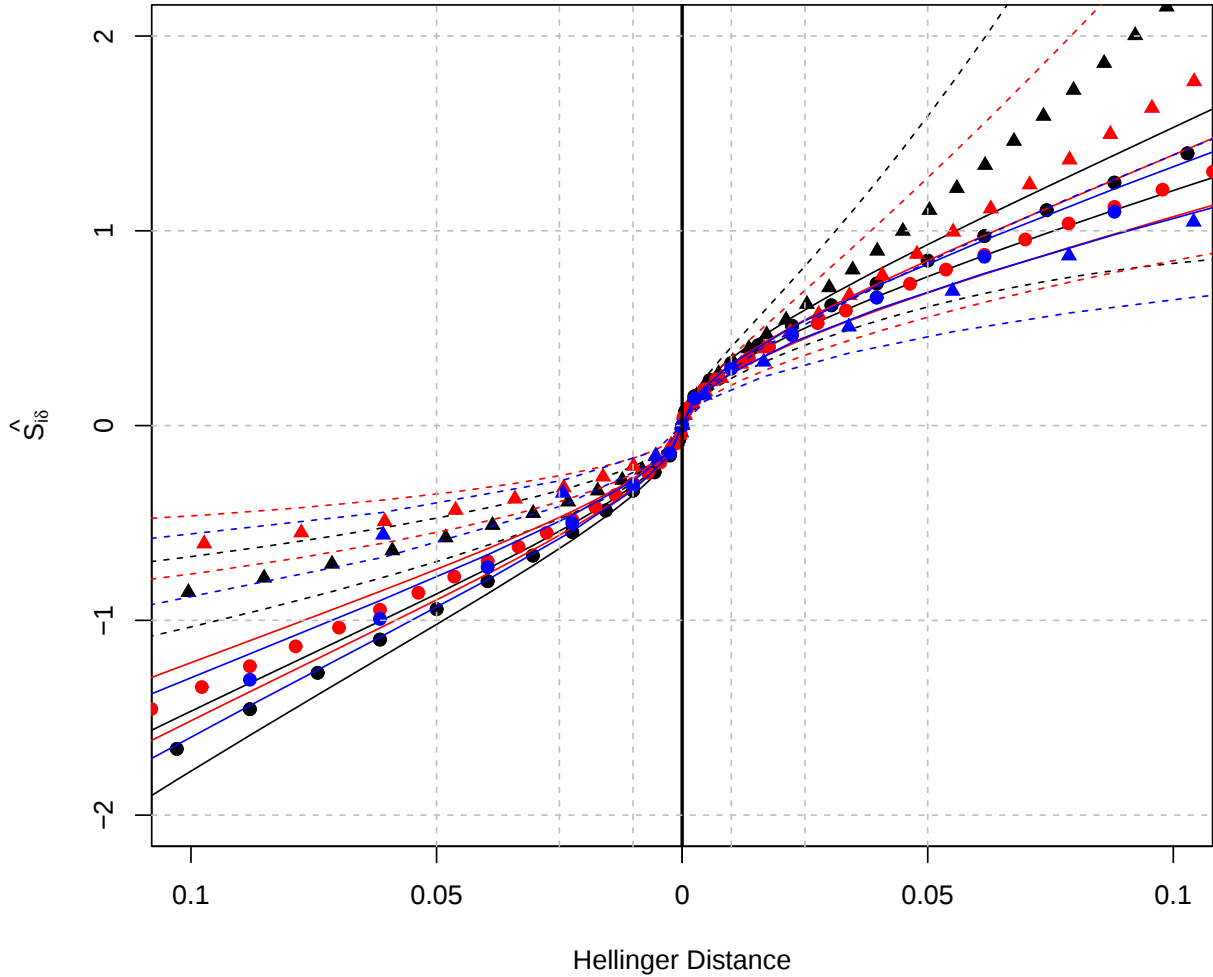


Figure 3.27: Parameters perturbation on the hyperplane with different spreads case. Dots are for means, triangle for the standard deviations. Black is for X_1 , red is for X_3 and blue is for X_5 .

	$X_i \sim \mathcal{N}(\mu = 0, \sigma = 2)$		$X_i \sim \mathcal{N}(\mu = 0, \sigma = 6)$		$X_i \sim \mathcal{N}(\mu = 0, \sigma = 10)$	
	$\mu \sigma = 2$	$\sigma \mu = 0$	$\mu \sigma = 6$	$\sigma \mu = 0$	$\mu \sigma = 10$	$\sigma \mu = 0$
$H^2(X_i, X_{i\delta}) = 0$	0	2	0	6	0	10
$H^2(X_i, X_{i\delta}) = 0.01$	0.400/-0.400	1.736/2.299	1.193/-1.193	5.208/6.897	1.989/-1.989	8.679/11.521
$H^2(X_i, X_{i\delta}) = 0.025$	0.634/-0.634	1.597/2.499	1.898/-1.898	4.790/7.496	3.163/-3.163	7.985/12.526
$H^2(X_i, X_{i\delta}) = 0.05$	0.900/-0.900	1.451/2.748	2.695/-2.695	4.353/8.245	4.492/-4.492	7.255/13.784
$H^2(X_i, X_{i\delta}) = 0.1$	1.281/-1.281	1.262/3.158	3.839/-3.839	3.785/9.475	6.398/-6.398	6.308/15.853

Table 3.12: Hellinger distance in function of the parameter perturbation

This figure leads to several comments and needs to be interpreted with table 3.12. Increasing any parameter leads to an increase of the failure probability whereas diminishing any parameter leads to a reduction of the failure probability. When increasing the parameters, indices are badly separated. One can however see that the confidence intervals associated to the means are narrower than the ones

associated to the standard deviations. On the other side of the graph, when reducing the parameters, an "equivalent" (in the H^2 sense) reduction of the mean has more impact (on the reduction of the failure probability) than a reduction of the standard deviations. The confidence intervals (for the means and for the standard deviations) are well separated. In all cases, the confidence intervals prevent from concluding that any variable is more influential than another. However, the indices for the first variable (black) seem a bit lower than the one associated to the other inputs in the decreasing case.

Conclusion and discussion When shifting the mean with small perturbations, all the variables are ranked with the same importance. We must insist that this result is obtained in shifting the mean including the spread of the various inputs in their respective perturbation. All the variables seem to have the same influence when shifting their quantiles. Similarly, a parameter perturbation does not allow to say that a variable is more influential than another - but this might be caused by numerical noise. Supplementary numerical experiments must be conducted on this topic.

3.4.6 Tresholded Ishigami function

A modified (thresholded) version of the Ishigami function will be considered in this subsection, as defined in Appendix B.2. Remind that all variables are independent Uniform with support $[-\pi, \pi]$. Finally, the failure probability is roughly $\hat{P} = 5.89 \times 10^{-3}$.

3.4.6.1 Importance factors

The algorithm FORM converges to an incoherent design point $(6.03, 0.1, 0)$ in 50 function calls, giving an approximate probability of $\hat{P}_{FORM} = 0.54$. The importance factors are displayed in Table 3.13. The bad performance of FORM is expected given that the failure domain consists in six separate domains and that the function is highly non-linear, leading to optimization difficulties. The design point is aberrant, therefore the importance factors results for SA are incorrect. Notice that the user is not warned that the result is incorrect.

Variable	X_1	X_2	X_3
Importance factor	$1e^{-17}$	1	0

Table 3.13: Importance factors for Ishigami function

3.4.6.2 Sobol' indices

The first-order and total indices are displayed in Table 3.14 which is a reproduction of Table 1.8. The following commentary is also coming from Chapter 1.

Index	S_1	S_2	S_3	S_{T1}	S_{T2}	S_{T3}
Estimation	0.018	0.007	0.072	0.831	0.670	0.919

Table 3.14: Sobol' indices estimation for the thresholded Ishigami function

The first order indices are close to 0. The variable with the most influence on its own is X_3 , explaining 7% of the output variance. Total indices state that all the variables are of high influence. A variable ranking can be made using the total indices, ranking X_3 with the highest influence, then X_1 and then X_2 . Figure B.1 allows to understand the meaning of the total indices. Each variable

“causes” the failure event on a restricted portion of its support. On the other hand, the knowledge of a single variable does not allow to explain the variance of the indicator, thus the weakness of first-order indices. The fact that the failure points are grouped in narrow strips can only be explained by the 3 variables together, thus the high third order index.

3.4.6.3 DMBRSI

The method presented throughout this chapter is applied on the thresholded Ishigami function. As previously, a MC sample of size 10^5 is used to estimate both the failure probability and the indices with all the perturbations. There are 574 failing points therefore the failure probability is estimated by $\hat{P} = 5.74 \times 10^{-3}$. The order of magnitude here is quite good. As for the hyperplane test case, a mean shifting and a variance shifting are applied at first, followed by a quantile perturbation. The parameters perturbation case is then discussed.

Mean shifting For the mean shifting (see Equation (3.15)), the variation domain for δ ranges from -3 to 3 with 60 points - numerical consideration forbidding to choose a shifted mean closer to the endpoints. The results of the estimation of the indices $\hat{S}_{i\delta}$ are plotted in Figure 3.28.

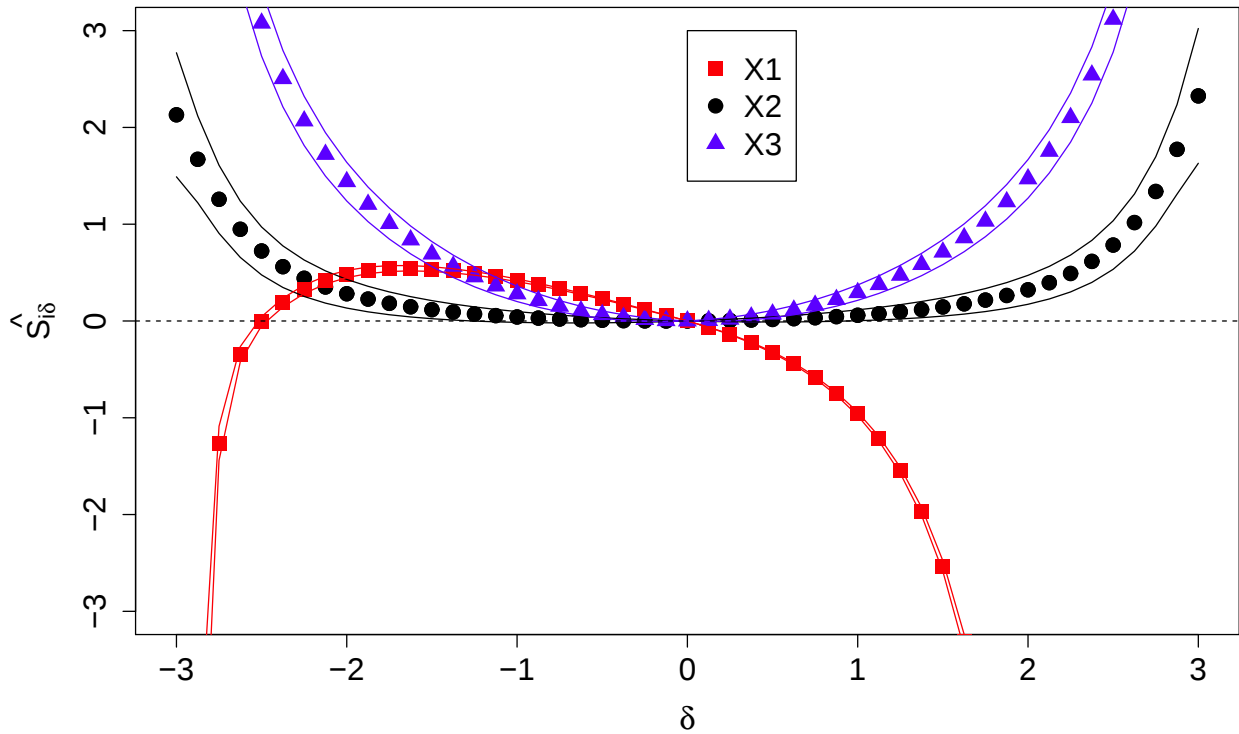


Figure 3.28: Estimated indices $\hat{S}_{i\delta}$ for the thresholded Ishigami function with a mean shifting

A perturbation of the mean for X_2 and X_3 will increase the failure probability, though the impact for the same mean perturbation is stronger for X_3 ($\hat{S}_{3,-3}$ and $\hat{S}_{3,3}$ approximately equal respectively 9.5 and 10, Figure 3.28). On the other hand, the indices concerning X_1 show that a mean shift between -1 and -2 increases the failure probability, whereas an increasing of the mean or a large decreasing strongly diminishes the failure probability ($\hat{S}_{1,3}$ approximately equals -7.10^{11}). Therefore, Figure 3.28 leads to two conclusions. First, the failure probability can be strongly reduced

when increasing the mean of the first variable X_1 (this is also provided by Figure B.1 wherein all failure points have a negative value of X_1). Second, any change in the mean for X_2 or X_3 will lead to an increase of the failure probability. The confidence intervals are well separated, except in the -1 to 1 zone. One can notice that the confidence interval associated to X_2 contains 0 between values of δ from -1.5 to 1.5 , thus the associated indices might be null in these case. This has to be taken into account when assessing the relative importance of X_2 .

Variance shifting For variance shifting, the variation domain for V_{per} ranges from 1 to 5 with 40 points. Let us recall that the original variance is $\text{Var}[X_i] = \pi^2/3 \simeq 3.29$. The modified pdf when shifting the variance and keeping the same expectation is proportional to a truncated Gaussian when decreasing the variance. When increasing the variance, the perturbed distribution is a symmetrical distribution with 2 modes close to the endpoints of the support (see Figure 3.3). The results of the estimation of the indices $\widehat{S}_{i,V_{\text{per}}}$ are plotted in Figure 3.29. The upper figure is a zoom where the $\widehat{S}_{i,V_{\text{per}}}$ axis lies into $[-0.5, 0.5]$. The lower figure shows almost the whole range variation for $\widehat{S}_{i,V_{\text{per}}}$. The curves cross for the value of V_{per} that corresponds to the original variance, namely $\pi^2/2$.

Figure 3.29 (upper part) shows that a change in the variance has little effect on X_2 and X_1 , though the change is of opposite effect on the failure probability. However, considering that the indices $\widehat{S}_{2,V_{\text{per},i}}$ and $\widehat{S}_{1,V_{\text{per},i}}$ lie between -0.4 and 0.4 , one can conclude that the variance of these variables are not of great influence on the failure probability. On the other hand, Figure 3.29 (lower part) shows that any reduction of $\text{Var}[X_3]$ strongly decreases the failure probability, and that an increase of the variance slightly increases the failure probability. This is relevant with the expression of the failure surface, as X_3 is fourth powered and multiplied by the sinus of X_1 . A variance decreasing as formulated gives a distribution concentrated around 0. Decreasing $\text{Var}[X_3]$ shrinks the concerned term in $G(\mathbf{X})$. Therefore it reduces the failure probability. The confidence intervals associated to X_3 are broadly separated from the others.

Quantile shifting First, the 5th percentile is perturbed and the result is displayed in Figure 3.30.

This graph shows that for variable X_1 , an increase of the weight of the right-hand tail diminishes the failure probability and a decrease of the weight affects positively the failure probability. It is the opposite for variable X_2 and X_3 : an increase of the weight of the left-hand tail increases the failure probability and a decrease of the weight decreases the failure probability. The effect is stronger for variable X_3 .

Then, the first quartile is perturbed. The results of the experiment are plotted in Figure 3.31.

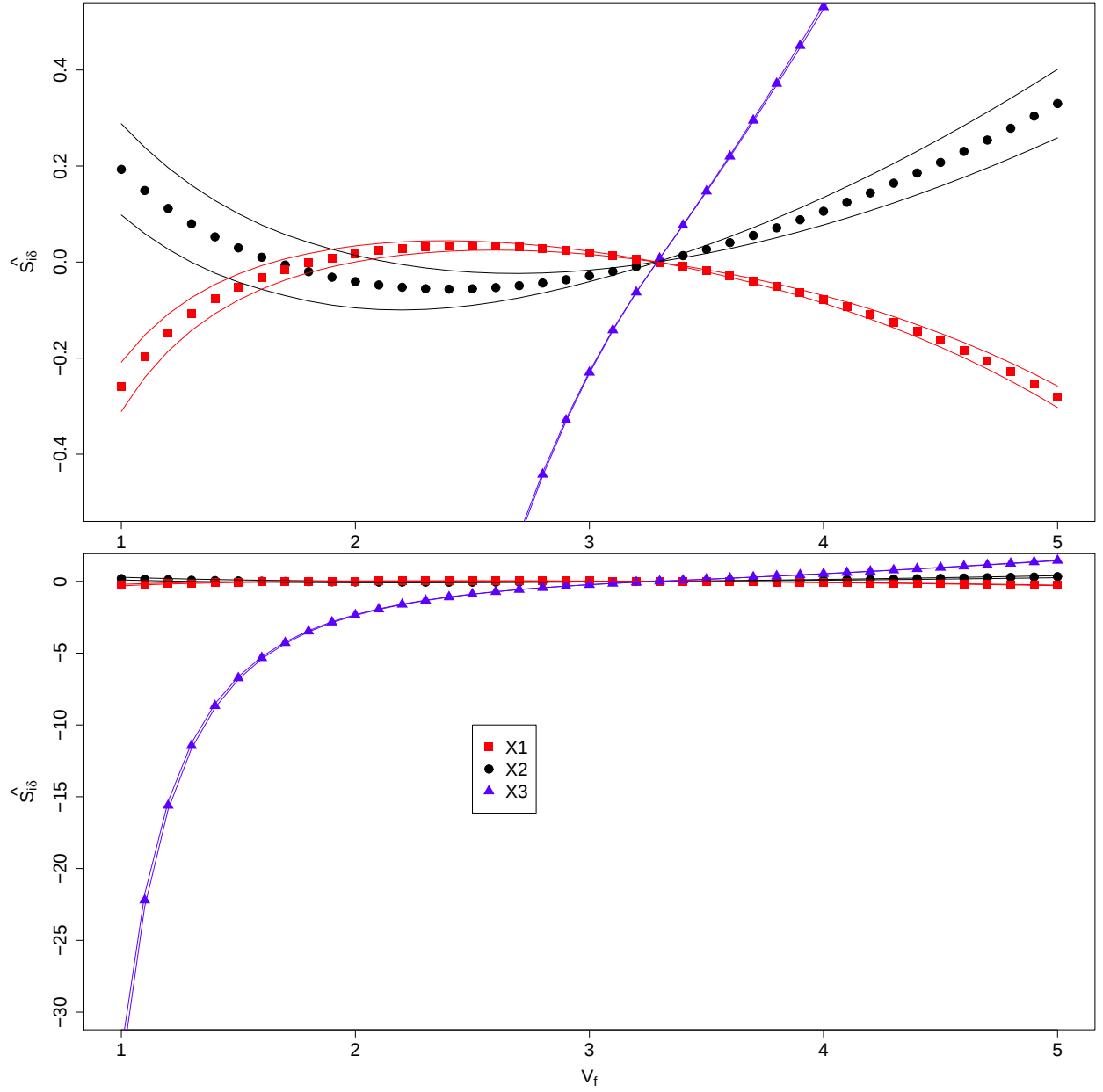
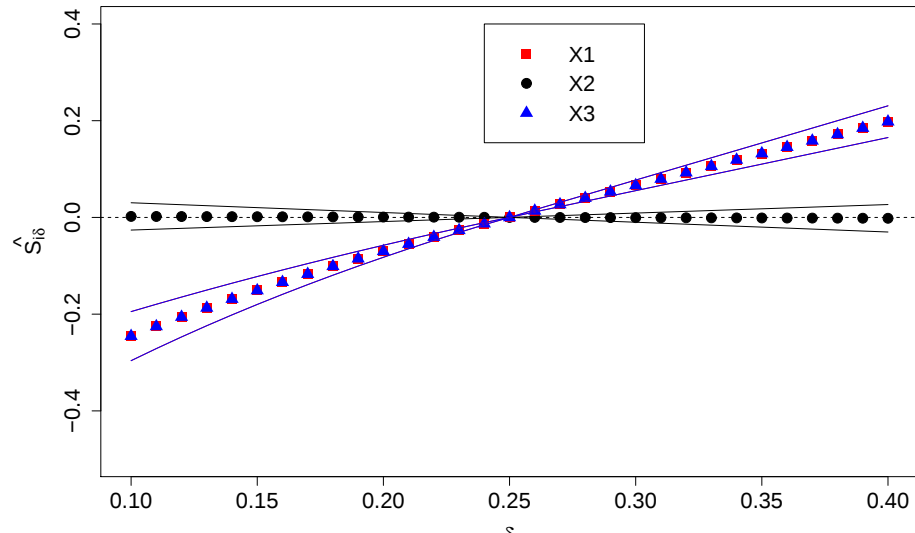


Figure 3.29: Estimated indices $\widehat{S}_{i, V_{\text{per}}}$ for the thresholded Ishigami function with a variance shifting



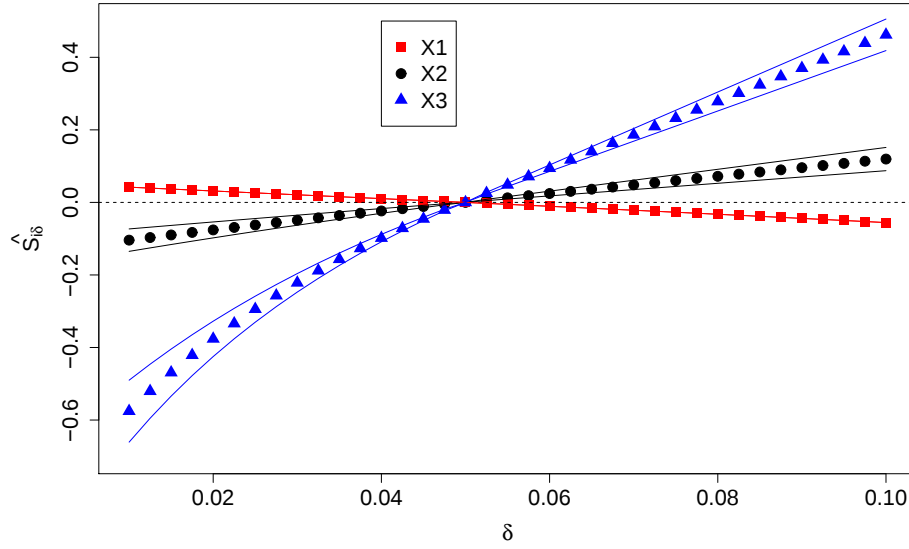


Figure 3.30: 5th percentile perturbation on the thresholded Ishigami test case

This graph shows that a 1st quartile perturbation of variable X_2 has no effect on the failure probability, for the considered range of variation. It also shows that variables X_1 and X_3 behave the same when the 1st quartile is perturbed: an increase of the weight of the left-hand tail increases the failure probability and a decrease of the weight decreases the failure probability.

It is interesting to note that the impact of the 5%-quantile perturbation of X_1 produces a different effect than a perturbation on the 1st quartile. It means that the relationship established for the monotonic case is not valid in this non-monotonic case.

The median is perturbed next and the results are shown in Figure 3.32.

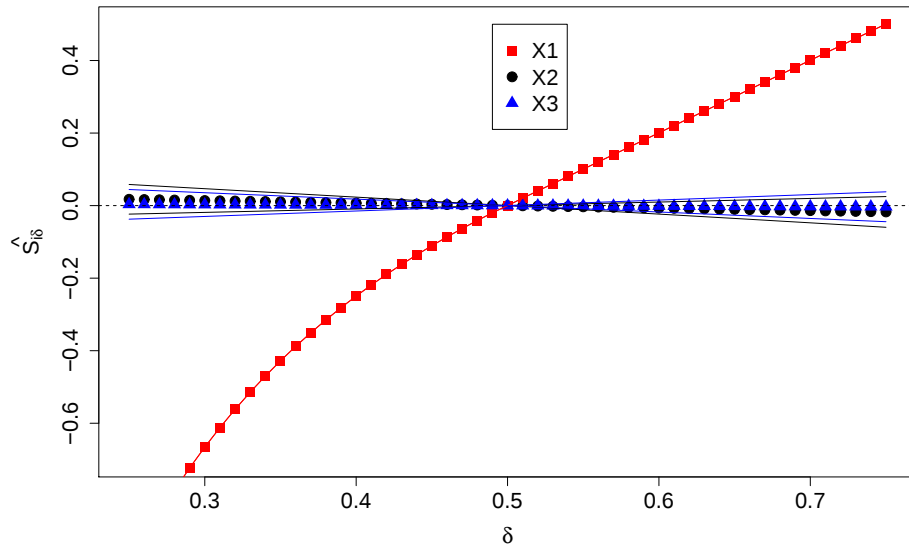


Figure 3.32: Median perturbation on the thresholded Ishigami test case

As it comes to a median perturbation, only variable X_1 produces effects. A decrease (increase) of the weight of the left-hand tail reduces (increases) the failure probability. 0 is included within the confidence intervals for variables X_2 and X_3 .

The third quartile is perturbed next and the results are displayed in Figure 3.33.

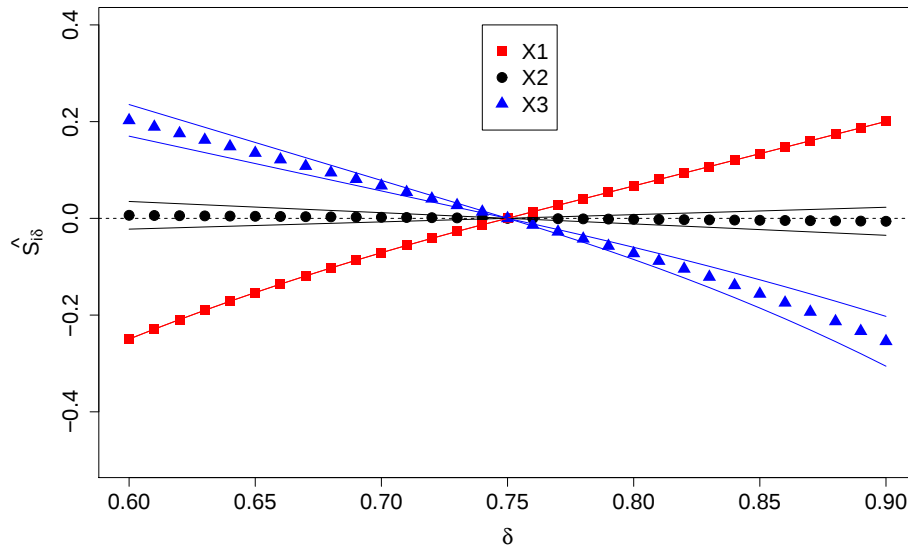


Figure 3.33: 3rd quartile perturbation on the thresholded Ishigami test case

An increase of the weight of the right-hand tail of variable X_1 increases the failure probability whereas it reduces the failure probability for variable X_3 , with the same order of magnitude. The effect is reversed when decreasing the weight. A perturbation of the third quartile of variable X_2 has no effect on the failure probability.

Finally, the 95th percentile is perturbed and the results are displayed in Figure 3.34.

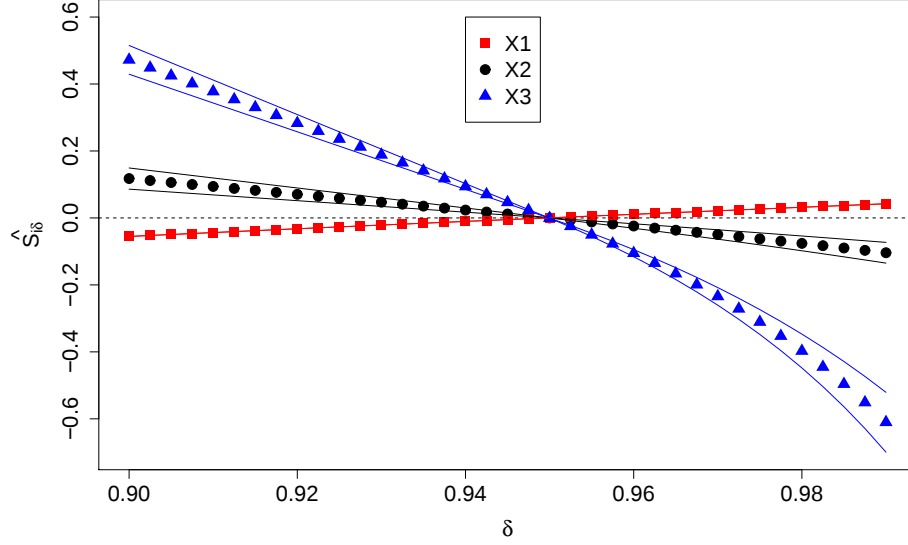


Figure 3.34: 95th percentile perturbation on the thresholded Ishigami test case

This last figure shows the higher influence of the right-hand quantile of X_3 over the two other variables. Precisely, increasing the weight of the 95%-quantile (which is equivalent to decreasing the weight of the right-hand tail) reduces the failure probability for variables X_2 and X_3 whereas the failure probability increases for X_1 . The effect is the opposite when decreasing the weight of the 95%-quantile.

This non-monotonic case shows that it is important to test several configurations of quantile perturbation before assessing the importance or non-influence of a variable.

Parameters perturbation The methodology presented in subsection 3.3.2 is tested here. The model is driven by 6 parameters: a minimum and maximum boundaries for each Uniform input. Here, we must stress a limitation of the method. The parameters of the inputs define their support. Yet, due to the conditions in Lemma 3.2.1, the support of the perturbed input cannot be broader than the one of the initial input. On this test case, this amounts to saying that the parameters perturbations can only lead to a support reduction, i.e. increasing the minimum and diminishing the maximum. Specifically, the parameters are perturbed so that the minimum varies from $-\pi$ to 0 and the maximum varies from π to 0. The result of such perturbations is presented in Figure 3.35 and Figure 3.36. 95% confidence intervals are provided as well. The amplitude of the perturbation given the Hellinger distance is given in Table 3.15.

At first in this figure we focus on small perturbations of the parameters, so that the deviation is no broader than 0.1 in Hellinger distance (refer to Table 3.15 for the equivalent in terms of parameters). On the right-hand of the graph are plotted (as triangles) the indices corresponding to an increase of the minimum bound of the inputs. On the left-hand of the graph are plotted (as dots) the indices corresponding to a decrease of the maximal bound of the inputs.

It can be seen that the indices are symmetrical. Increasing (diminishing) the minimum (maximum) for variable X_1 slightly increases the failure probability. On the other hand, increasing (diminishing) the minimum (maximum) for variable X_2 slightly decreases the failure probability. However shifting the parameters of variable X_3 produce the following effects: increasing its minimum until 2.771 (Hellinger distance 0.06) diminishes the failure probability (almost dividing it by

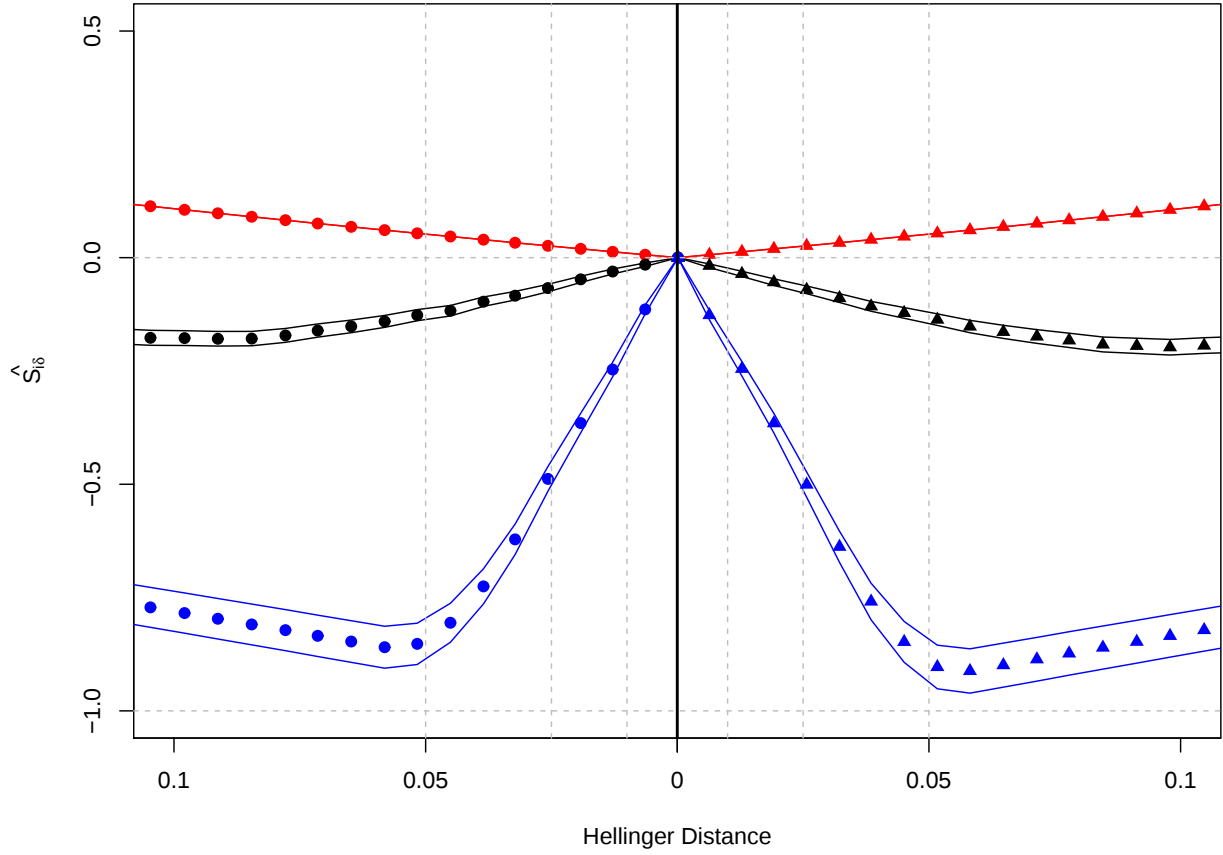


Figure 3.35: Parameters perturbation on the thresholded Ishigami test cas. Triangles correspond to a minimum bound, dots to a maximum bound. X_1 is plotted in red, X_2 in black and X_3 in blue.

2). Then, an increasing of the minimum is reflected by a slightly lower diminution of the failure probability. The effect is symmetrical when decreasing the minimum of variable X_3 .

Figure 3.36 focuses on large perturbations of the parameters (at most, the minimum and the maximum worth 0). This figure essentially shows that an increase of the minimum of variable X_1 strongly diminishes the failure probability. On the other hand, a decrease of the minimum of variable X_1 slightly increases the failure probability. When dealing with variable X_2 , the symmetry of the effects can be seen. When increasing the minimum, it diminishes the failure probability at first then it increases it. Finally, setting the minimum (or maximum) to 0 has no impact on the failure probability. Concerning variable X_3 , the attenuation of the decrease in failure probability described in Figure 3.35 goes on until the minimum (maximum) worth 0 - the impact on the failure probability is then null.

From Figures 3.36, 3.35 and Table 3.15, it can be concluded that the most influential parameters when dealing with small perturbations are the ones related to X_3 . When dealing with large perturbation of parameters, the minimum of X_1 is the most influential parameter. This is confirmed by Figure B.1.

Conclusion and discussion This non-linear case has shown that:

- When dealing with a mean perturbation, the failure probability can be strongly reduced when

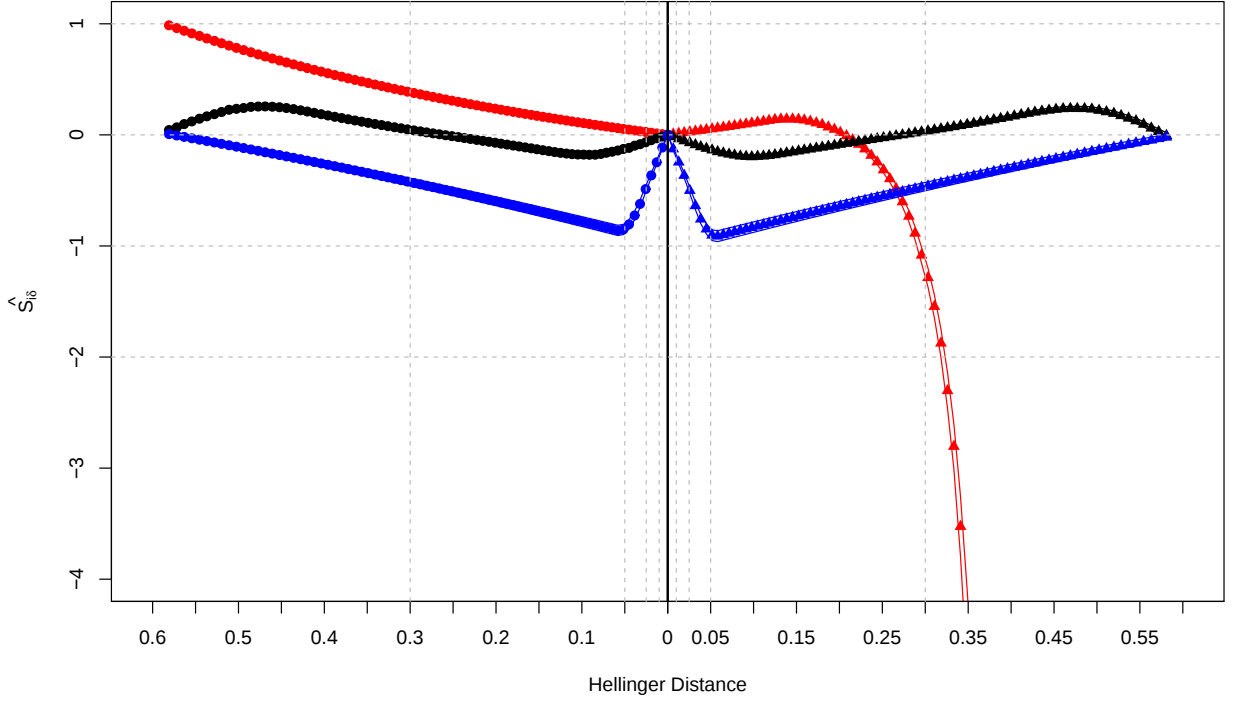


Figure 3.36: Parameters perturbation on the thresholded Ishigami test case. Triangles correspond to a minimum bound, dots to a maximum bound. X_1 is plotted in red, X_2 in black and X_3 in blue.

	$X_i \sim \mathcal{U}(\min = -\pi, \max = \pi)$	
	min max = π	max min = $-\pi$
$H^2(X_i, X_{i\delta}) = 0$	$-\pi$	π
$H^2(X_i, X_{i\delta}) = 0.01$	-3.079	3.079
$H^2(X_i, X_{i\delta}) = 0.025$	-2.985	2.985
$H^2(X_i, X_{i\delta}) = 0.05$	-2.832	2.832
$H^2(X_i, X_{i\delta}) = 0.1$	-2.529	2.529
$H^2(X_i, X_{i\delta}) = 0.3$	-1.398	1.398

Table 3.15: Hellinger distance in function of the parameter perturbation

increasing the mean of X_1 . Any change in the mean for X_2 or X_3 will lead to an increase of the failure probability.

- When dealing with a variance perturbation, any reduction of $\text{Var}[X_3]$ strongly decreases the failure probability. The impact of the other variables is negligible in this case.
- When dealing with a quantile perturbation, it is important to test several configurations before assessing the importance or non-influence of a variable. In particular, the influence of the median of X_1 can be noticed, altogether with the tails of X_3 . X_2 has a smaller influence.
- When perturbing the parameters, a limitation of the method has been highlighted (constraint on the support of the perturbed density). The various influences of the parameters have been noticed, especially the broad influence of the minimum of X_1 when dealing with large perturbations, and the parameters conducting X_3 when dealing with small perturbations.

Additionally, we argue that it is of prime importance to keep in mind the shape of the perturbed density when interpreting the figures.

3.4.7 Flood test case

This test case has been described in Appendix B.3. As stressed in the appendix, the inputs follows different distributions and the failure probability is roughly $\hat{P} = 7.88 \times 10^{-4}$.

3.4.7.1 Importance factors

The algorithm FORM converges to a design point $(1.72, -2.70, 0.55, -0.18)$ in 52 function calls, giving an approximate probability of $\hat{P}_{FORM} = 5.8 \times 10^{-4}$. The importance factors are displayed in Table 3.16.

Variable	Q	K_s	Z_v	Z_m
Importance factor	0.246	0.725	0.026	0.003

Table 3.16: Importance factors for the flood case

FORM assesses that K_s is of extremely high influence, followed by Q that is of medium influence. Z_v has a very weak influence and Z_m is negligible. It can be noticed that the estimated failure probability is twice as small as the one estimated with crude MC, but remains in the same order of magnitude.

3.4.7.2 Sobol' indices

The first-order and total indices are displayed in Table 3.17 which is a reproduction of Table 1.9. The Sobol' indices are estimated with 2 samples of size 10^6 , using the Saltelli [87] method. The total number of function evaluations is 6×10^6 .

Index	S_Q	S_{K_s}	S_{Z_v}	S_{Z_m}	S_{TQ}	S_{TK_s}	S_{TZ_v}	S_{TZ_m}
Estimation	0.019	0.251	0	0	0.746	0.976	0.248	0.115

Table 3.17: Estimated Sobol' indices for the flood case

Considering the first order indices, Z_v and Z_m are of null influence on their own. Q is considered to have a minimal influence (2% of the variance of the indicator function) by itself, and K_s explains 25% of the variance on its own. When considering the total indices, it can be noticed that both Z_v and Z_m have a weak impact on the failure probability. On the other hand, Q has a major influence on the failure probability. K_s total index is close to one, therefore K_s explains (with or without any interaction with other variables) almost all the variance of the failure function.

Let us compare the informations provided by the Sobol' indices with the information provided by the importance factors. One cannot conclude from the total Sobol' indices that Z_m is not influential whereas the importance factors assess that this variable is of negligible influence. Additionally, the total Sobol' index associated to K_s and Q state that both these variables are of high influence whereas the importance factors state that K_s is of high influence and Q is of medium influence.

3.4.7.3 DMBRSI

Notice that the different inputs follow various distributions, thus the question of "equivalent" perturbation arises. Due to this non-similarity of the distributions, only a (modified) mean shift, a quantile shift and a parameter shift will be applied on this test case. It has been discussed further in 3.3.1.3. Additionally, a numerical trick is used to deal with truncated distributions, as stressed in Appendix D.4.

Mean shifting The choice has been made to shift the mean relatively to the standard deviation, hence including the spread of the various inputs in their respective perturbation. So for any input, the original distribution is perturbed so that its mean is the original's one plus δ times its standard deviation, δ ranging from -1 to 1 with 40 points.

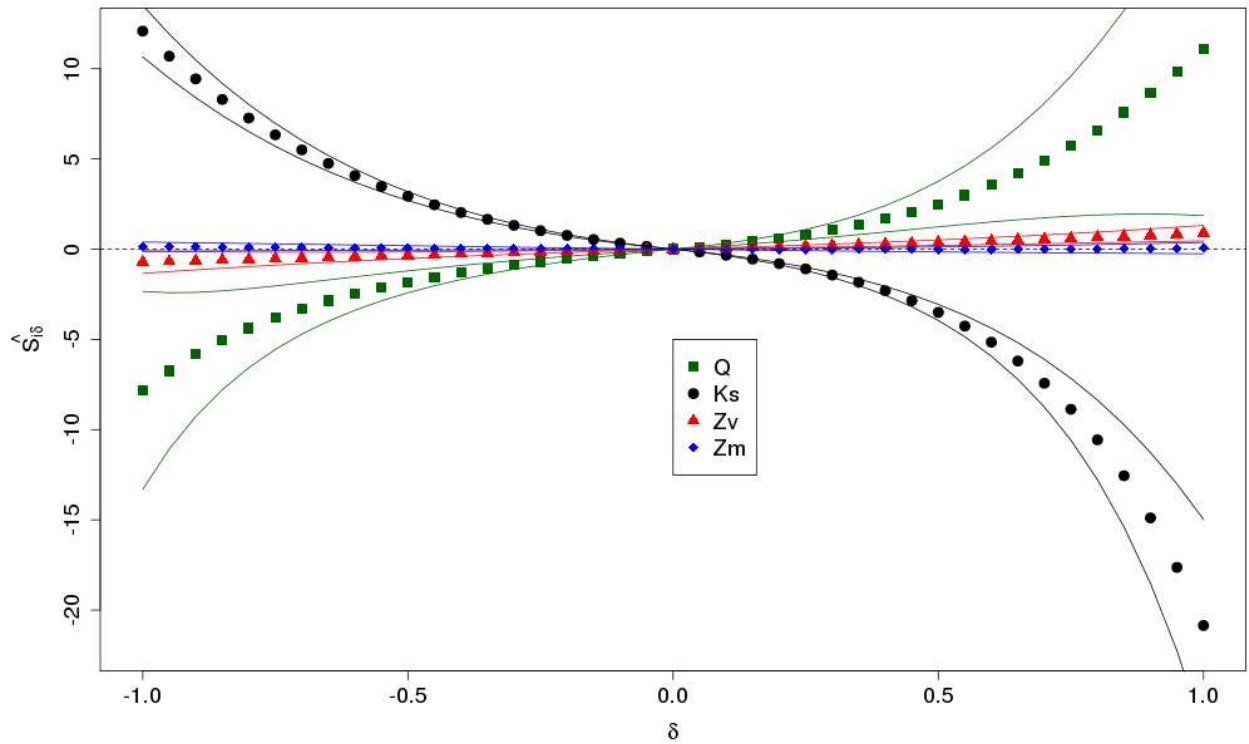
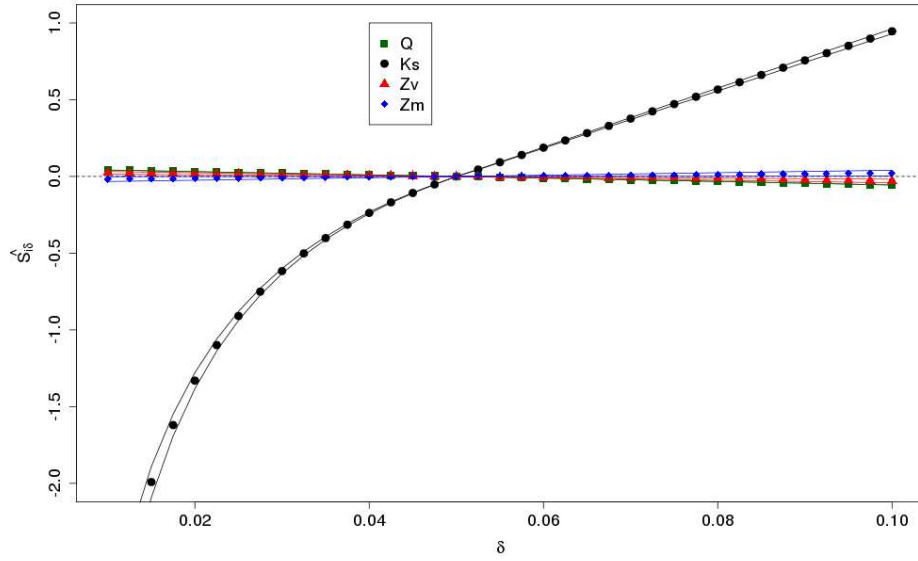


Figure 3.37: Estimated indices $\hat{S}_{i\delta}$ for the flood case with a mean perturbation

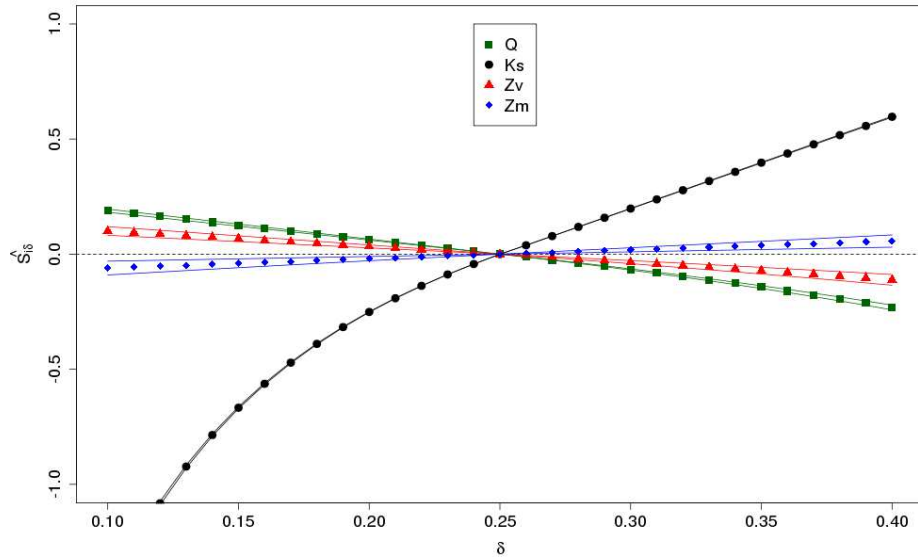
Figure 3.37 assesses that an increasing of the mean of the inputs increases the failure probability slightly for Z_v , strongly for Q , and diminishes it slightly for Z_m and strongly for K_s . This goes the opposite way when decreasing the mean. In terms of absolute modification, K_s and Q are of same magnitude, even if K_s has a slightly stronger impact. On the other hand, the effects of mean perturbation on Z_m and Z_v are negligible. The CI associated to Q and K_s are well separated from the others, except in a $\delta = -.3$ to $.3$ zone. The confidence intervals associated to Z_v and Z_m overlap. Thus even though the indices seem to have different values, it is not possible to conclude with certainty about the influence of those variables.

Quantile shifting The first quantile to be perturbed is the extreme left-hand tail, namely the 5%-quantile. The result of such a perturbation for all the variables is plotted in Figure 3.38.

Figure 3.38: 5th percentile perturbation on the flood case

When it comes to a left-hand tail perturbation, the influence of K_s over the three other variables is preponderant. In particular, a reduction of the weight of the 5th percentile to 0.015 leads to a division by 3 of the failure probability.

The 1st quartile is then perturbed and the results are plotted in Figure 3.39.

Figure 3.39: 1st quartile perturbation on the flood case

Once again when perturbing the left-hand tail, the influence of K_s is larger than the influence of the other variables.

The median of the input distributions is then perturbed, the resulting indices are plotted in Figure 3.40.

The influence of K_s is weaker than in the two previous figures, as K_s and Q have a similar

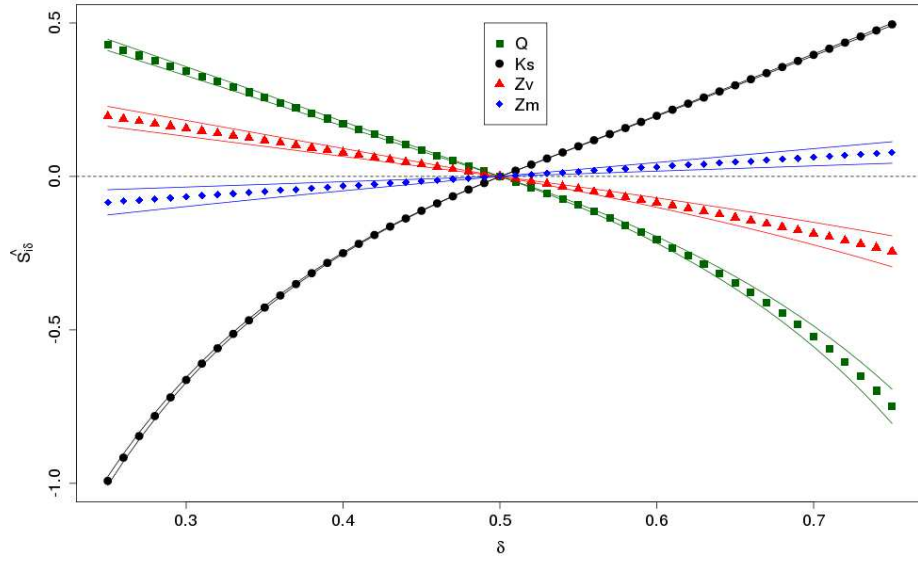


Figure 3.40: Median perturbation on the flood case

influence (although the effects of a median perturbation of these variables is reversed). Z_m has less impact on the failure probability than Z_v , when dealing with a median perturbation.

The third quartile is then perturbed and the indices are plotted in Figure 3.41.

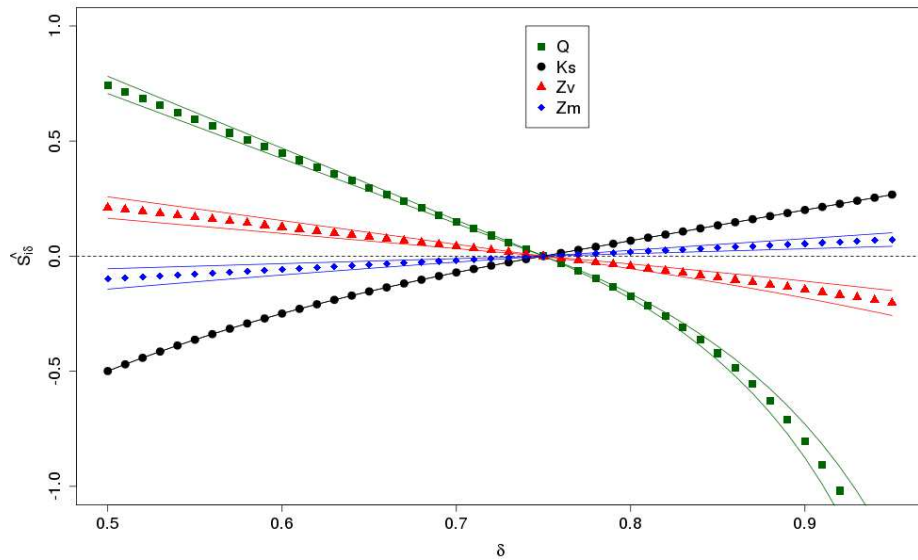
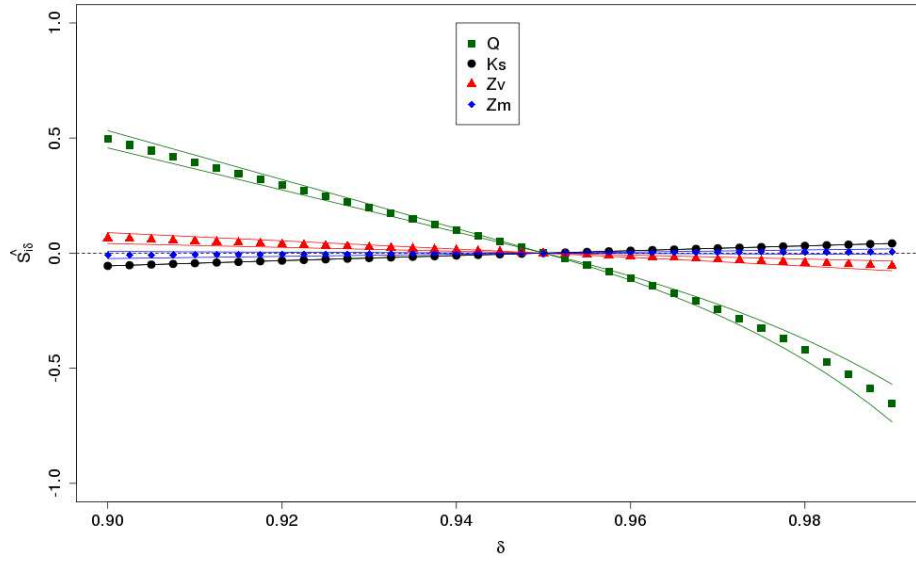


Figure 3.41: 3rd quartile perturbation on the flood case

Increasing the weight of the right hand tail (that is to say decreasing the weight of the 3rd quartile) increases the failure probability for Q and Z_v whereas it reduces the probability for Z_m and K_s . The magnitude of influence is the following: Q has most influence, then K_s and Z_v have almost the same influence, then comes Z_m .

Finally, the extreme right-hand tail is perturbed, this comes to a perturbation on the 95th percentile. Results are plotted in Figure 3.42.

Figure 3.42: 95th percentile perturbation on the flood case

This last graph shows the strongest influence of Q when perturbing extreme right-hand quantiles. More precisely, increasing the weight of the right-hand tail of Q increases the failure probability whereas it is the opposite when decreasing this weight. The impact of the other variables is much smaller.

As a conclusion, we would say that the practitioner needs to be careful when modelling the right-hand tail of Q and the left-hand tail of K_s , as the failure probability is sensitive to a variation of these two quantities. Additionally, the code seems to behave in a monotonic fashion (the indices of a given variable have the same sign all along the interval of variation).

Parameters perturbation The model is driven by 12 parameters:

- a location parameter, a scale parameter and a minimum for Q ;
- a mean, a standard deviation and a minimum for K_s ;
- a minimum, a maximum and a mode for Z_v ;
- a minimum, a maximum and a mode for Z_m .

However on this case we decide to perturb only the parameters that do not affect the support of the densities, namely the location, the scale, the mean, the standard deviation and the two modes. These parameters are perturbed and the estimated indices are plotted in function of the Hellinger distance in Figure 3.43 as explained in Figure 3.8. 95% confidence intervals are provided as well.

Table 3.18, presenting the relationship between the parameter perturbation and the Hellinger distance, is needed to interpret Figure 3.43.

Increasing the parameters value increases the failure probability when dealing with the standard deviation of K_s , the scale, the location of Q and the mode of Z_v . It decreases the failure probability when dealing with the mode of Z_m and the mean of K_s . The effect on the failure probability are reversed when decreasing the value of the parameters. The perturbation of the parameters produces a large perturbation of the failure probability for the parameters associated to K_s and for the scale

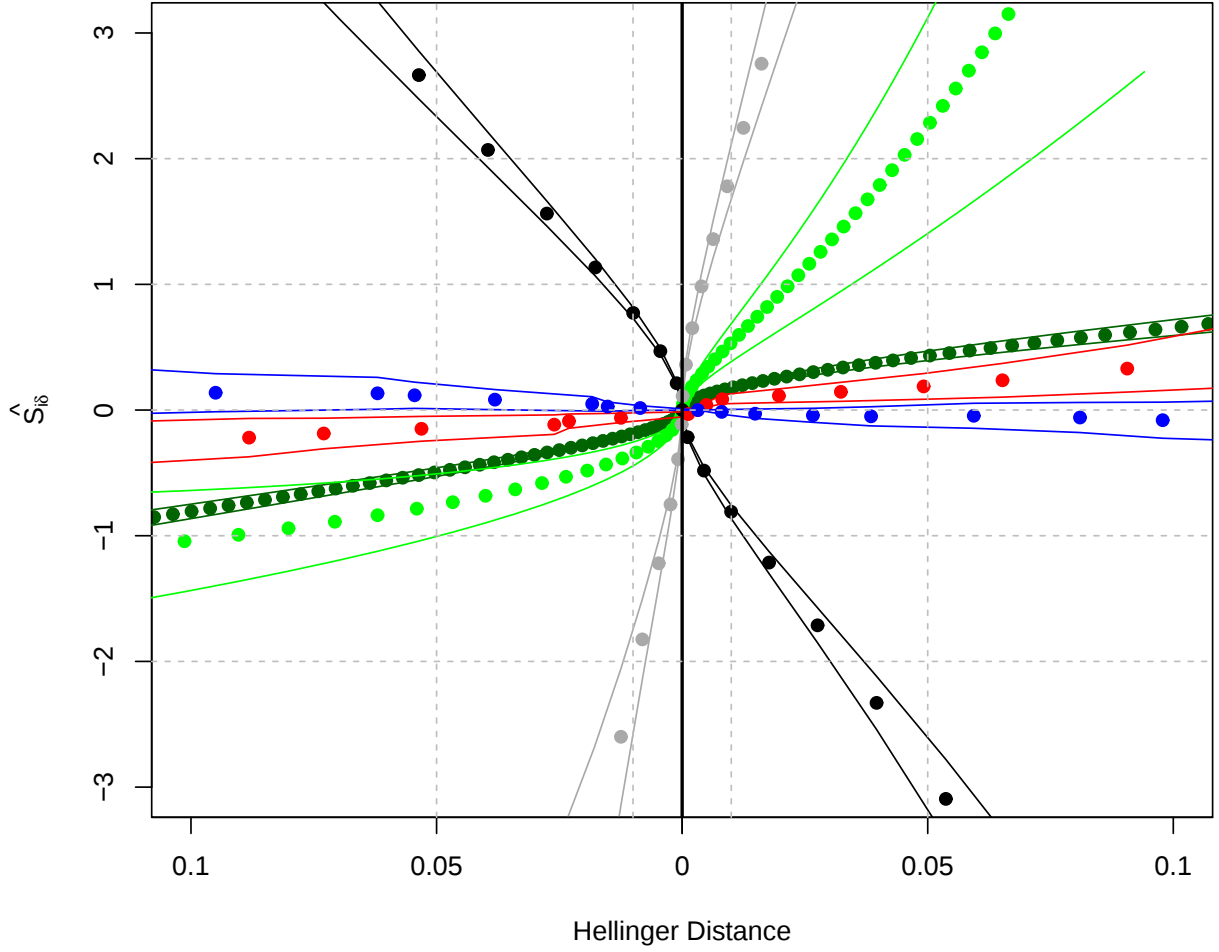


Figure 3.43: Parameters perturbation on the flood test case. The indices corresponding to Q are plotted in green: dark green for the location parameter and light green for the scale parameter. The indices corresponding to K_s are plotted as follows: black for the mean, dark grey for the standard deviation. The indices of the mode of Z_v are plotted in red while the ones corresponding to the mode of Z_m are plotted in blue.

parameter of Q . The impact on the failure probability is moderate when perturbing the location of Q , and is quasi-null when perturbing the modes of Z_v and Z_m .

It is thus of prime importance to model correctly the parameters conducting K_s , and the scale parameter of Q .

Conclusion and discussion On this test case, we can conclude the following:

- In terms of mean perturbation, the indices associated to K_s and Q have a high value.
- The quantile perturbation has shown that the right-hand tail of Q and the left-hand tail of K_s are particularly influential on the failure probability. Additionally, the code seems to behave in a monotonic fashion.

	$X_i \sim \mathcal{G}_T(\text{loc} = 1013, \text{scale} = 558, \text{min} = 0)$		$X_i \sim \mathcal{N}_T(\mu = 30, \sigma = 7.5, \text{min} = 1)$	
	loc scale = 558, min = 0	scale loc = 1013, min = 0	$\mu \sigma = 7.5, \text{min} = 1$	$\sigma \mu = 30, \text{min} = 1$
$H^2(X_i, X_{i\delta}) = 0$	1013	558	30	7.5
$H^2(X_i, X_{i\delta}) = 0.01$	893/1128	478/661	28.49/31.50	6.51/8.65
$H^2(X_i, X_{i\delta}) = 0.025$	820/1194	437/736	27.62/32.38	5.99/9.42
$H^2(X_i, X_{i\delta}) = 0.05$	732/1269	395/838	26.62/33.38	5.44/10.40
$H^2(X_i, X_{i\delta}) = 0.1$	590/1377	342/1021	25.19/34.81	4.73/12.08

	$X_i \sim \mathcal{T}(a = 49, b = 51, c = 50)$	$X_i \sim \mathcal{T}(a = 54, b = 56, c = 55)$
	$c a = 49, b = 51$	$c a = 54, b = 56$
$H^2(X_i, X_{i\delta}) = 0$	50	55
$H^2(X_i, X_{i\delta}) = 0.01$	49.79/50.21	54.79/55.21
$H^2(X_i, X_{i\delta}) = 0.025$	49.65/50.35	54.65/55.35
$H^2(X_i, X_{i\delta}) = 0.05$	49.49/50.51	49.49/50.51
$H^2(X_i, X_{i\delta}) = 0.1$	49.26/50.74	49.26/50.74

Table 3.18: Hellinger distance in function of the parameter perturbation

- The parameters perturbation has demonstrated that the parameters of K_s and the scale parameter of Q impact most the output.

This more realistic test case has shown that the DMBRSI provide several complementary informations.

3.5 Improving the DMBRSI estimation

This chapter has presented a new SA methodology based on density perturbations. For the sake of simplicity, we have considered a crude Monte-Carlo framework. However, this consideration might be unrealistic when dealing with real application cases where the number of function calls is limited. We thus propose in this Section to improve the DMBRSI estimation with importance sampling (Section 3.5.1) and with subset simulation (Section 3.5.2).

3.5.1 Coupling DMBRSI with importance sampling

3.5.1.1 Estimating P_f with IS

Denoting $\tilde{f}$ a d -dimensional importance density such that $\text{Supp}(\tilde{f}) \supseteq \text{Supp}(f)$. Suppose one has an i.i.d. N-sample with pdf $\tilde{f}$, denoted $\mathbf{x}^n$ with n going from 1 to N .

The failure probability P_f can be estimated with Importance Sampling method (see Section 1.2.1.3) and the associated estimator with N function calls is:

$$\hat{P}_{NIS} = \frac{1}{N} \sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f(\mathbf{x}^n)}{\tilde{f}(\mathbf{x}^n)}. \quad (3.32)$$

One can show that:

$$\text{Var} [\hat{P}_{NIS}] = \frac{1}{N} \text{Var}_{\tilde{f}} \left[1_{\{G(\mathbf{X}) < 0\}} \frac{f(\mathbf{X})}{\tilde{f}(\mathbf{X})} \right] = \frac{1}{N} \left(\int 1_{\{G(\mathbf{x}) < 0\}} \frac{f^2(\mathbf{x})}{\tilde{f}(\mathbf{x})} d\mathbf{x} - P_f^2 \right) \quad (3.33)$$

NB : the variance reduction from IS is not straightforward, one should compare $\text{Var}_{\tilde{f}} \left[1_{\{G(\mathbf{X}) < 0\}} \frac{f(\mathbf{X})}{\tilde{f}(\mathbf{X})} \right]$ and $\text{Var}_f [1_{\{G(\mathbf{X}) < 0\}}]$ to conclude, as stressed in Section 1.2.1.3.

3.5.1.2 Estimating $P_{i\delta}$ with IS

Let us recall that

$$P_{i\delta} = \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(\mathbf{x}) d\mathbf{x}. \quad (3.34)$$

Thus the expression of $P_{i\delta}$ using IS is:

$$P_{i\delta} = \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} \frac{f(\mathbf{x}) \tilde{f}(\mathbf{x})}{\tilde{f}(\mathbf{x})} d\mathbf{x}. \quad (3.35)$$

Then supposing one has an i.i.d. N-sample with pdf $\tilde{f}$, denoted $\mathbf{x}^n$ as previously, one can estimate $P_{i\delta}$ with:

$$\hat{P}_{i\delta NIS} = \frac{1}{N} \sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)} \frac{f(\mathbf{x}^n)}{\tilde{f}(\mathbf{x}^n)}. \quad (3.36)$$

It is straightforward that the expectation of $\hat{P}_{i\delta NIS}$ is $P_{i\delta}$.

One is obviously interested in the variance of such an estimate, therefore one has:

$$\text{Var}_{\tilde{f}} \left[1_{\{G(\mathbf{X}) < 0\}} \frac{f_{i\delta}(X_i)}{f_i(X_i)} \frac{f(\mathbf{X})}{\tilde{f}(\mathbf{X})} \right] = \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}^2(x_i)}{f_i^2(x_i)} \frac{f^2(\mathbf{x})}{\tilde{f}(\mathbf{x})} d\mathbf{x} - P_{i\delta}^2. \quad (3.37)$$

Then:

$$\text{Var} [\hat{P}_{i\delta NIS}] = \frac{1}{N} \left(\int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}^2(x_i)}{f_i^2(x_i)} \frac{f^2(\mathbf{x})}{\tilde{f}(\mathbf{x})} d\mathbf{x} - P_{i\delta}^2 \right) \quad (3.38)$$

3.5.1.3 Asymptotic results

Proposition 3.5.1 *Assume the usual conditions*

- (i) $\text{Supp}(f_{i\delta}) \subseteq \text{Supp}(f_i)$,
- (ii) $\text{Supp}(\tilde{f}) \supseteq \text{Supp}(f)$
- (iii) $\int_{\text{Supp}(f_i)} \frac{f_{i\delta}^2(x)}{f_i(x)} dx < \infty$,

then

$$\hat{P}_{i\delta NIS} \xrightarrow[N \rightarrow \infty]{} P_{i\delta} \quad (3.39)$$

and

$$\sqrt{N} \left(\hat{P}_{i\delta NIS} - P_{i\delta} \right) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, \sigma_{i\delta}^2). \quad (3.40)$$

One has:

$$\sigma_{i\delta}^2 = \text{Var}_{\tilde{f}} \left[1_{\{G(\mathbf{X}) < 0\}} \frac{f_{i\delta}(X_i)}{f_i(X_i)} \frac{f(\mathbf{X})}{\tilde{f}(\mathbf{X})} \right] = \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}^2(x_i)}{f_i^2(x_i)} \frac{f^2(\mathbf{x})}{\tilde{f}(\mathbf{x})} d\mathbf{x} - P_{i\delta}^2. \quad (3.41)$$

This comes from Van der Vaart [98], 2.17.

$\sigma_{i\delta}^2$ can be consistently estimated by:

$$\hat{\sigma}_{i\delta N}^2 = \frac{1}{N} \sum_{n=1}^N \left[1_{\{G(\mathbf{x}^n) < 0\}} \frac{f_{i\delta}^2(x_i^n)}{f_i^2(x_i^n)} \frac{f^2(\mathbf{x}^n)}{\tilde{f}(\mathbf{x}^n)} - \hat{P}_{i\delta NIS}^2 \right]. \quad (3.42)$$

Proposition 3.5.2

$$\sqrt{N} \begin{pmatrix} \hat{P}_{NIS} \\ \hat{P}_{i\delta NIS} \end{pmatrix} - \begin{pmatrix} P_f \\ P_{i\delta} \end{pmatrix} \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, \Sigma_{i\delta IS}) \quad (3.43)$$

where:

$$\Sigma_{i\delta IS} = \begin{pmatrix} \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f^2(\mathbf{x})}{f(\mathbf{x})} d\mathbf{x} - P_f^2 & \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} \frac{f^2(\mathbf{x})}{f(\mathbf{x})} d\mathbf{x} - P_f P_{i\delta} \\ \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} \frac{f^2(\mathbf{x})}{f(\mathbf{x})} d\mathbf{x} - P_f P_{i\delta} & \int 1_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}^2(x_i)}{f_i^2(x_i)} \frac{f^2(\mathbf{x})}{f(\mathbf{x})} d\mathbf{x} - P_{i\delta}^2 \end{pmatrix}. \quad (3.44)$$

This comes according to Van der Vaart [98], 2.18.

We propose the following estimator for $\Sigma_{i\delta IS}$:

$$\widehat{\Sigma_{Ni\delta IS}} = \begin{pmatrix} \frac{1}{N} \left(\sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f^2(\mathbf{x}^n)}{\hat{f}(\mathbf{x}^n)} \right) - \hat{P}_{NIS}^2 & \frac{1}{N} \left(\sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f^2(\mathbf{x}^n)}{\hat{f}(\mathbf{x}^n)} \frac{f_{i\delta}(x_i)}{\hat{f}_i(x_i)} \right) - \hat{P}_{NIS} \hat{P}_{i\delta NIS} \\ \frac{1}{N} \left(\sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f^2(\mathbf{x}^n)}{\hat{f}(\mathbf{x}^n)} \frac{f_{i\delta}(x_i)}{\hat{f}_i(x_i)} \right) - \hat{P}_{NIS} \hat{P}_{i\delta NIS} & \frac{1}{N} \left(\sum_{n=1}^N 1_{\{G(\mathbf{x}^n) < 0\}} \frac{f^2(\mathbf{x}^n)}{\hat{f}(\mathbf{x}^n)} \frac{f_{i\delta}^2(x_i)}{\hat{f}_i^2(x_i)} \right) - \hat{P}_{i\delta NIS}^2 \end{pmatrix}. \quad (3.45)$$

Proposition 3.5.3 *Introducing the function $s(x, y) = (\frac{y}{x} - 1) 1_{\{y > x\}} + (1 - \frac{x}{y}) 1_{\{x > y\}}$, denoting:*

- (i) $S_{i\delta} = s(P_f, P_{i\delta})$
- (ii) $\hat{S}_{Ni\delta IS} = s(\hat{P}_{NIS}, \hat{P}_{i\delta NIS})$.

As s is differentiable in $(P, P_{i\delta})$ (see Proposition 3.2.2), one has:

$$\sqrt{N} (\hat{S}_{Ni\delta IS} - S_{i\delta}) \xrightarrow[N \rightarrow \infty]{\mathcal{L}} \mathcal{N}(0, d_s^T \Sigma_{i\delta IS} d_s). \quad (3.46)$$

The proof lies in Theorem 3.1 in Van der Vaart [98].

3.5.2 Coupling DMBRSI with subset simulation

We refer to Section 1.2.3 for more details about subset simulation. The aim of the current section is to show that it is possible to use the results of a subset simulation algorithm to estimate the quantity $P_{i\delta}$, the perturbed failure probability (see Equation 3.1).

Let us imagine, for the sake of clarity, a two-step subset where the levels are fixed in advance. Let us denote by $A, B, 0$ the thresholds to cross at the algorithm's steps, with $A > B > 0$.

We have $P_A = \int 1_{\{G(x) \leq A\}} f(x) dx$; $P_B = \int 1_{\{G(x) \leq B\}} f(x) dx$ and $P_f = \int 1_{\{G(x) \leq 0\}} f(x) dx$.

Additionally, let us remind that

$$P_{i\delta} = \int 1_{\{G(x) \leq 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(x) dx = \mathbb{E}[1_{\{G(X_{i\delta}) \leq 0\}}] = P(G(X_{i\delta}) \leq 0) \quad (3.47)$$

The algorithm starts with N points $x^{(j),1}, j = 1 \dots N$ distributed according to f , the original density. P_A can be estimated by:

$$\hat{P}_A = \frac{1}{N} \sum 1_{\{G(x^{(j),1}) \leq A\}} \quad (3.48)$$

where one has:

$$\mathbb{E} [\widehat{P}_A] = \int 1_{\{G(x) \leq A\}} f(x) dx = P_A. \quad (3.49)$$

Then, after a mutation/selection step, one has N points $x^{(j),2}, j = 1 \dots N$ distributed according to $f(x|A) = \frac{1_{\{G(x) \leq A\}} f(x)}{P_A}$. The following estimator is proposed for $P_{B|A} = \int 1_{\{G(x) \leq B|G(x) \leq A\}} f(x) dx = \frac{\int 1_{\{G(x) \leq B \cap G(x) \leq A\}} f(x) dx}{\int 1_{\{G(x) \leq A\}} f(x) dx} = \frac{\int 1_{\{G(x) \leq B\}} f(x) dx}{\int 1_{\{G(x) \leq A\}} f(x) dx} = \frac{P_B}{P_A}$:

$$\frac{\widehat{P}_B}{\widehat{P}_A} = \frac{1}{N} \sum 1_{\{G(x^{(j),2}) \leq B\}}. \quad (3.50)$$

One has:

$$\mathbb{E} \left[\frac{\widehat{P}_B}{\widehat{P}_A} \right] = \int 1_{\{G(x) \leq B\}} \frac{1_{\{G(x) \leq A\}} f(x)}{P_A} dx = \frac{\int 1_{\{G(x) \leq B\}} f(x) dx}{P_A} = \frac{P_B}{P_A} \quad (3.51)$$

After a second mutation/selection step, one has N points $x^{(j),3}, j = 1 \dots N$ distributed according to $f(x|B) = \frac{1_{\{G(x) \leq B\}} f(x)}{P_B}$. The following estimator is proposed for $P_{0|B} = \int 1_{\{G(x) \leq 0|G(x) \leq B\}} f(x) dx$:

$$\widehat{P}_{0|B} = \frac{1}{N} \sum 1_{\{G(x^{(j),3}) \leq 0\}}. \quad (3.52)$$

One can check that:

$$\mathbb{E} [\widehat{P}_{0|B}] = \int 1_{\{G(x) \leq 0\}} \frac{1_{\{G(x) \leq B\}} f(x)}{P_B} dx = \frac{\int 1_{\{G(x) \leq 0\}} f(x) dx}{P_B} = \frac{P}{P_B} \quad (3.53)$$

Finally, $P_f = P_A \times P_{B|A} \times P_{0|B,A}$. Yet $B \Rightarrow A$ thus $P_{0|B,A} = P_{0|B}$. P is estimated by:

$$\widehat{P} = \widehat{P}_A \widehat{P}_{B|A} \widehat{P}_{0|B}$$

Considering $\widehat{P}_A, \widehat{P}_{B|A}$ et $\widehat{P}_{0|B}$ as realisation of independent random variables¹ one has:

$$\mathbb{E} [\widehat{P}] = \mathbb{E} [\widehat{P}_A] \mathbb{E} [\widehat{P}_{B|A}] \mathbb{E} [\widehat{P}_{0|B}] = P_A \times \frac{P_B}{P_A} \times \frac{P}{P_B} = P.$$

Then, it is observed that:

$$P_{i\delta} = \frac{P_{i\delta}}{P_B} \frac{P_B}{P_A} P_A$$

Considering the N points $x^{(j),3}, j = 1..N$ distributed according to $f(x/B) = \frac{1_{\{G(x) \leq B\}} f(x)}{P_B}$. $\frac{P_{i\delta}}{P_B}$ is estimated by:

$$\frac{\widehat{P}_{i\delta}}{\widehat{P}_B} = \frac{1}{N} \sum 1_{\{G(x^{(j),3}) \leq 0\}} \frac{f_{i\delta}(x_i^{(j),3})}{f_i(x_i^{(j),3})}.$$

One can check that:

$$\mathbb{E} \left[\frac{\widehat{P}_{i\delta}}{\widehat{P}_B} \right] = \int 1_{\{G(x) \leq 0\}} \frac{1_{\{G(x) \leq B\}} f(x)}{P_B} \frac{f_{i\delta}(x_i)}{f_i(x_i)} dx = \frac{1}{P_B} \int 1_{\{G(x) \leq 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(x) dx = \frac{P_{i\delta}}{P_B}.$$

¹This is not the case in reality, the mutation step is just performed several times

Considering $\widehat{\frac{P_{i\delta}}{P_B}}, \widehat{\frac{P_B}{P_A}}$ et $\widehat{P_A}$ as realisation of independent random variables ² one has:

$$\mathbb{E} \left[\widehat{P_{i\delta}} \right] = \mathbb{E} \left[\widehat{\frac{P_{i\delta}}{P_B}} \right] \mathbb{E} \left[\widehat{\frac{P_B}{P_A}} \right] \mathbb{E} \left[\widehat{P_A} \right] = P_{i\delta}.$$

Conclusion To couple DMBRSI and subset simulation, one just has to perturb the points coming from the last step of the subset. However, the variance of $\widehat{P_{i\delta}}$ is intractable so far. This will be the object of further researches.

3.6 Discussion and conclusion

3.6.1 Conclusion on the DMBRSI method

The method presented in this chapter gives relevant complementary information in addition of traditional SA methods applied to a reliability problem. Traditional SA methods provide variable ranking, whereas the proposed method provides an indication on the variation in the probability of failure given the variation of parameter δ . This is useful when the practitioner is interested on which configurations of the problem lead to an increase of the failure probability. This might also be used to assess the conservatism of a problem, if every variations of the input lead to decrease in the probability of failure. Additionally, it has three advantages:

- the ability for the user to set the most adapted constraints considering his/her problem/objective.
- The MC framework allowing to use previously done function calls, thus limiting the CPU cost of the SA, and allowing the user to test several perturbations.
- They are easy to interpret.

We argue that with an adapted perturbation, this method can fulfill the presented reliability engineer's objective (see Section 3.3.3 for further discussions on this topic). From this point of view, the DMBRSI are a good alternative to FORM/SORM's importance factors (as they can provide wrong results, see the Ishigami case) and to Sobol' indices (as they are costly and non-informative).

3.6.2 Equivalent perturbation

The question of "equivalent" perturbation arises from cases where all inputs are not identically distributed. Indeed, problems may emerge when some inputs are defined on infinite intervals and when other inputs are defined on finite intervals (such as uniform distributions). We have proposed three ways to deal with these problems:

- perform a mean perturbation relatively to the standard deviation, hence including the spread of the various inputs in their respective perturbation;
- perform a quantile shifting;
- perform a parameters perturbation.

²This is not the case in reality

3.6.3 Support perturbation

In most examples given throughout this chapter, the perturbations of the inputs left the support of those variables unperturbed. However, a support modification has been tested on the Ishigami case where the parameters defining the support have been perturbed. Yet, we stress that given the estimation method (reverse importance sampling), it is mandatory that the support of the perturbed density is included in the support of the original density. Thus one cannot perturb the inputs so that the perturbed support is wider than the original one.

3.6.4 Further work

Most of the further work will be devoted to adapting the estimator of the indices $S_{i\delta}$ in term of variance reduction and of number of function calls. Further work will be made with importance sampling methods (test the proposed estimators). The adaptation of estimators using subset simulation must also be done.

A perturbation based on an entropy constraint might also be proposed. The differential entropy of a distribution can be seen as a quantification of uncertainty (Auder *et al.* [6]). Thus an example of (non-linear) constraint on the entropy can be:

$$-\int f_{i\delta}(x) \log f_{i\delta}(x) dx = -\delta \int f_i(x) \log f_i(x) dx.$$

Yet further computations have to be made to obtain a tractable solution of the KL minimization problem under the above constraint.

Another avenue worth exploring would be to change the metrics/divergences. That would amount to change the D in equation 3.9 (choice was made to take KLD); and to take another distance than Hellinger's in the parameter perturbation context. This has to be tested.

3.6.5 Acknowledgements

We thank all the coauthors of the paper [63] for their help on this chapter. Part of this work has been backed by French National Research Agency (ANR) through COSINUS program (project COSTA BRAVA noANR-09-COSI-015). We also thank Dr. Daniel Busby (IFP EN) for several discussions and Emmanuel Remy (EDF R&D) for proofreading. We also thank two anonymous reviewers for their helpful comments. All the statistical parts of this work have been performed within the R environment, including the *mistral* and *sensitivity* packages.

Chapter 4

Application to the CWNR case

4.1 Introduction

This fourth chapter presents the application of some of the developed methods to the CWNR case. This numerical model has been presented in the outline of the thesis, page 24. Remind that this black-box model provided the initial motivation for this thesis.

The software interfacing is done using the Open TURNS [2] software that manages the probabilist part of the analysis. A wrapper calls the model when necessary. Concerning the sensitivity analysis part, post-processing of the data obtained is done using the R software.

In this thesis, focus has been set on SA methods that are separated from the sampling step (see Chapter 2), Chapter 3), thus the separation between the estimation of P_f and the sensitivity analysis. To estimate P_f , the failure probability, FORM (see Section 1.2.2.2) method and crude Monte-Carlo (see Section 1.2.1.1) have been used. Crude Monte-Carlo is considered to be the reference method in this chapter. Importance sampling (see Section 1.2.1.3) was available but was not used due to the lack of knowledge to set the importance densities. Subset simulation (see Section 1.2.3) was also available but was not used due to the fact that the Open TURNS module only provides an estimation for P_f and not the sampling points.

The sensitivity analysis part then focuses on three methods: first, importance factors (see Section 1.3.2.2) are derived from the FORM sampling. Then, random forests (see Section 2.2) are built on the MC sample and sensitivity measures are obtained. Finally, DMBRSI (see Chapter 3) are used. Several perturbations (mean, quantile and parameters) are proposed.

Sobol' indices (see Section 1.4) are not tested in this chapter due to the limited information provided and the high computational cost. $\delta_i^{SS}(A_k)$ indices (see Section 2.3) are not used in this chapter since a sampling scheme from subset simulation was not available.

This chapter is divided in three main sections, focusing respectively on random input of dimension 3 (Section 4.2), dimension 5 (Section 4.3) and dimension 7 (Section 4.4). Notice that the smaller the dimension of the input, the more penalizing the case (since non-probabilised variables are set to penalizing values). Thus the failure probability diminishes as the dimensionality grows. A final section (Section 4.5) concludes.

4.2 Three variables case

In this first section, three variables are probabilised. Table 1 is partially reproduced in Table 4.1 to indicate which distributions follow the variables. Table A.1 is a reminder of the inputs' densities.

Random var.	Distribution	Parameters
Thickness (m)	Uniform	$a = 0.0075, b = 0.009$
h (m)	Weibull	$a = 0.02, \text{scale} = 0.00309, \text{shape} = 1.8$
Ratio height/length	Lognormal	$a = 0.02, \ln(b) = -1.53, \ln(c) = 0.55$

Table 4.1: Distributions of the random physical variables of the CWNR model - 3 variables

4.2.1 Estimating P_f

4.2.1.1 Crude Monte-Carlo

A Crude Monte-Carlo (MC) estimation has been performed, with a sample of size 10000. 683 points were failing points thus the failure probability is estimated by:

$$\hat{P}_f = 0.0683.$$

This will be considered as the reference result. The sampling scheme will be used to build random forests (Section 4.2.2.1) and DMBRSI (Section 4.2.2.2).

4.2.1.2 FORM

FORM has been used. 52 function calls have been done. However the estimated failure probability is here of:

$$\hat{P}_{FORM} = 3.19 \times 10^{-16},$$

which is several orders of magnitude beneath the reference value. The results of FORM are not trustworthy in this case, therefore no sensitivity analysis will be performed with FORM in this case. Notice that the user is not warned that the FORM results are wrong. This is a major drawback of this technique.

4.2.2 Sensitivity Analysis

4.2.2.1 Random Forests

The methodology presented in Section 2.2 is used along this section. A forest of 500 trees is fitted on the MC sample. The reference value $\lfloor \sqrt{d} \rfloor$ is used as the number of variables randomly selected at each step. In this case, it means that 1 variable is selected as $d = 3$.

Variable	Thickness	h	Ratio
Index	0.01448048	0.10574811	0.02529668

Table 4.2: MDA index - 3 variables

MDA From Table 4.2, it can be inferred that the most influential variable is h , with 5 times as much influence as the secondly important variable, namely the ratio. Finally comes the thickness with an index twice as small as the one of the ratio. However from the numerical results of Section 2.2, it can be stated that the thickness has some influence on its own, according to the MDA indices.

Variable	Thickness	h	Ratio
Index	103.473	998.0205	169.5899

Table 4.3: Gini importance - 3 variables

Gini importance Table 4.3 assesses that the variable ranking is not modified when switching the measure. The index of h is more than 5 times higher than the one of the ratio, which is almost twice as large as the one of the thickness. However, due to the fact that here the Gini importance is used, it cannot be certain that the thickness has an influence on its own.

Model validation The confusion matrix (on the out-of-bag samples) of the forest is presented in Table 4.4.

		Observed		Class prediction error
		0	1	
Predicted	0	9299	18	0.001931952
	1	43	640	0.062957540

Table 4.4: Confusion matrix of the forest - 3 variables

It can be seen that the class prediction error is around 30 times bigger for the failing points than for the safe points. This is much less than in the tests of Section 2.2, but the model is still uneven.

4.2.2.2 DMBRSI

The methodology presented in Chapter 3 is used here. Due to the non-similarity of the distributions, a mean shift, a quantile shift and a parameter shift will be applied on this test case. It has been discussed further in Section 3.3.1.3, and the flood case (Section 3.4.7) might be used as an example.

Mean shifting First, the mean is shifted relatively to the standard deviation. Thus for any input, the original distribution is perturbed so that its mean is the original's one plus δ times its standard deviation, δ ranging from -1 to 1 with 40 points. The result is plot in Figure 4.1.

Figure 4.1 shows two tendencies. First the thickness and the ratio behave as follows: increasing the mean of these variables slightly decreases the failure probability whereas decreasing their mean slightly increases the failure probability. The effect is a little bit stronger for the thickness, but the confidence intervals are not well separated thus it is difficult to conclude with certainty on the relative influence of these two variables. On the other hand, increasing the mean of h increases the failure probability and decreasing the mean of h strongly decreases the failure probability. The effect is much stronger for h than it is for the two other variables.

Quantile shifting The first quantile to be perturbed is the extreme left-hand tail, namely the 5%-quantile. The result of such a perturbation for all the variables is plotted in Figure 4.2.

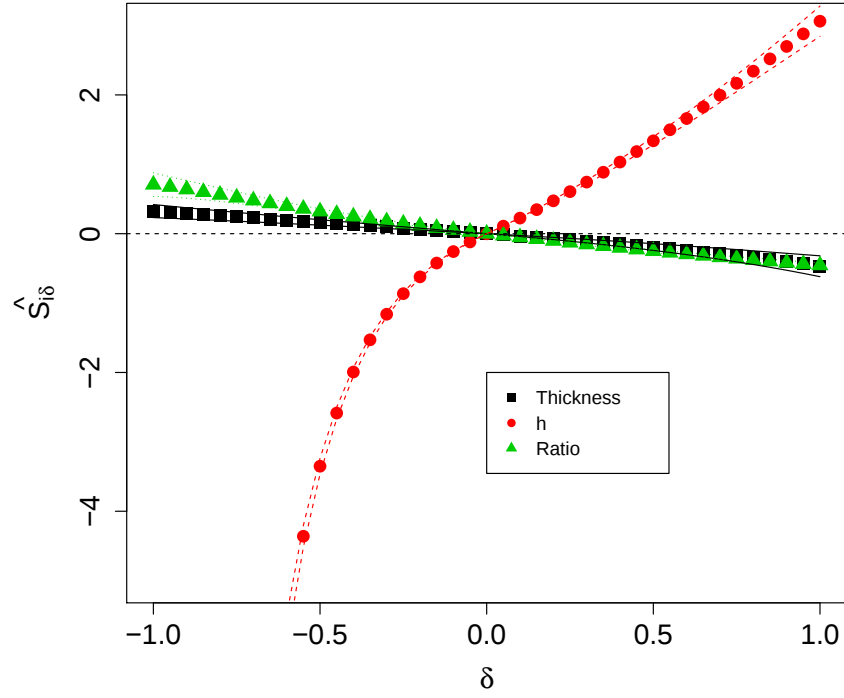


Figure 4.1: Estimated indices $\hat{S}_{i\delta}$ for the CWNR case with a mean perturbation - 3 variables

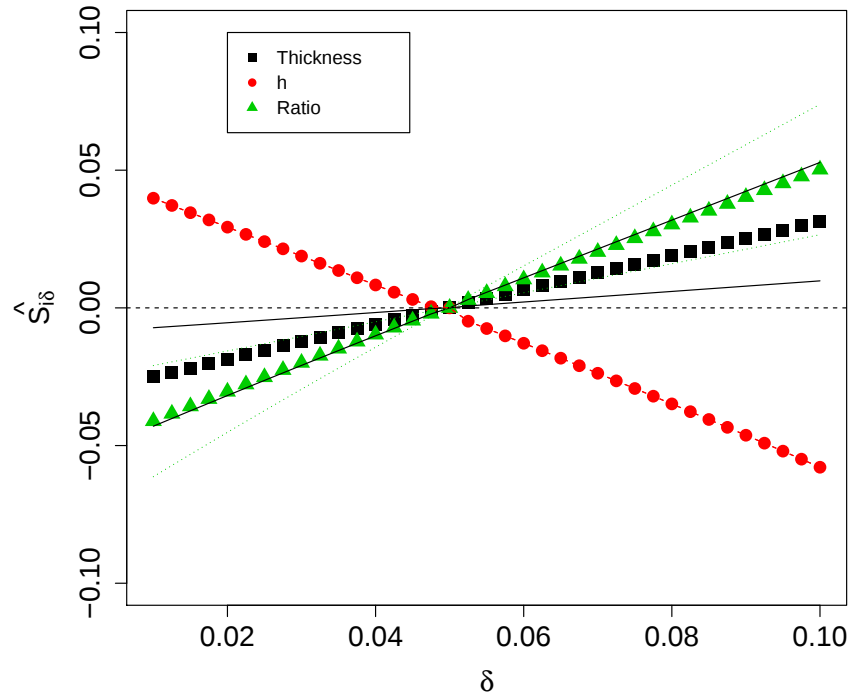


Figure 4.2: 5th percentile perturbation on the CWNR case - 3 variables

This graph shows that a quantile weight reduction for the thickness and the ratio diminishes the failure probability, whereas it increases the failure probability for h . The effect is reversed when

increasing the weight of the quantile. The influence is of the same order of magnitude for the three variables, with a slightly smaller influence for the ratio. However, the confidence intervals for the ratio and the thickness are not well separated.

The 1st quartile is then perturbed and the results are plotted in Figure 4.3.

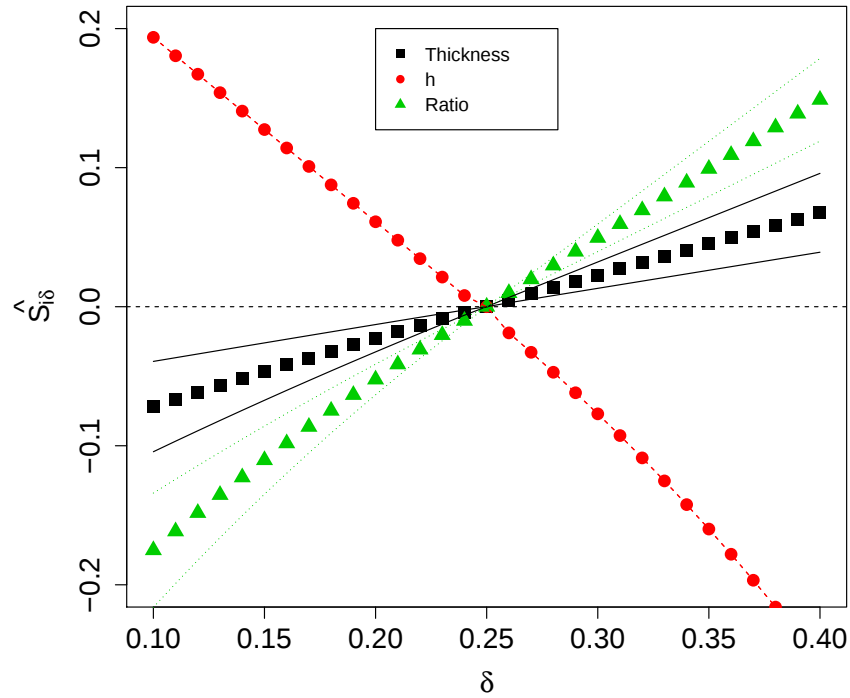


Figure 4.3: 1st quartile perturbation on the CWNR case - 3 variables

When perturbing less extreme values of the left-hand tail, the results are similar. In particular, the influences are of the same order of magnitude yet h has a larger influence than the thickness, which has a larger influence than the ratio. The confidence intervals are separated.

The median of the input distributions is then perturbed, the resulting indices are plotted in Figure 4.4.

When perturbing the median, tendencies are similar to the two previous graphs. The influence of h is larger than the influence of the other variables. The thickness has a larger influence than the ratio. Confidence intervals are well separated.

The third quartile is then perturbed and the indices are plotted in Figure 4.5.

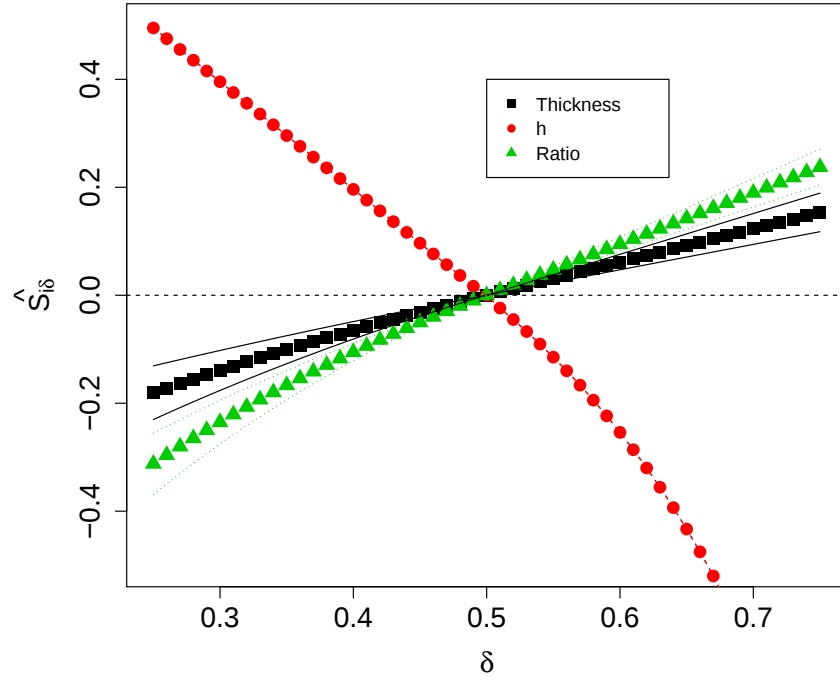


Figure 4.4: Median perturbation on the CWNr case - 3 variables

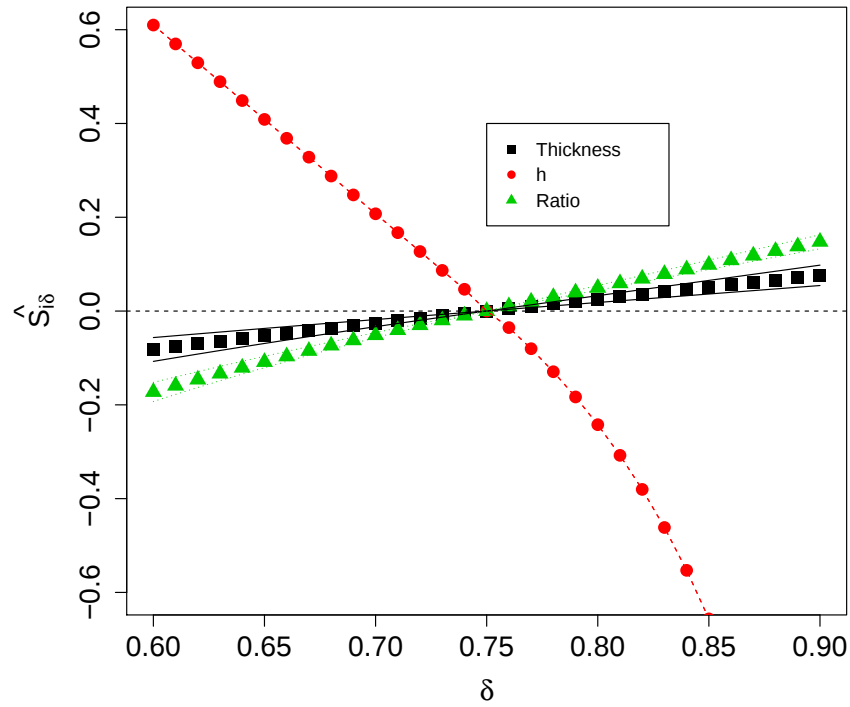


Figure 4.5: 3rd quartile perturbation on the CWNr case - 3 variables

Tendencies are similar to the three previous graphs. The influence of h is much larger than the influence of the thickness and of the ratio. Confidence intervals are well separated.

Finally, the extreme right-hand tail is perturbed, this comes to a perturbation on the 95th percentile. Results are plotted in Figure 4.6.

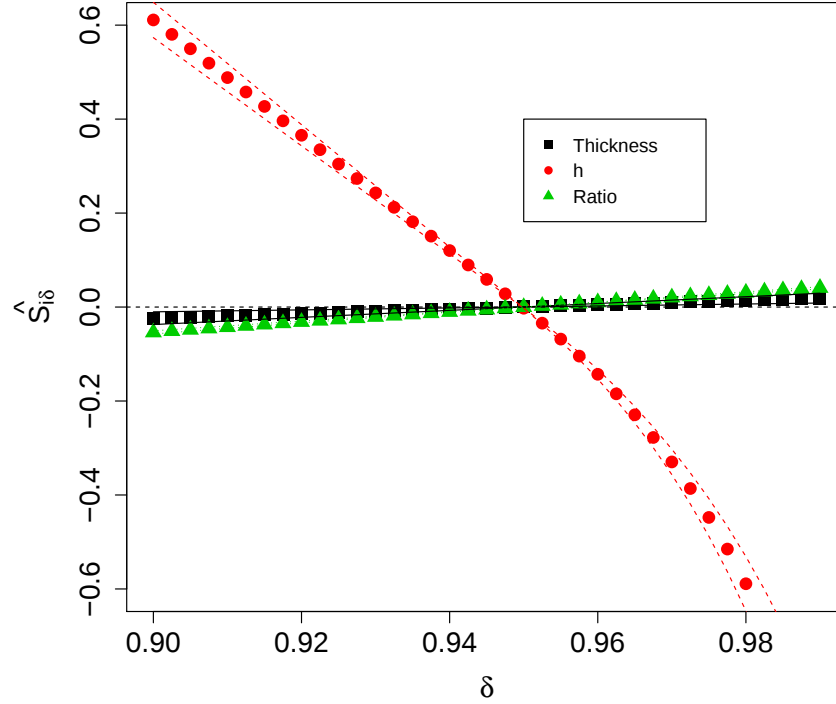


Figure 4.6: 95th percentile perturbation on the CWN case - 3 variables

The influence of h over the two other variables is tremendous. This variable is much more sensitive to a right-hand tail perturbation than the thickness and the ratio.

As a conclusion, the practitioner needs to be careful when modelling the right-hand tail of h . The left-hand tail of the three variables is equally important, but the indices are much smaller than for the right-hand tail. Additionally, the code seems to behave in a monotonic fashion.

Parameters shifting 6 parameters will be perturbed on this case:

- a minimum and a maximum for the thickness;
- a scale and a shape for h ;
- a mean of the logarithm (meanlog) and a standard deviation of the logarithm (sdlog) for the ratio.

These parameters are perturbed¹ and the estimated indices are plotted in function of the Hellinger distance in Figure 4.7 as explained in Figure 3.8. 95% confidence intervals are provided as well.

First, the two parameters driving the thickness bear a small influence with respect to the others. Diminishing the maximum of the thickness increases slightly the failure probability whereas increasing its minimum slightly diminishes the failure probability. Second, the scale of h has the largest influence over the model. Increasing it largely increases the failure probability whereas diminishes it diminishes in a tremendous way the failure probability. The confidence intervals get broader yet

¹notice that the minimum of the thickness is only increased and the maximum is decreased

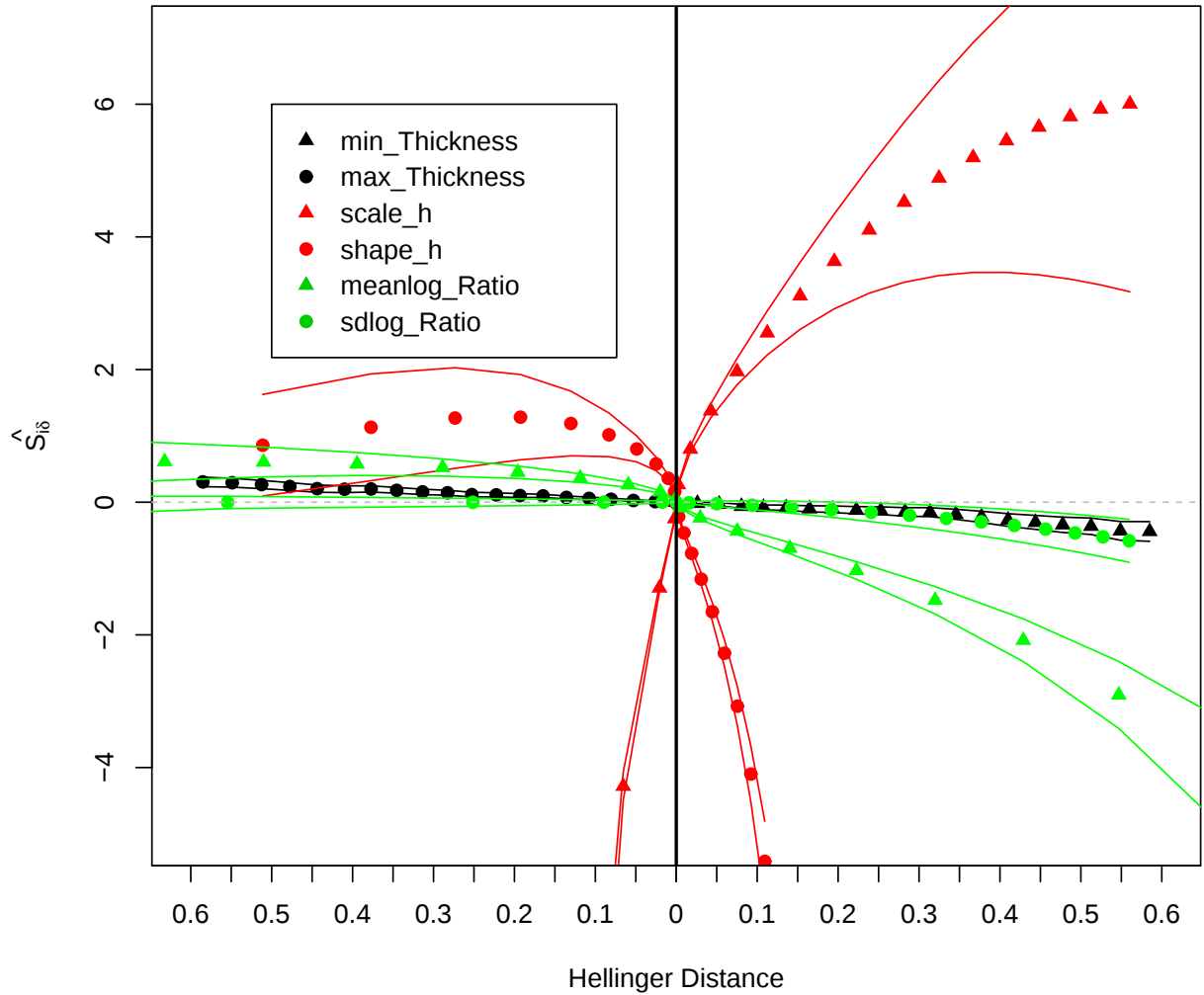


Figure 4.7: Parameters perturbation on the CWNr case - 3 variables

stay well separated from the others. Third, increasing the shape of h strongly diminishes the failure probability. Decreasing the shape of h increases the failure probability. The effect of this augmentation is not linear, as the growing tendency seems to vanish when decreasing strongly the shape. This is an interesting result. Then, diminishing the meanlog of the ratio increases slightly the failure probability whereas increasing it slightly diminishes the failure probability. Finally, the sdlog of the ratio behaves in a similar manner, yet with a smaller influence. The final ranking of the parameters in terms of influence is: the scale, the shape, the sdlog. Other parameters bear a quasi-null influence. These results are consistent with the ones provided by the mean and the quantile perturbation.

4.2.2.3 Conclusion

On the three variables CWNr case, the following can be concluded:

- The ranking provided by the forest is h , ratio then thickness.
- In terms of mean perturbation, the indices associated to h have a high (absolute) value whereas the ones associated to the two other variables are much smaller.

- The quantile perturbation has shown that the right-hand tail of h has the more impact on the failure probability. The left-hand tail of the three variables is equally important. Additionally, the code seems to behave in a monotonic fashion.
- The parameters perturbation has demonstrated that the model is mostly driven by the scale and the shape of h and by the sdlog of the ratio.

4.3 Five variables case

In this section, five variables are probabilised. Table 1 is partially reproduced in Table 4.5 to remind which distributions follows the variables. Table A.1 is a reminder of the inputs' densities.

Random var.	Distribution	Parameters
Thickness (m)	Uniform	$a = 0.0075, b = 0.009$
h (m)	Weibull	$a = 0.02, \text{scale} = 0.00309, \text{shape} = 1.8$
Ratio height/length	Lognormal	$a = 0.02, \ln(b) = -1.53, \ln(c) = 0.55$
Azimuth flaw ($^\circ$)	Uniform	$a = 0, b = 360$
Altitude (mm)	Uniform	$a = -5096, b = -1438$

Table 4.5: Distributions of the random physical variables of the CWNR model - 5 variables

4.3.1 Estimating P_f

4.3.1.1 Crude Monte-Carlo

A Crude Monte-Carlo (MC) estimation has been performed, with a sample of size 10^5 . Only 81 points were failing points thus the failure probability is estimated by:

$$\hat{P}_f = 0.00081.$$

This will be considered as the reference result. The sampling scheme will be used to build random forests (Section 4.3.2.1) and DMBRSI (Section 4.3.2.2).

4.3.1.2 FORM

FORM has been used. 106 function calls have been done. However the estimated failure probability is here of:

$$\hat{P}_{FORM} = 6.28 \times 10^{-2},$$

which is two orders of magnitude above the reference value (the failure probability is overestimated). The results of FORM are not trustworthy here, therefore no sensitivity analysis will be performed with FORM in this case.

4.3.2 Sensitivity Analysis

4.3.2.1 Random Forests

The methodology presented in Section 2.2 is used along this section. A forest of 500 trees is fitted on the MC sample. 2 variables are selected at each step of the tree building.

Variable	Thickness	h	Ratio	Azimuth	Altitude
Index	3.99×10^{-5}	5.16×10^{-4}	5.52×10^{-5}	3.91×10^{-4}	3.19×10^{-4}

Table 4.6: MDA index - 5 variables

MDA Indices are quite close to 0, as if no variable was influential. The variables with the strongest indices are h , the azimuth and the altitude.

Variable	Thickness	h	Ratio	Azimuth	Altitude
Index	19.79398	53.43655	22.38101	37.72627	28.28982

Table 4.7: Gini importance - 5 variables

Gini importance Indices are smaller than in the tests. The ranking provided is the following: h , azimuth, altitude, the ratio and the thickness.

Model validation The confusion matrix (on the out-of-bag samples) of the forest is presented in Table 4.8.

		Observed		Class prediction error
		0	1	
Predicted	0	99917	2	2×10^{-5}
	1	59	22	0.73

Table 4.8: Confusion matrix of the forest - 5 variables

It can be seen that the class prediction error for the failure points is above 0.7. The fitted model is then unusable. No conclusion should be drawn from this forest, therefore the rankings provided above are not to be considered. This lack of quality of the fitted model is a major drawback of the method.

4.3.2.2 DMBRSI

The methodology presented in Chapter 3 is used here. Due to the non-similarity of the distributions, a mean shift, a quantile shift and a parameter shift will be applied on this test case.

Mean shifting First, the mean is shifted relatively to the standard deviation. Thus for any input, the original distribution is perturbed so that its mean is the original's one plus δ times its standard deviation, δ ranging from -1 to 1 with 40 points. The result is plot in Figure 4.8.

Figure 4.8 shows three different behaviours. First the thickness and the ratio behave as is the three variables case: increasing the mean of these variables slightly decreases the failure probability whereas decreasing their mean slightly increases the failure probability. The effect is a little bit stronger for the thickness when increasing the mean, while it is a little bit stronger for the ratio when decreasing the mean. The confidence intervals are not well separated here. Then, increasing the mean of h increases the failure probability and decreasing the mean of h strongly decreases the failure probability. The behaviour is the same for the altitude with a smaller impact. Finally,

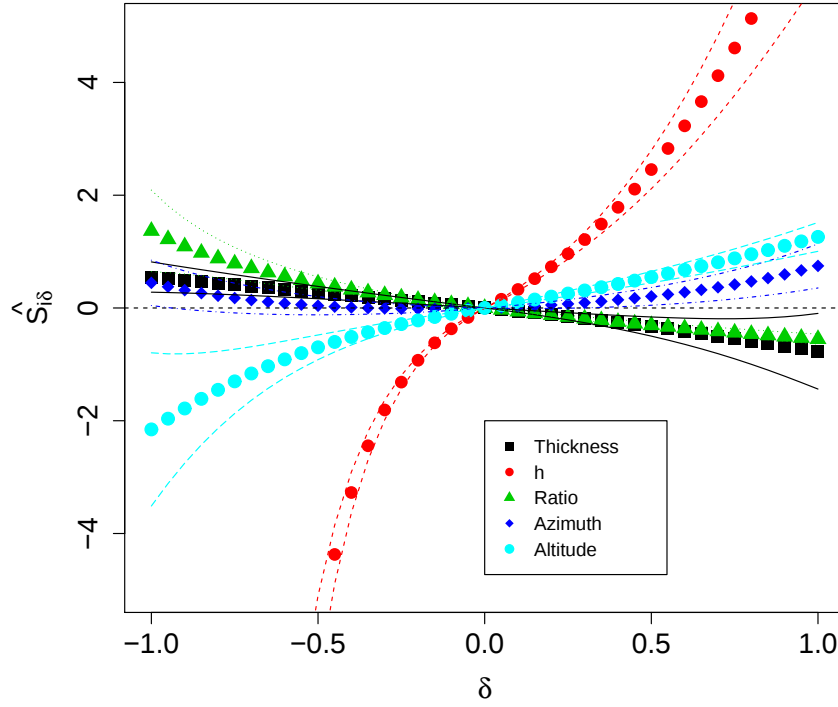


Figure 4.8: Estimated indices $\hat{S}_{i\delta}$ for the CWN case with a mean perturbation - 5 variables

increasing or decreasing the mean of the azimuth slightly increases the failure probability. The two more influential variables here are h and the altitude, yet it has to be noticed that h is of primary importance.

Quantile shifting The first quantile to be perturbed is the extreme left-hand tail, namely the 5%-quantile. The result of such a perturbation for all the variables is plotted in Figure 4.9.

This graph shows two opposite behaviours. First, decreasing the weight of the 5th percentile decreases the failure probability for the thickness, the ratio and the azimuth. For these variables, increasing the weight of the considered quantile increases the failure probability. Then, the behaviour is reversed for h and the altitude. Concerning the variable ranking, the azimuth has the more influence, while the altitude and h have the same small influence. The ratio has a larger influence than the thickness, but the confidence intervals are not well separated here.

The 1st quartile is then perturbed and the results are plotted in Figure 4.10.

When perturbing less extreme values of the left-hand tail, the behaviours are similar, but the order of influence is modified. In particular, the azimuth that was the most influential variable in Figure 4.9 is now the less influential. Then comes the thickness, and the three remaining variables have an equivalent influence.

The median of the input distributions is then perturbed, the resulting indices are plotted in Figure 4.11.

When perturbing the median, tendencies are similar to the two previous graphs for the thickness, h , the ratio and the altitude. However, the tendency is modified for the azimuth: increasing the weight of the median slightly decreases the failure probability whereas decreasing the weight increases

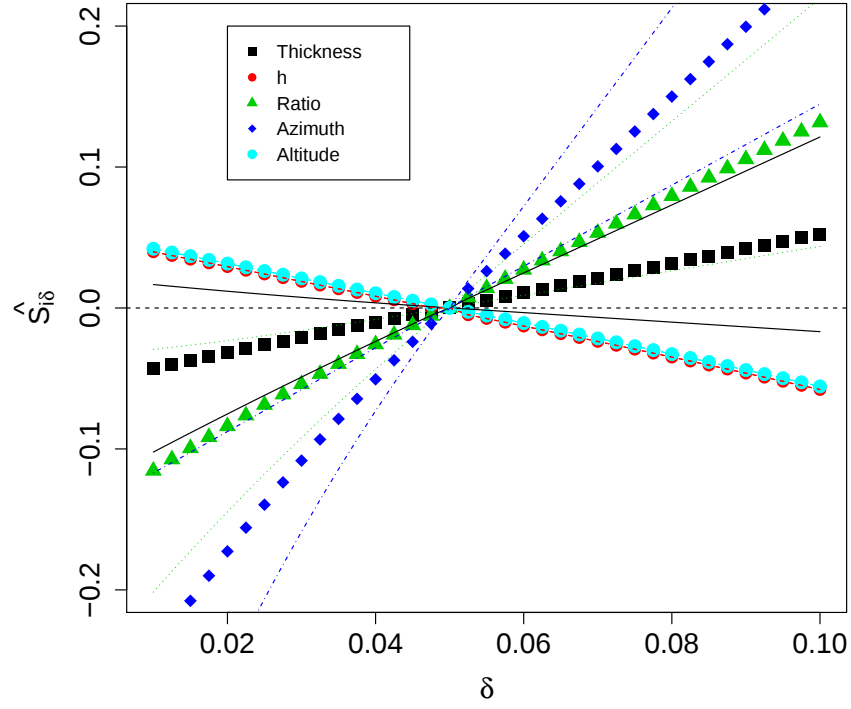


Figure 4.9: 5th percentile perturbation on the CWNR case - 5 variables

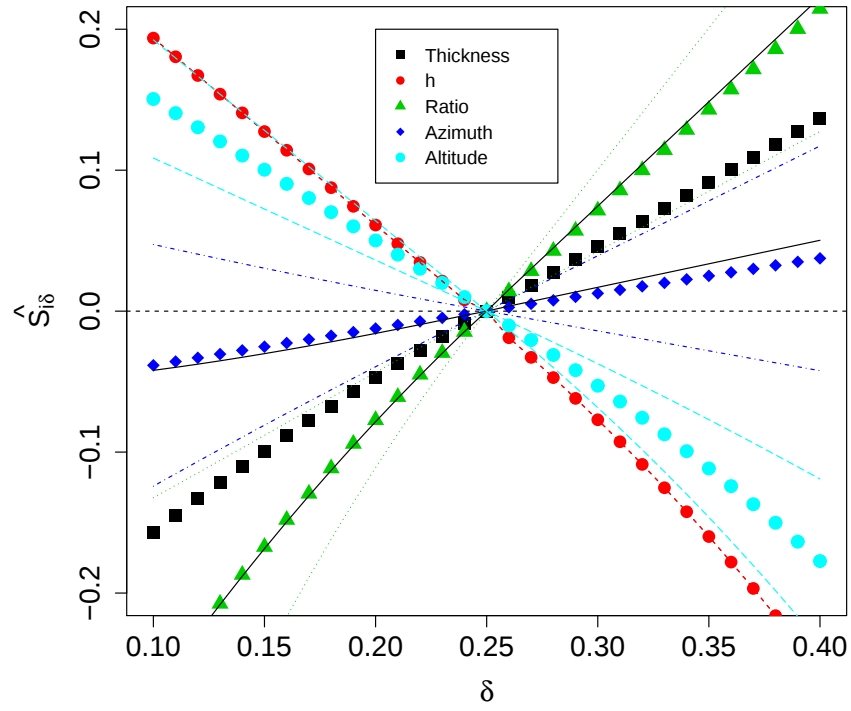


Figure 4.10: 1st quartile perturbation on the CWNR case - 5 variables

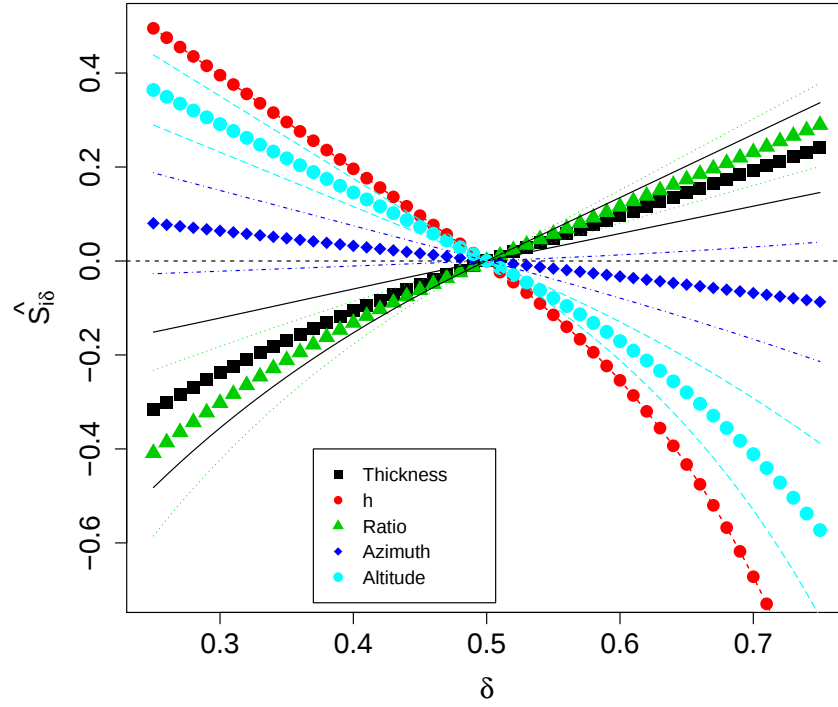


Figure 4.11: Median perturbation on the CWN case- 5 variables

the failure probability. The influence of h is the largest, then comes the altitude, followed by the ratio and the thickness. The azimuth has the smallest influence.

The third quartile is then perturbed and the indices are plotted in Figure 4.12.

Tendencies are similar to the previous graphs. The influence of h and of the altitude is larger than the one of the other variables. Confidence intervals are well separated except for the ratio and the thickness.

Finally, the extreme right-hand tail is perturbed, this comes to a perturbation on the 95th percentile. Results are plotted in Figure 4.13.

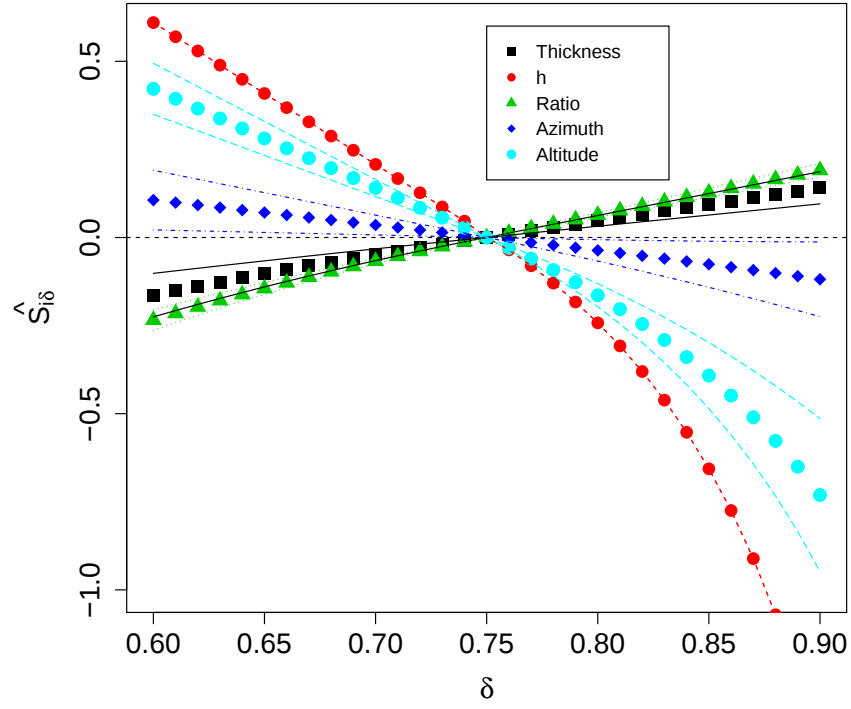


Figure 4.12: 3rd quartile perturbation on the CWNR case - 5 variables

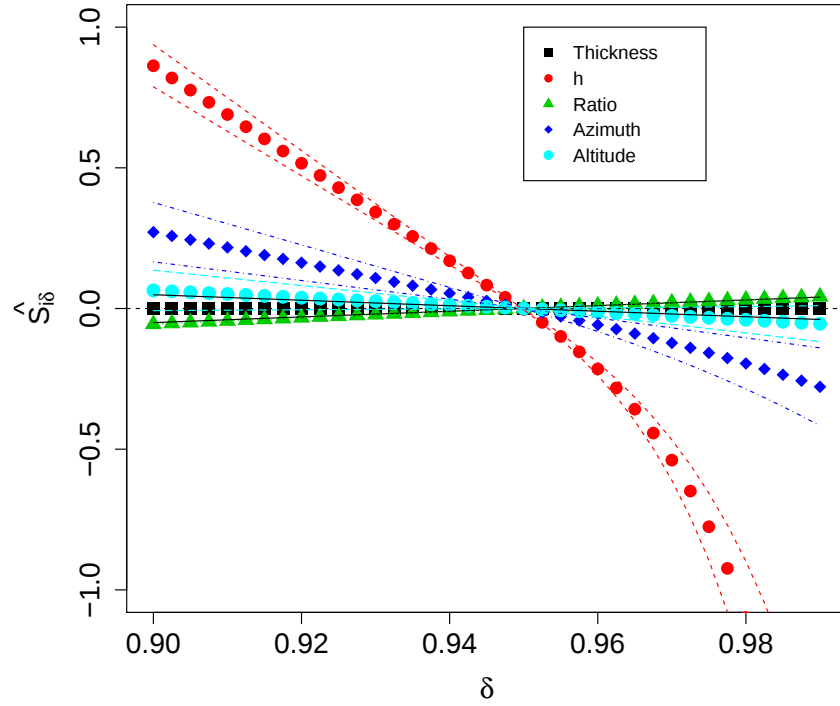


Figure 4.13: 95th percentile perturbation on the CWNR case - 5 variables

The influence of h over the other variables is tremendous. The azimuth is also more influential

than the three remaining variables.

As a conclusion, the practitioner needs to be careful when modelling the right-hand tail of h , and the tails of the azimuth. In terms of value of the indices, the right-hand tails have much more impact than the left-hand tails. Additionally, this analysis revealed the non-monotonic behaviour of the azimuth.

Parameters shifting 10 parameters will be perturbed on this case:

- a minimum and a maximum for the thickness;
- a scale and a shape for h ;
- a mean of the logarithm (meanlog) and a standard deviation of the logarithm (sdlog) for the ratio.
- a minimum and a maximum for the azimuth;
- a minimum and a maximum for the altitude;

These parameters are perturbed so that the support is not increased: the minimums are only increased and the maximums are decreased. The estimated indices are plotted in function of the Hellinger distance in Figure 4.14 as explained in Figure 3.8. 95% confidence intervals are provided as well.

Due to the large number of parameters perturbed, the image is difficult to read. However, the influence of the parameters driving h (plotted in red) is tremendous. The indices associated to the scale are larger than the ones associated to the shape, however the width of the confidence intervals grows quite large, thus it is difficult to conclude on these two parameters. Then, the maximum of the altitude seems to have the most influence over the failure probability. Diminishing the maximum of the altitude leads to a decrease of the failure probability. It is followed by the meanlog of the ratio. The indices associated with other parameters are too noisy and stacked around 0.

4.3.2.3 Conclusion

On the five variables CWNRR case, the following can be concluded:

- The ranking provided by the forest is not to be considered as the model is badly fitted.
- In terms of mean perturbation, the indices associated to h have the highest (absolute) value. Then comes the altitude, followed by the ratio, the azimuth and the thickness. Notice that the relative influence of the ratio, the azimuth and the thickness is hardly separable.
- The quantile perturbation has shown that the right-hand tail of h , and the tails of the azimuth are more influential than the tails of others variable. The right-hand tails have much more impact than the left-hand tails though. Additionally, this analysis revealed the non-monotonic behaviour of the azimuth.
- The parameters perturbation has demonstrated that the model is mostly driven by the scale and the shape of h . Then, the maximum of the altitude seems to have the most influence over the failure probability, followed by the meanlog of the ratio.

It is noticeable that the ranking differs from the three variables case, yet the dimension of the flaw h is still the most influential variable. Additionally, it seems interesting to notice that the altitude is influential, but mostly in the right-hand tail (see Figure 4.12).

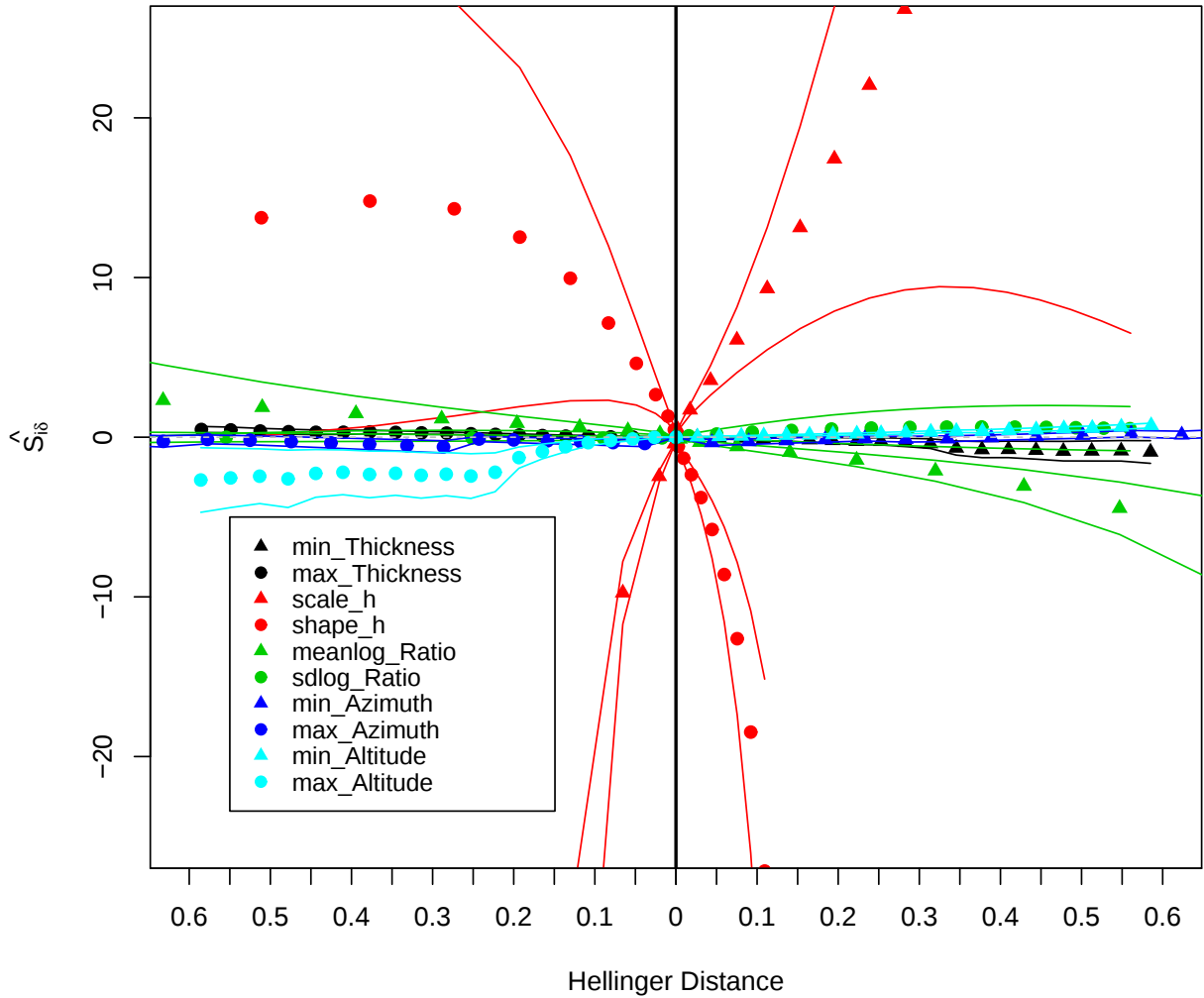


Figure 4.14: Parameters perturbation on the CNWR case - 5 variables

4.4 Seven variables case

In this section, seven variables are probabilised. Table 1 is reproduced in Table 4.9 to remind which distributions follows the variables. Table A.1 is a reminder of the inputs' densities.

Random var.	Distribution	Parameters
Thickness (m)	Uniform	$a = 0.0075, b = 0.009$
h (m)	Weibull	$a = 0.02, \text{scale} = 0.00309, \text{shape} = 1.8$
Ratio height/length	Lognormal	$a = 0.02, \ln(b) = -1.53, \ln(c) = 0.55$
Azimuth flaw ($^{\circ}$)	Uniform	$a = 0, b = 360$
Altitude (mm)	Uniform	$a = -5096, b = -1438$
$\sigma\Delta TT$	Gaussian	$\mu = 0, \sigma = 1$
σRes	Gaussian	$\mu = 0, \sigma = 1$

Table 4.9: Distributions of the random physical variables of the CNWR model - 7 variables

4.4.1 Estimating P_f

4.4.1.1 Crude Monte-Carlo

A Crude Monte-Carlo (MC) estimation has been performed, with a sample of size 7×10^6 . Notice that this samples took several weeks to be computed. 468 points were failing points thus the failure probability is estimated by:

$$\hat{P}_f = 6.68 \times 10^{-5}.$$

This will be considered as the reference result.

4.4.1.2 FORM

FORM has been used. 183 function calls have been done. The estimated failure probability is here of:

$$\hat{P}_{FORM} = 4.23 \times 10^{-7},$$

which is two orders of magnitude under the reference value (the failure probability is underestimated). The results of FORM are not trustworthy in this case.

4.4.2 Sensitivity Analysis

4.4.2.1 Random Forests

The methodology presented in Section 2.2 is used along this section. We tried to fit a forest of 500 trees on the MC sample which dimension was 7×7000000 , with 2 variables selected at each step of the tree building. However the fitting step failed due to the size of the sample (as in Section 2.2.5.2, paragraph "increasing the sample size").

4.4.2.2 DMBRSI

The methodology presented in Chapter 3 is used here. Due to the non-similarity of the distributions, a mean shift, a quantile shift and a parameter shift will be applied on this test case.

Mean shifting First, the mean is shifted relatively to the standard deviation. Thus for any input, the original distribution is perturbed so that its mean is the original's one plus δ times its standard deviation, δ ranging from -1 to 1 with 40 points. The result is plot in Figure 4.15.

Three different behaviours can be observed. When increasing the mean of h , of the altitude and of $\sigma\Delta TT$ it increases the failure probability while when decreasing their means it decreases the failure probability. The effect is reversed for the thickness, the ratio and σRes . Finally, increasing or decreasing the mean of the azimuth slightly increases the failure probability. In terms of amplitude, three variables differentiate themselves from the others: h , $\sigma\Delta TT$ and σRes . Others variables have a smaller influence and their confidence intervals contains 0.

Quantile shifting The first quantile to be perturbed is the extreme left-hand tail, namely the 5%-quantile. The result of such a perturbation for all the variables is plotted in Figure 4.16.

It first should be notices that the indices for h and for $\sigma\Delta TT$ coincide. This graph shows two opposite behaviours. First, decreasing the weight of the 5th percentile decreases the failure probability for the thickness, the ratio, the azimuth and for σRes . For these variables, increasing the weight of the considered quantile increases the failure probability. Then, the behaviour is reversed for h , the altitude and $\sigma\Delta TT$. Concerning the variable ranking, σRes has the more influence. Then

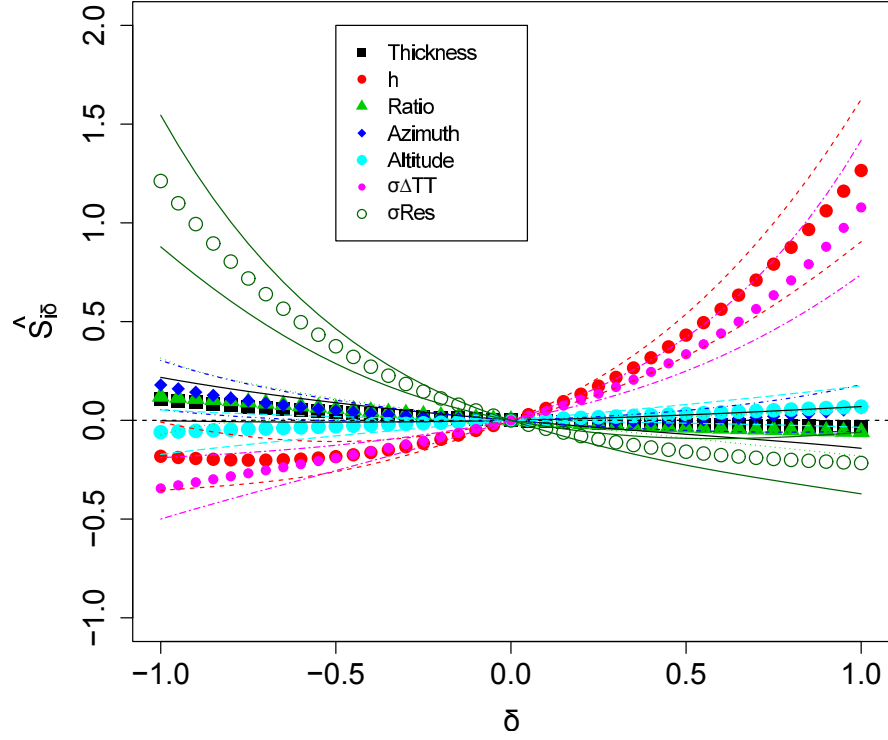


Figure 4.15: Estimated indices $\widehat{S}_{i\delta}$ for the CWNr case with a mean perturbation - 7 variables

comes the azimuth that has a medium influence, while the rest of the variables have the same small influence.

The 1st quartile is then perturbed and the results are plotted in Figure 4.17.

The indices for h and for the altitude coincide. When perturbing less extreme values of the left-hand tail, the behaviour are similar, but the order of influence is modified. The azimuth that was an influential variable in Figure 4.16 is now the less influential. The two most influential variables are $\sigma\Delta TT$ and σRes .

The median of the input distributions is then perturbed, the resulting indices are plotted in Figure 4.18.

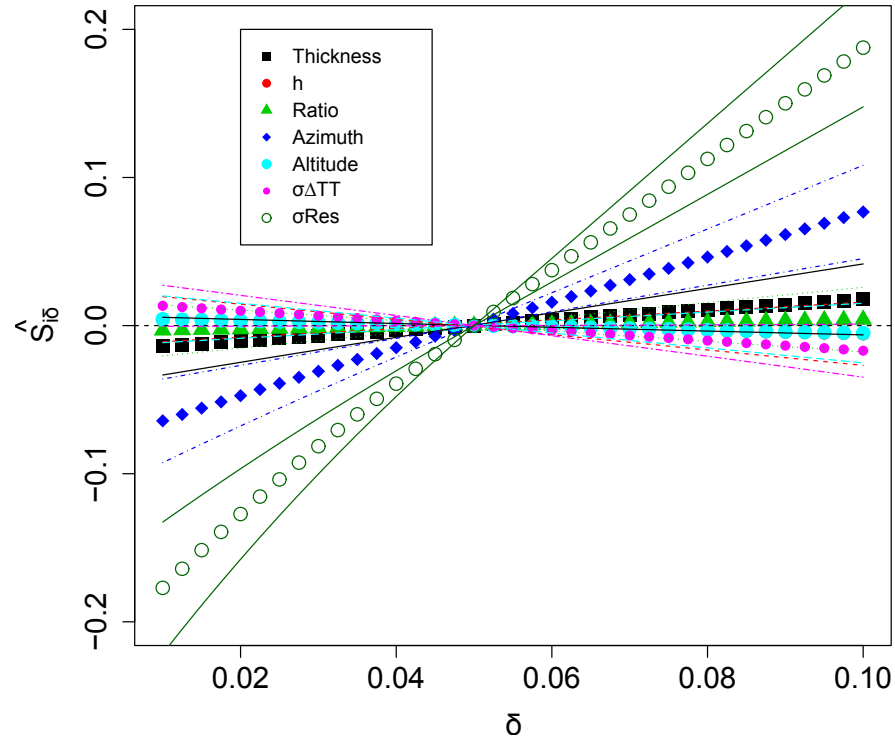
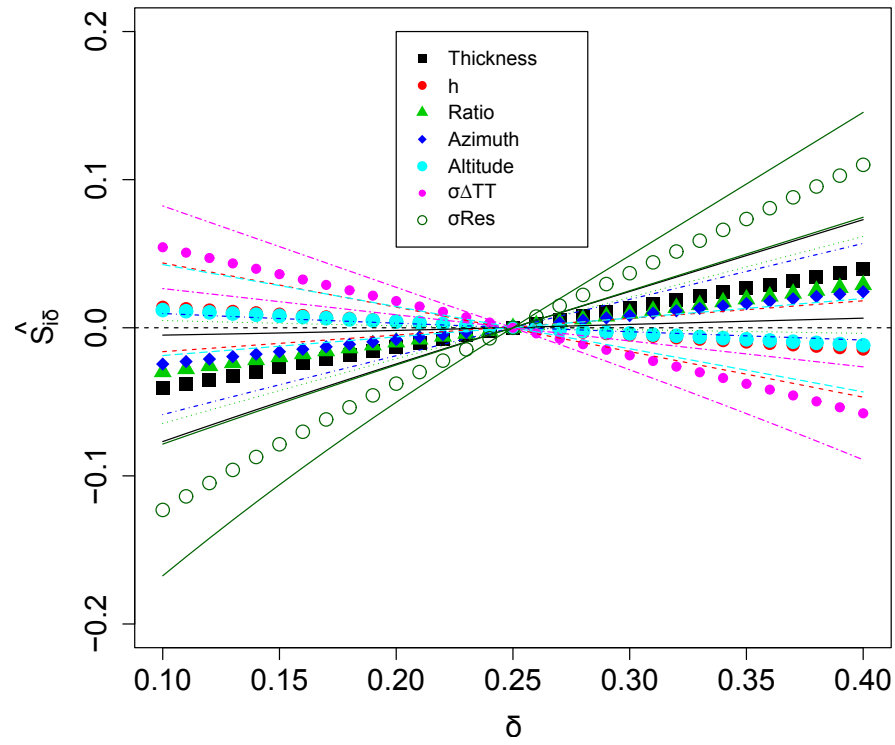
When perturbing the median, tendencies are similar to the two previous graphs for all the variables but the azimuth. Indeed increasing or decreasing the weight of the median for this variable does not impact the failure probability. The influence of $\sigma\Delta TT$ is the largest, followed by h and σRes that have a similar impact (but a different behaviour). Then comes the ratio and the thickness. The two other variables have a small to null impact.

The third quartile is then perturbed and the indices are plotted in Figure 4.19.

Tendencies are similar to the previous graphs except for the altitude and the thickness. The influence of h and of $\sigma\Delta TT$ is larger than the one of the other variables. The impact of the ratio and of σRes is similar.

Finally, the extreme right-hand tail is perturbed, this comes to a perturbation on the 95th percentile. Results are plotted in Figure 4.20.

This figure shows clearly the impact of the following variables (for which increasing the weight of the quantile decreases the failure probability), ordered by influence: h , $\sigma\Delta TT$, the azimuth. The others variables have a small to null impact.

Figure 4.16: 5th percentile perturbation on the CWN case - 7 variablesFigure 4.17: 1st quartile perturbation on the CWN case - 7 variables

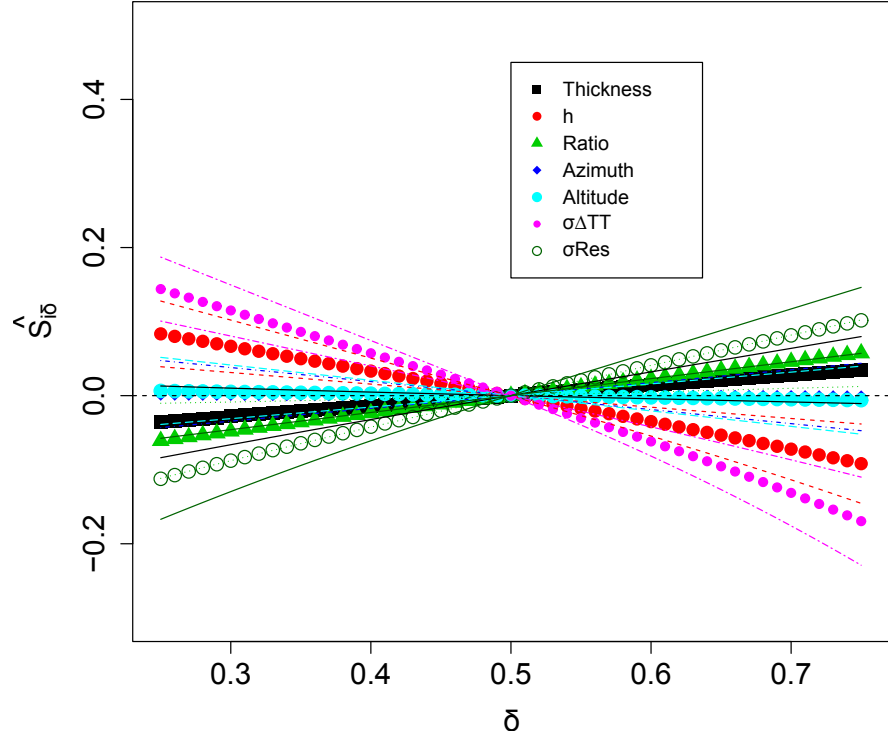


Figure 4.18: Median perturbation on the CWNr case- 7 variables

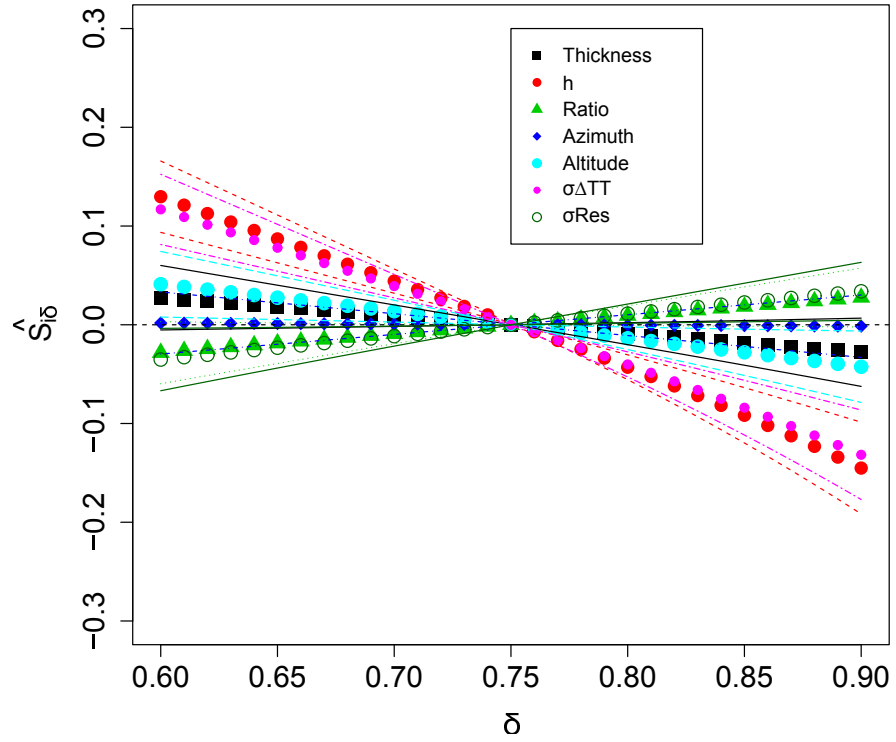
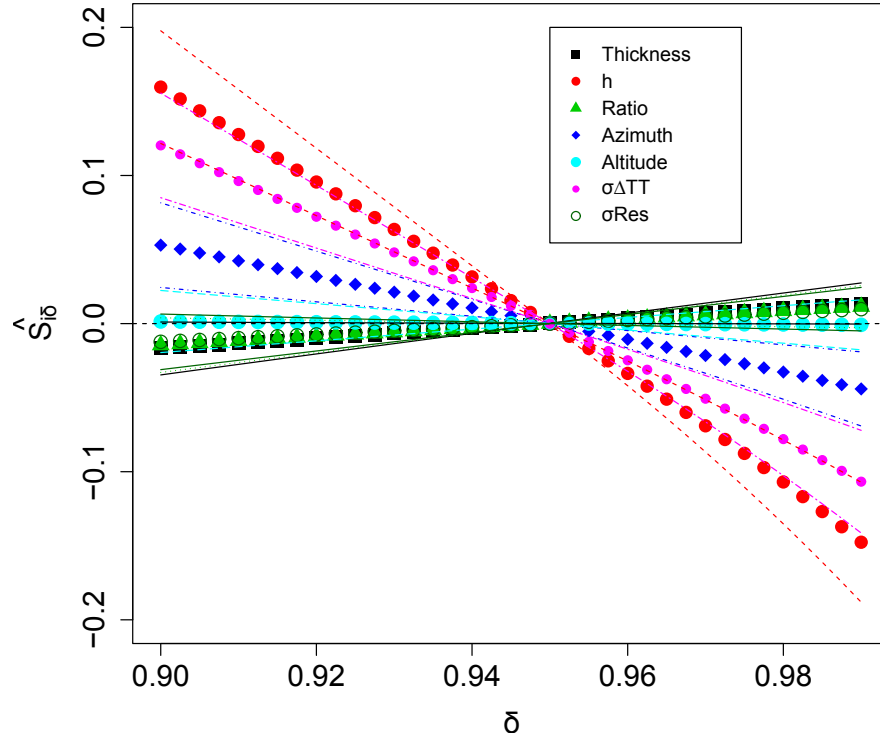


Figure 4.19: 3rd quartile perturbation on the CWNr case - 7 variables

Figure 4.20: 95th percentile perturbation on the CWN case - 7 variables

As a conclusion, the practitioner needs to be careful when modelling the right-hand tail of h and $\sigma\Delta TT$ altogether with the left-hand tail of σRes . The tails of the azimuth need caution too. Additionally, this analysis revealed the non-monotonic behaviour of the azimuth for the 7 variables case.

Parameters shifting 14 parameters will be perturbed on this case:

- a minimum and a maximum for the thickness;
- a scale and a shape for h ;
- a mean of the logarithm (meanlog) and a standard deviation of the logarithm (sdlog) for the ratio.
- a minimum and a maximum for the azimuth;
- a minimum and a maximum for the altitude;
- a mean and a standard deviation for $\sigma\Delta TT$;
- a mean and a standard deviation for σRes .

These parameters are perturbed so that the support is not increased: the minimums are only increased and the maximums are decreased. The estimated indices are plotted in function of the Hellinger distance in Figure 4.21 as explained in Figure 3.8. 95% confidence intervals are provided as well.

Due to the large number of parameters perturbed, the image is very difficult to read.

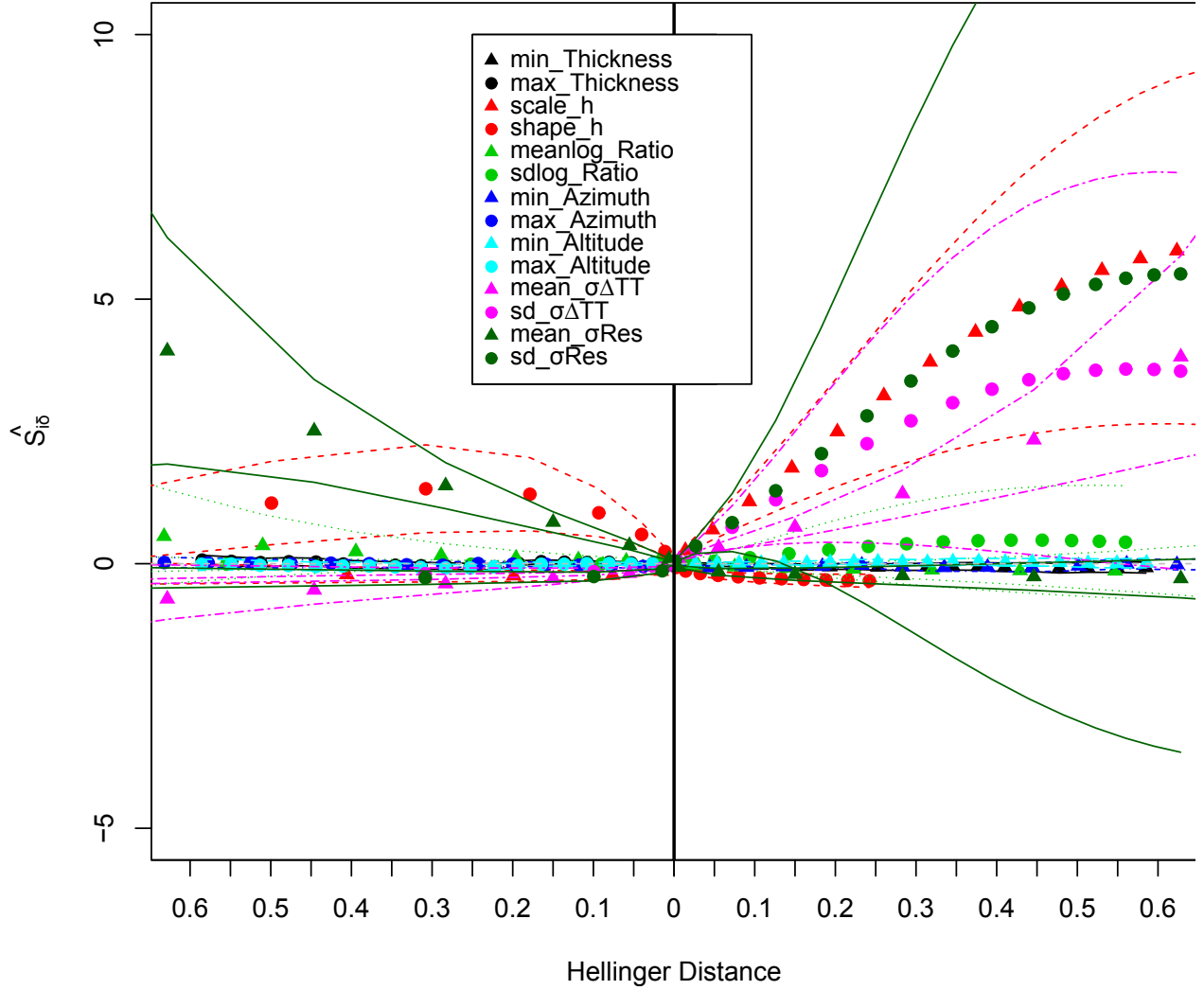


Figure 4.21: Parameters perturbation on the CNWR case - 7 variables

However, the influence of the parameters driving h (plotted in red), of $\sigma\Delta TT$ and of σRes is tremendous. The impact of a parameter perturbation for other variables is much smaller. In particular for h , increasing the scale or decreasing the shape increases the failure probability. Concerning $\sigma\Delta TT$, increasing the mean and the standard deviation increases the failure probability while decreasing these parameters has a much smaller impact. Finally when shifting the parameters of σRes it can be seen that decreasing the mean or increasing the standard deviation strongly increases the failure probability. However the width of the confidence intervals grows quite large.

4.4.2.3 Conclusion

On the seven variables CWNr case, the following can be concluded:

- The forest model could not be fitted due to the size of the sample.
- In terms of mean perturbation, the indices associated to h , $\sigma\Delta TT$ and σRes have the highest (absolute) value.

- The quantile perturbation has shown that the right-hand tail of h and $\sigma\Delta TT$ altogether with the left-hand tail of σRes and the tails of the azimuth are more influential than the tails of others variable. Additionally, this analysis revealed the non-monotonic behaviour of the azimuth.
- The parameters perturbation has demonstrated that the model is mostly driven by the parameters of h of $\sigma\Delta TT$ and of σRes . This confirms the conclusion of the mean perturbation.

It is noticeable that the ranking differs from the three and the five variables case. However the dimension of the flaw h is still an influential variable.

4.5 Conclusion

Concerning the P_f estimation part, the MC method is still the reference method on an industrial code. The major drawback is of course the computational time needed. FORM is wrong in all the cases and should not be used.

Concerning the sensitivity analysis part, the random forest technique provides questionable results, since the fitted models are uneven or bad. This method is inconclusive at the moment.

DMBRSI seems an adapted method to perform sensitivity analysis on a failure event. Several tunings for several problems have been tested. However, if a single graph had to be provided to decision makers, we would present the mean perturbation one, as it carries most of the information.

In all the configuration studied, h is a priority variable. This is also the case for $\sigma\Delta TT$ and σRes in the 7 variables case.

The improvement perspectives of this study are:

- to combine subset simulation with the DMBRSI. To do so, an implementation of subset simulation that provides the sampling scheme must be performed;
- to combine importance sampling with the DMBRSI, now that the priority zones are known.

Conclusion

Summary and contributions

This thesis' first objective was to perform a sensitivity analysis on a black-box model, the CWNR case. Because the quantity of interest is a (small) failure probability, appropriate methods had to be used. Thus this thesis focused on two fields: structural reliability in one hand, and sensitivity analysis on the other hand.

First step was a bibliographical chapter (Chapter 1). This chapter aimed at clarifying the main existing techniques to estimate a failure probability (Section 1.2) and the main sensitivity analysis methods (Section 1.3). Then one of the most used sensitivity analysis technique (Sobol' indices) was tested on reliability toy-cases (Section 1.4). Sobol' indices applied to a failure indicator have highlighted a capacity to distinguish the non-influential from the influential variables. However, tests have shown that the following configuration -low first-order indices, high total order indices- is often present. Therefore the information provided by such indices is limited and may only confirm that all the variables interact to cause the failure event. Next, a moment-independent method (Borgonovo's δ_i indices) was tested on reliability toy-cases (Section 1.5). However, the produced indices were rather small with a positive bias in the estimations. The conclusion is that moment independent techniques are not adapted within the reliability context. A synthesis of the tested methods was proposed in Section 1.6. Finally, a discussion on the meaning and objectives of sensitivity analysis when dealing with failure probabilities, that we argue might be of use for the practitioner, was conducted in Section 1.7.

The conclusion of this bibliographical chapter is that there is a need for new sensitivity analysis methods in the reliability context. The next two chapters aimed at reaching this objective.

The second chapter focused on sensitivity analysis techniques with a variable ranking objective. Two sensitivity analysis methods were presented, thought as by-products of two sampling techniques (Monte-Carlo and subset simulation). The first part of the chapter (Section 2.2) was devoted to importance measures derived from random forests.

Reminders on specific binary classifiers (random trees) were proposed altogether with a review on stabilisation methods, including random forests. The importance measures (Gini importance and Mean Decrease Accuracy importance) were elicited. Then a bibliographical step was performed on the "sensitivity analysis using random forests" theme. Then the importance measures have been tested on reliability toy-cases. The conclusions were that the Mean Decrease Accuracy importance indices seemed more adapted since the Gini importance indices could affect a non-null importance to a non-influential variable. However, it must be stressed that the fitted models' quality is not satisfying. Indeed, from the imbalance of the classes in the original sample, there is a tendency in getting "weak" predictors that make much more prediction error on the minority class. This is a problem when drawing conclusions on sensitivity analysis with these types of models. The second

part of the chapter (Section 2.3) proposed a new sensitivity measure based upon the departure, at each step of a subset method, between each input original density and the density given the subset reached. Several tunings of the departure can be used. However this sensitivity analysis method gives a similar information that the one provided by the Sobol' indices on the failure indicator.

The third chapter presented an original sensitivity analysis method, called Density Modification Based Reliability Sensitivity Indices (DMBRSI). This sensitivity index is based upon input density modification, and is adapted to failure probabilities. The proposed indices reflect the impact of an input density modification on the failure probability. One needs to differentiate the proposed index and the perturbations. The indices are independent of the perturbation in the sense that the practitioner can set the perturbation adapted to his/her problem. The sensitivity index can be computed using the sole set of simulations that has already been used to estimate the failure probability, thus limiting the number of calls to the numerical model.

First, the indices and their theoretical properties have been presented in Section 3.2, altogether with the estimation methodology. For the sake of simplicity, a Monte-Carlo sampling scheme was considered. Second, Section 3.3 dealt with several perturbation methodologies. These perturbations can be classified into two main families: Kullback-Leibler minimization methods and parameter perturbations methods. The behaviour of the indices was examined in Section 3.4 through numerical simulations. In Section 3.5, it was proposed to improve the DMBRSI estimation with importance sampling and with subset simulation.

This chapter presented an original method designed for failure probabilities. One of the main advantage is the possibility to modify the perturbation applied without new calls to the model. However a major drawback persists: when there are too many parameters to perturb, the results may be complicated to interpret.

The fourth chapter presented the application of some of the developed methods to the CWNR case. Remind that this black-box model provided the initial motivation for this thesis.

To estimate P_f , two methods were used: crude Monte-Carlo and FORM. It appeared that FORM was wrong in every case, thus Monte-Carlo stays the reference method.

The sensitivity analysis part then focused on two methods: random forests (Chapter 2), and DMBRSI (Chapter 3). Sobol' indices (see Section 1.4) were not tested in this chapter due to the limited information provided and their high computational cost. $\delta_i^{SS}(A_k)$ indices (see Section 2.3) were not used either since a sampling scheme from subset simulation was not available.

This chapter is divided in three main sections, focusing respectively on random input of dimension 3, dimension 5 and dimension 7. Notice that the smaller the dimension of the input, the more penalizing the case (since non-probabilised variables are set to penalizing values). Thus the failure probability diminishes as the dimensionality grows.

DMBRSI appeared as an adapted method to perform sensitivity analysis on a failure event. In all the configurations studied, h (the dimension of the flaw) is a priority variable. This is also the case for $\sigma\Delta TT$ and σRes in the 7 variables case.

Future avenues for research and application

The methods presented in Chapter 2 can be improved. Specifically, there is a need to improve the binary classifiers (random forests). The MDA indices when using subset simulation must be implemented. Another perspective of improvement, when using the $\delta_i^{SS}(A_k)$ indices, is to conduct a work

including the copula theory.

The DMBRSI introduced in Chapter 3 have several ways of improvement. Most of the further work will be devoted to adapting the estimator of the indices $S_{i\delta}$ in terms of variance reduction and of number of function calls. The adaptation of estimators using subset simulation must also be done. A perturbation based on an entropy constraint might also be proposed. Yet further computations have to be made to obtain a tractable solution of the KL minimization problem. Another avenue worth exploring would be to change the metrics/divergences. That would amount to change the D in equation 3.9 (choice was made to take KLD); and to take another distance than Hellinger's in the parameter perturbation context. Another avenue might be the introduction of a structural dependency between the marginals of the input vector, and to perturb this dependency *via* the copula theory.

Further work can be done in Chapter 4. The main improvement perspectives of this study is to use subset simulation, to improve the estimation of P_f and to reduce the computational time. A coupling with the random forests *via* adapted MDA indices might be of interest as well. This could also allow the use of the c.d.f. departure measures $\delta_i^{SS}(A_k)$. Still to reduce the variance of the estimators, importance sampling must be tested.

Broader perspectives have to be considered. In particular, the use of sequential methods coupled with meta-models (Bect *et al.* [9]) is to be tested.

Recently, Fort *et al.* [35] introduced a new sensitivity index as a generalisation of Sobol' indices. They propose an adapted contrast function for each statistical purpose. It is interesting to notice that the contrast adapted to a threshold exceed is presented. This index then has to be tested and compared with DMBRSI in further work.

Communications

Publications

- P. Lemaître, E. Sergienko, A. Arnaud, N. Bousquet, F. Gamboa, and B. Iooss. Density modification based reliability sensitivity analysis. *Journal of Statistical Computation and Simulation*, In press, 2014
- E. Sergienko, P. Lemaître, A. Arnaud, D. Busby and F. Gamboa. Reliability sensitivity analysis based on probability distribution perturbation with application to CO2 storage. *Accepted with minor reviews in Structural Safety*, 2014

Oral presentations

- P. Lemaître and A. Arnaud. Hiérarchisation des sources d'incertitudes vis à vis d'une probabilité de dépassement de seuil - Une méthode basée sur la pondération des lois. In *Proceedings des 43 èmes Journées de Statistique*, Tunis, Tunisia, June 2011.
- P. Lemaître. Analyse de sensibilité pour des probabilités de dépassement de seuil. In *Proceedings of GdR MASCOT NUM*, Bruyères-le-Châtel, France, March 2012.
- P. Lemaître, E. Sergienko, F. Gamboa and B. Iooss. A global sensitivity analysis method for reliability based upon density modification. In *Proceedings of SIAM Conference on Uncertainty Quantification*, Raleigh, North Carolina USA, April 2012.
- A.L. Popelin, A. Dutfoy and P. Lemaître. Open TURNS: Open source Treatment of Uncertainty, Risk 'N Statistics. In *Proceedings of SIAM Conference on Uncertainty Quantification*, Raleigh, North Carolina USA, April 2012.
- P. Lemaître, A. Arnaud and B. Iooss. Sensitivity analysis methods for a failure probability. In *Proceedings of Lambda Mu 18*, Tours, France, October 2012.

Poster

- P. Lemaître, E. Sergienko, A. Arnaud, N. Bousquet, F. Gamboa, and B. Iooss. Sensitivity analysis method for failure probability. *Poster presented at SAMO 2013*, Nice, France, June 2013.

Software developments

- P. Lemaître. Density modification based reliability sensitivity indices (DMBRSI Function). *R Sensitivity Package*: <http://cran.r-project.org/web/packages/sensitivity/>

Bibliography

- [1] T. Abdo and R. Rackwitz. A new beta-point algorithm for large time-invariant and time-variant reliability problems. *Reliability and Optimization of Structural Systems*, 90:1–11, 1990.
- [2] G. Andrianov, S. Burriel, S. Cambier, A. Dutfoy, I. Dutka-Malen, E. De Rocquigny, B. Sudret, P. Benjamin, R. Lebrun, F. Mangeant, et al. Open TURNS, an open source initiative to treat uncertainties, risks' n statistics in a structured industrial approach. In *ESREL'2007 safety and reliability conference*, Stavenger, Norway, 2007.
- [3] K. J. Archer and R. V. Kimes. Empirical characterization of random forest variable importance measures. *Computational Statistics & Data Analysis*, 52(4):2249–2260, 2008.
- [4] S.K. Au and J.L. Beck. Estimation of small failure probabilities in high dimensions by subset simulation. *Probabilistic Engineering Mechanics*, 16(4):263–277, 2001.
- [5] B. Auder, A. De Crecy, B. Iooss, and M. Marques. Screening and metamodeling of computer experiments with functional outputs. application to thermal-hydraulic computations. *Reliability Engineering & System Safety*, 107:122–131, 2012.
- [6] B. Auder and B. Iooss. Global sensitivity analysis based on entropy. In *Safety, Reliability and Risk Analysis - Proceedings of the ESREL 2008 Conference*, pages 2107–2115. CRC Press, september 2008.
- [7] M. Balesdent, J. Morio, and J. Marzat. Recommendations for the tuning of rare event probability estimators. *Submitted*, 2011.
- [8] R.J. Beckman and M.D McKay. Monte-Carlo estimation under different distributions using the same simulation. *Technometrics*, 29(2):153–160, 1987.
- [9] J. Bect, D. Ginsbourger, L. Li, V. Picheny, and E. Vazquez. Sequential design of computer experiments for the estimation of a probability of failure. *Statistics and Computing*, 22(3):1–21, 2011.
- [10] P. Bernardara, E. de Rocquigny, N. Goutal, A. Arnaud, and G. Passoni. Uncertainty analysis in flood hazard assessment: hydrological and hydraulic calibration. *Canadian Journal of Civil Engineering*, 37(7):968–979, 2010.
- [11] P. Besse. Apprentissage statistique & data mining. *Institut de Mathématiques de Toulouse, Laboratoire de Statistique et Probabilités UMR CNRS C*, 2008.
- [12] G. Biau, L. Devroye, and G. Lugosi. Consistency of random forests and other averaging classifiers. *The Journal of Machine Learning Research*, 9:2015–2033, 2008.

- [13] E. Borgonovo. A new uncertainty importance measure. *Reliability Engineering & System Safety*, 92(6):771–784, 2007.
- [14] E. Borgonovo, W. Castaings, and S. Tarantola. Moment independent importance measures: new results and analytical test cases. *Risk Analysis*, 31(3):404–428, 2011.
- [15] L. Breiman. Bagging predictors. *Machine learning*, 24(2):123–140, 1996.
- [16] L. Breiman. Random forests. *Machine learning*, 45(1):5–32, 2001.
- [17] L. Breiman, J. Friedman, C. J. Stone, and R. A. Olshen. *Classification and regression trees*. Chapman & Hall/CRC, 1984.
- [18] K. Breitung. Asymptotic approximations for multinormal integrals. *Journal of Engineering Mechanics*, 110:357, 1984.
- [19] B. Briand. *Construction d’arbres de discrimination pour expliquer les niveaux de contamination radioactive des végétaux*. PhD thesis, Université Montpellier II, France, 2008.
- [20] B. Briand, G. R. Ducharme, V. Parache, and C. Mercat-Rommens. A similarity measure to assess the stability of classification trees. *Computational Statistics & Data Analysis*, 53(4):1208–1217, 2009.
- [21] Y. Caniou. *Analyse de sensibilité globale pour les modèles imbriqués et multi-échelles*. PhD thesis, Université Blaise Pascal-Clermont-Ferrand II, France, 2012.
- [22] C. Cannaméla. *Apport des méthodes probabilistes dans la simulation du comportement sous irradiation du combustible à particules*. PhD thesis, Université Paris VII, France, 2007.
- [23] F. Cérou, P. Del Moral, T. Furon, and A. Guyader. Sequential Monte Carlo for rare event estimation. *Statistics and Computing*, 22(3):795–808, 2012.
- [24] S.-H. Cha. Comprehensive survey on distance/similarity measures between probability density functions. *International Journal of Mathematical Models and Methods in Applied Sciences*, 1(4):300–307, 2007.
- [25] T.M. Cover and J.A. Thomas. Elements of information theory 2nd edition (Wiley series in telecommunications and signal processing). 2006.
- [26] I. Csiszár. I-divergence geometry of probability distributions and minimization problems. *The Annals of Probability*, 3(1):146–158, 1975.
- [27] R.I. Cukier, C.M. Fortuin, K.E. Shuler, A.G. Petschek, and J.H. Schaibly. Study of the sensitivity of coupled reaction systems to uncertainties in rate coefficients. I theory. *The Journal of Chemical Physics*, 59:3873, 1973.
- [28] 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.
- [29] F. Dagnegger. Tree stability diagnostics and some remedies for instability. *Statistics in medicine*, 19(4):475–491, 2000.
- [30] E. de Rocquigny, N. Devictor, and S. Tarantola. *Uncertainty in industrial practice: a guide to quantitative uncertainty management*. Wiley, 2008.

-
- [31] A. Der Kiureghian and T. Dakessian. Multiple design points in first and second-order reliability. *Structural Safety*, 20(1):37–49, 1998.
 - [32] O. Ditlevsen. Generalized second moment reliability index. *Journal of Structural Mechanics*, 7(4):435–451, 1979.
 - [33] V. Dubourg. *Méta-modèles adaptatifs pour l’analyse de fiabilité et l’optimisation sous contrainte fiabiliste*. PhD thesis, Université Blaise Pascal - Clermont-Ferrand II, France, 2011.
 - [34] A. Dutfoy and R. Lebrun. Le test du maximum fort: une façon efficace de valider la qualité d’un point de conception. *Actes du 18ème Congrès Français de Mécanique (Grenoble 2007)*, 2007.
 - [35] J.-C. Fort, T. Klein, and N. Rachdi. New sensitivity analysis subordinated to a contrast. *arXiv preprint arXiv:1305.2329*, 2013.
 - [36] Y. Freund, R. E. Schapire, et al. Experiments with a new boosting algorithm. In *Machine Learning*, pages 148–156, 1996.
 - [37] H. C. Frey, A. Mokhtari, and T. Danish. Evaluation of selected sensitivity analysis methods based upon applications to two food safety process risk models. *Dept. of Civil, Construction, and Environmental Eng., North Carolina State Univ., Raleigh, NC*, 2003.
 - [38] H. C. Frey, A. Mokhtari, and J. Zheng. Recommended practice regarding selection, application, and interpretation of sensitivity analysis methods applied to food safety risk process models. *US Department of Agriculture*. <http://www.ce.ncsu.edu/risk/Phase3Final.pdf>, 2004.
 - [39] R. Genuer. *Forêts aléatoires: aspects théoriques, sélection de variables et applications*. PhD thesis, Université Paris Sud-Paris XI, France, 2010.
 - [40] B. Ghattas. *Importance des variables dans les méthodes CART*. Report Universités d’Aix-Marseille II et III, France, 2000.
 - [41] A. Gille Genest. *Utilisation des méthodes numériques probabilistes dans les applications au domaine de la fiabilité des structures*. PhD thesis, Université Paris VI, 1999.
 - [42] Z. Govindarajulu. *Nonparametric inference*. World Scientific, 2007.
 - [43] A.M. Hasofer and N.C. Lind. Exact and invariant second-moment code format. *Journal of the Engineering Mechanics Division*, 100(1):111–121, 1974.
 - [44] T. Hastie, R. Tibshirani, J. Friedman, and J. Franklin. The elements of statistical learning: data mining, inference and prediction. *The Mathematical Intelligencer*, 27(2):83–85, 2005.
 - [45] T.C. Hesterberg. Estimates and confidence intervals for importance sampling sensitivity analysis. *Mathematical and Computer Modelling*, 23(8):79–85, 1996.
 - [46] W. Hoeffding. A class of statistics with asymptotically normal distributions. *Annals of Mathematical Statistics*, 19:293–325, 1948.
 - [47] T. Homma and A. Saltelli. Importance measures in global sensitivity analysis of nonlinear models. *Reliability Engineering & System Safety*, 52(1):1–17, 1996.

- [48] G. M. Hornberger and R.C. Spear. Approach to the preliminary analysis of environmental systems. *J. Environ. Manage.*, 12(1), 1981.
- [49] B. Iooss. Revue sur l'analyse de sensibilité globale de modèles numériques. *Journal de la Société Française de Statistique*, 152(1):1–23, 2011.
- [50] B. Iooss, F. Van Dorpe, and N. Devictor. Response surfaces and sensitivity analyses for an environmental model of dose calculations. *Reliability Engineering & System Safety*, 91(10):1241–1251, 2006.
- [51] T. Ishigami and T. Homma. An importance quantification technique in uncertainty analysis for computer models. In *Uncertainty Modeling and Analysis, 1990. Proceedings., First International Symposium on*, pages 398–403. IEEE, 1990.
- [52] J. Jakubowicz, S. Lefebvre, F. Maire, and E. Moulines. Detecting aircraft with a low-resolution infrared sensor. *Image Processing, IEEE Transactions on*, 21(6):3034–3041, 2012.
- [53] A. Janon, T. Klein, A. Lagnoux-Renaudie, M. Nodet, and C. Prieur. Asymptotic normality and efficiency of two Sobol index estimators. *ESAIM: Probability and Statistics*, doi:10.1051/ps/2013040, 2013.
- [54] M. Jansen. Analysis of variance designs for model output. *Computer Physics Communications*, 117(1):35–43, 1999.
- [55] B. Krykacz-Hausmann. Epistemic sensitivity analysis based on the concept of entropy. *Proceedings of SAMO2001*, pages 31–35, 2001.
- [56] A. Lagnoux. Rare event simulation. *Probability in the Engineering and Informational Sciences*, 20(1):45, 2006.
- [57] R. Lebrun and A. Dutfoy. Do Rosenblatt and Nataf isoprobabilistic transformations really differ? *Probabilistic Engineering Mechanics*, 24(4):577–584, 2009.
- [58] R. Lebrun and A. Dutfoy. A generalization of the Nataf transformation to distributions with elliptical copula. *Probabilistic Engineering Mechanics*, 24(2):172–178, 2009.
- [59] R. Lebrun and A. Dutfoy. An innovating analysis of the Nataf transformation from the copula viewpoint. *Probabilistic Engineering Mechanics*, 24(3):312–320, 2009.
- [60] M. Lemaire. *Structural reliability*. Wiley-ISTE, 2010.
- [61] M. Lemaire, A. Chateauneuf, and J.C. Mitteau. *Structural reliability*. Wiley Online Library, 2009.
- [62] P. Lemaître and A. Arnaud. Hiérarchisation des sources d'incertitudes vis à vis d'une probabilité de dépassement de seuil - une méthode basée sur la pondération des lois. In *Proceedings des 43 èmes Journées de Statistique*, Tunis, Tunisie, June 2011.
- [63] P. Lemaître, E. Sergienko, A. Arnaud, N. Bousquet, F. Gamboa, and B. Iooss. Density modification based reliability sensitivity analysis. *Journal of Statistical Computation and Simulation*, 2014, In press.
- [64] L. Li. *Sequential Design of Experiments to Estimate a Probability of Failure*. PhD thesis, Supelec, France, 2012.

-
- [65] H. Liu, W. Chen, and A. Sudjianto. Relative entropy based method for probabilistic sensitivity analysis in engineering design. *Journal of Mechanical Design*, 128:326, 2006.
- [66] H. O. Madsen, S. Krenk, and N. C. Lind. *Methods of structural safety*. Courier Dover Publications, 2006.
- [67] D. Makowski, C. Naud, M.H. Jeuffroy, A. Barbottin, and H. Monod. Global sensitivity analysis for calculating the contribution of genetic parameters to the variance of crop model prediction. *Reliability Engineering & System Safety*, 91(10-11):1142–1147, 2006.
- [68] S. Mishra, N. E. Deeds, and B. S. RamaRao. Application of classification trees in the sensitivity analysis of probabilistic model results. *Reliability Engineering & System Safety*, 79(2):123–129, 2003.
- [69] S. Mishra, William H. Statham, Neil E. Deeds, and Banda S. RamaRao. Measures of uncertainty importance for monte carlo simulation result. In *Proceedings of Probabilistic Safety Assessment*, 2002.
- [70] H. Monod, C. Naud, and D. Makowski. Uncertainty and sensitivity analysis for crop models. *D. Wallach, D. Makowski et J. Jones, editors: Working with Dynamic Crop Models, chapter 3*, pages 55–100, 2006.
- [71] M.D. Morris. Factorial sampling plans for preliminary computational experiments. *Technometrics*, 33(2):161–174, 1991.
- [72] M. Munoz Zuniga. *Méthodes stochastiques pour l'estimation contrôlée de faibles probabilités sur des modèles physiques complexes. Application au domaine nucléaire*. PhD thesis, Université Paris VII, France, 2011.
- [73] M. Munoz Zuniga, J. Garnier, E. Remy, and E. de Rocquigny. Adaptive directional stratification for controlled estimation of the probability of a rare event. *Reliability Engineering & System Safety*, 92(12):1691–1712, 2011.
- [74] R.B. Nelsen. *An introduction to copulas*. Springer Verlag, 2006.
- [75] H. Niederreiter. Quasi-Monte Carlo methods and pseudo-random numbers. *Journal of American Mathematical Society*, 84(6):957–1041, 1978.
- [76] F. Pappenberger, K. J. Beven, M. Ratto, and P. Matgen. Multi-method global sensitivity analysis of flood inundation models. *Advances in Water Resources*, 31(1):1–14, 2008.
- [77] F. Pappenberger, I. Iorgulescu, and K. J. Beven. Sensitivity analysis based on regional splits and regression trees (SARS-RT). *Environmental Modelling & Software*, 21(7):976–990, 2006.
- [78] C. K Park and K. Ahn. A new approach for measuring uncertainty importance and distributional sensitivity in probabilistic safety assessment. *Reliability Engineering & System Safety*, 46(3):253–261, 1994.
- [79] R. Pastel. *Estimation de probabilités d'événements rares et de quantiles extrêmes. Applications dans le domaine aérospatial*. PhD thesis, Université de Rennes I, France, 2012.
- [80] D Pollard. *Asymptopia*. Book in progress, 2000.

- [81] J. R. Quinlan. Simplifying decision trees. *International Journal of Man-Machine Studies*, 27(3):221–234, 1987.
- [82] R. Rackwitz and B. Flessler. Structural reliability under combined random load sequences. *Computers & Structures*, 9(5):489–494, 1978.
- [83] H. Rinne. *The Weibull distribution: a handbook*. CRC Press, 2010.
- [84] M. Rosenblatt. Remarks on a multivariate transformation. *The Annals of Mathematical Statistics*, 23(3):470–472, 1952.
- [85] R.Y. Rubinstein. *Simulation and the Monte Carlo method*. Wiley Series in Probability and Mathematical Statistics, 1981.
- [86] L. Ruey-Hsia. *Instability of decision tree classification algorithms*. PhD thesis, University of Illinois, 2001.
- [87] A. Saltelli. Making best use of model evaluations to compute sensitivity indices. *Computer Physics Communications*, 145(2):280–297, 2002.
- [88] A. Saltelli, P. Annoni, I. Azzini, F. Campolongo, M. Ratto, and S. Tarantola. Variance based sensitivity analysis of model output. Design and estimator for the total sensitivity index. *Computer Physics Communications*, 181(2):259 – 270, 2010.
- [89] A. Saltelli, S. Tarantola, F. Campolongo, and M. Ratto. *Sensitivity analysis in practice: a guide to assessing scientific models*. John Wiley & Sons Inc, 2004.
- [90] A. Saltelli, S. Tarantola, and K.P.S. Chan. A quantitative model-independent method for global sensitivity analysis of model output. *Technometrics*, 41(1):39–56, 1999.
- [91] M. Sauvé and C. Tuleau-Malot. Variable selection through CART. *arXiv:1101.0689*, 2012.
- [92] I.M. Sobol. Sensitivity analysis for non-linear mathematical models. *Mathematical Modelling and Computational Experiment*, 1(1):407–414, 1993.
- [93] I.M. Sobol. Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Mathematics and Computers in Simulation*, 55(1-3):271–280, 2001.
- [94] I.M. Sobol, S. Tarantola, D. Gatelli, S.S. Kucherenko, and W. Mauntz. Estimating the approximation errors when fixing unessential factors in global sensitivity analysis. *Reliability Engineering and System Safety*, 92:957–960, 2007.
- [95] C. Strobl, A. Boulesteix, A. Zeileis, and T. Hothorn. Bias in random forest variable importance measures: Illustrations, sources and a solution. *BMC bioinformatics*, 8(1):25, 2007.
- [96] S. Tarantola, D. Gatelli, and TA Mara. Random balance designs for the estimation of first order global sensitivity indices. *Reliability Engineering & System Safety*, 91(6):717–727, 2006.
- [97] J.-Y. Tissot and C. Prieur. Bias correction for the estimation of sensitivity indices based on random balance designs. *Reliability Engineering & System Safety*, 107:205–213, 2012.
- [98] A. W. Van der Vaart. *Asymptotic statistics*, volume 3. Cambridge University Press, 2000.
- [99] S. Varet. *Développement de méthodes statistiques pour la prédiction d’un gabarit de signature infrarouge*. PhD thesis, Université Toulouse III, France, 2010.

- [100] E. Volkova, B. Iooss, and F. Van Dorpe. Global sensitivity analysis for a numerical model of radionuclide migration from the rrc kurchatov institute radwaste disposal site. *Stochastic Environmental Research and Risk Assessment*, 22(1):17–31, 2008.
- [101] H. Weyl. Mean motion. *American Journal of Mathematics*, 60(4):889–896, 1938.
- [102] Y. Zhang and A. Der Kiureghian. Two improved algorithms for reliability analysis. In *Reliability and Optimization of Structural Systems*, pages 297–304. Proceedings of the 6th IFI PWG7, 1994.

Appendix A

Distributions formulas

Distribution	Parameters	pdf	Support
Uniform	a, b	$f(x) = \frac{1}{b-a}$	$[a, b]$
Weibull	a, b, c	$f(x) = \frac{c}{b} \left(\frac{x-a}{b}\right)^{c-1} \exp \left[-\left(\frac{x-a}{b}\right)^c\right]$	$x \geq a$
Lognormal	μ, σ	$f(x) = \frac{1}{x\sigma\sqrt{2\pi}} e^{-\frac{(\ln x - \mu)^2}{2\sigma^2}}$	$x > 0$
Gaussian	μ, σ	$f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$	$] -\infty, +\infty[$

Table A.1: Distributions of the random physical variables taken for the CWNr models.

Appendix B

Test cases

In the present subsection, usual sensitivity test cases will be presented. They will be used as benchmark cases for the sensitivity analysis methods. One should note that these test cases return binary values, failure or non-failure of the studied system. One should notice that the focus is set on the probability $P_f = \mathbb{P}(G(\mathbf{X}) \leq 0)$.

B.1 Hyperplane test case

For the first case, $\mathbf{X}$ is set to be a d -dimensional vector, with d independent marginals normally distributed. Unless otherwise mentioned (that is to say for the last case), one has $f_i \sim \mathcal{N}(0, 1)$ for $i = 1, \dots, d$. The failure function $G(\cdot)$ is defined as:

$$G(\mathbf{X}) = k - \sum_{i=1}^d a_i X_i \quad (\text{B.1})$$

where k is a threshold and $\mathbf{a} = (a_1, \dots, a_d)$ are the parameters of the model. One can see that the model is solely linear. What can be expected in terms of SA is that the influence of each variable on P_f depends on its coefficient, namely a_i . The greater the absolute value of the coefficient is, the bigger the expected influence is. One can, by adjusting k , set the failure probability P_f to a value of interest. An explicit expression for P_f can be given as the sum of the d variables behaves like a

Gaussian distribution with parameters 0 and standard deviation $\sqrt{\sum_{i=1}^d a_i^2}$, unless in the last case.

In table B.1 the usual test cases that will be employed throughout the document are detailed.

Number of variables	Values of a_i	Value of k	Value of P
4	$(1, -6, 4, 0)$	16	0.014
5	$a_i = 1 \ \forall i = 1 : 5$	6	0.0036
15	$a_i = 1 \ \forall i = 1 : 5$ $a_i = 0.2 \ \forall i = 6 : 10$ $a_i = 0 \ \forall i = 11 : 15$	6	0.00425
5	$\mathbf{a} = (\frac{1}{2}, \dots, \frac{1}{10})$	5	0.0036

Table B.1: Usual hyperplane test cases

In the first test case, with the specific values of $\mathbf{a}$, the influence of X_2 is greater than the influence of X_3 which is greater than X_1 's. X_4 has no impact on the output. It should be noted that X_1 and X_3 The aim of choosing one non-influential variable is to assess if the SA methods can identify this variable as non-influential on the failure probability.

In the second test case, with all the components equally influential, the aim is to assess or infirm the capability of the SA method to give the same importance to each input.

In the third case, the SA method is put to the test of determining the influential from the little-influential and non-influential variables.

In the last test case, the impact of having variables with the same importance, but distributed with a different spread is studied. Precisely, variables are such that $f_i \sim \mathcal{N}(0, \sigma = 2i)$ for $i = 1..5$. Thus given the a_i , the variables have the same impact on the failure probability. The aim of this test is to assess or infirm the capability of the SA method to give to each equally contributing variable the same importance, despite their different spread.

B.2 Tresholded Ishigami function

The Ishigami function (Ishigami [51]) is a common test case in SA since it has a complex expression, with interactions between the variables. A modified version of the Ishigami function will be considered here. A threshold is added to the value obtained with the regular expression and this is considered as the failure function. Therefore:

$$G(\mathbf{X}) = \sin(X_1) + 7 \sin^2(X_2) + 0.1 X_3^4 \sin(X_1) + k \quad (\text{B.2})$$

where $k = 7$. $\mathbf{X}$ is a 3-dimensional vector of independent marginals uniformly distributed on $[-\pi, \pi]$. In figure B.1, the failure points (where $G(\mathbf{x}) < 0$) are plotted in a 3-d scatterplot.

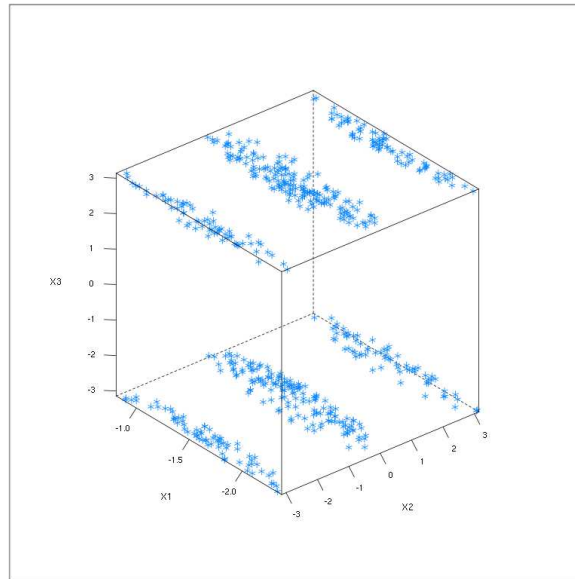


Figure B.1: Ishigami failure points from a MC sample

The failure probability here is roughly $\hat{P} = 5.89 \times 10^{-3}$ (estimated by Monte-Carlo technique, see section 1.2.1). The complex repartition of the failure points can be noticed. Those points lay in a zone defined by the negative values of X_1 , the extremal and mean values of X_2 (around $-\pi$, 0 and π), and the extremal values of X_3 (around $-\pi$ and π)

B.3 Flood case

The goal of this test case is to assess the risk of a flood over a dyke for the safety of industrial installations (Bernardara [10]). This comes down to model the level of a flood. As a function of hydraulic parameters, many of them being randomized to account for uncertainty. From a simplification of the Saint-Venant equation, a flood risk model is obtained.

The quantity of interest is the difference between the level of the dyke and the height of water. If this quantity is negative, the installation is flooded. Hydraulic parameters are the following: Q the flow rate, L the watercourse section length studied, B the watercourse width, K_s the watercourse bed friction coefficient (also called Strickler coefficient), Z_m and Z_v respectively the upstream and downstream bottom watercourse level above sea level and H_d the dyke height measured from the bottom of the watercourse bed. The water level model is expressed as:

$$H = \left(\frac{Q}{K_s B \sqrt{\frac{Z_m - Z_v}{L}}} \right)^{\frac{3}{5}}. \quad (\text{B.3})$$

Therefore the following quantity is considered:

$$G = H_d - (Z_v + H). \quad (\text{B.4})$$

Among the model inputs, the choice is made that the following variables are known precisely: $L = 5000$ (m), $B = 300$ (m), $H_d = 58$ (m), and the following are considered to be random. Q ($\text{m}^3 \cdot \text{s}^{-1}$) follows a positively truncated Gumbel distribution of parameters $a = 1013$ and $b = 558$ with a minimum value of 0. K_s ($\text{m}^{1/3} \cdot \text{s}^{-1}$) follows a truncated Gaussian distribution of parameters $\mu = 30$ and $\sigma = 7.5$, with a minimum value of 1. Z_v (m) follows a triangular distribution with minimum 49, mode 50 and maximum 51. Z_m (m) follows a triangular distribution with minimum 54, mode 55 and maximum 56.

The failure probability here is roughly $\hat{P} = 7.88 \times 10^{-4}$ (estimated by MC technique, see 1.2.1).

Appendix C

Isoprobabilistic transformations

Here, we briefly introduce the notion of copula, which is needed for the presentation of isoprobabilistic transformations. Copulas are a mathematical object describing the dependencies in a random vector without referring to the marginal distributions. Nelsen's monograph [74] presents such objects.

C.1 Presentation of the copulas

Definition C.1.1 A d dimensional function f is said d -increasing if:

$$\sum_{i_1=1}^2 \cdots \sum_{i_d=1}^2 (-1)^{i_1+\cdots+i_d} f(x_{1,i_1}, \dots, x_{d,i_d}) \geq 0$$

where $x_{j,1} = a_j$ and $x_{j,2} = b_j \forall j \in \{1, \dots, d\}$ and $a_j, b_j \in [0, 1], a_j \leq b_j \forall j \in \{1, \dots, d\}$

Definition C.1.2 A d -dimensional copula C is a d -dimensional cumulative distribution function defined over $[0, 1]^d$, whose marginal distributions are uniform over $[0, 1]$:

- C is d -increasing;
- for all $\mathbf{u} \in [0, 1]^d$ which have at least one component equal to 0, $C(\mathbf{u}) = 0$;
- for all $\mathbf{u} \in [0, 1]^d$ which have all their components equal to 1 except one, u_k , $C(\mathbf{u}) = u_k$.

Theorem C.1.1 (Sklar 1959)

Let F be a d -dimensional cumulative distribution function with $F_1, \dots, F_p$ the marginal distribution functions. There exists a d -dimensional copula, C , such that for all $\mathbf{x} \in \mathbb{R}^d$ we have:

$$F(x_1, \dots, x_p) = C(F_1(x_1), \dots, F_p(x_p)). \quad (\text{C.1})$$

If the marginal distributions $F_1, \dots, F_p$ are continuous, then the copula C is unique, otherwise it is uniquely determined over $\text{Im}(F_1) \times \cdots \times \text{Im}(F_p)$. In the continuous case, for all $\mathbf{u} \in [0, 1]^d$ we have:

$$C(\mathbf{u}) = F(F_1^{-1}(u_1), \dots, F_p^{-1}(u_p)) \quad (\text{C.2})$$

if absolutely continuous

$$f(\mathbf{x}) = c((F_1(x_1), \dots, F_p(x_p))) \prod_{i=1}^d f_i(x_i) \quad (\text{C.3})$$

with c the probability distribution function associated to C , f the probability distribution function associated to F and f_i the marginal distributions function associated to F .

Definition C.1.3 Let us denote $\mathcal{SO}_d(\mathbb{R})$ the rotation group over R^d and $\text{supp}(\mathbf{X})$ the set of the values that can be taken by a random vector $\mathbf{X}$. An isoprobabilistic transformation T of a d -dimensional random vector X is a diffeomorphism from $\text{supp}(\mathbf{X})$ into R^d such that the random vectors $\mathbf{U} = T(\mathbf{X})$ and $r\mathbf{U}$ have the same distribution for all $r \in \mathcal{SO}_d(\mathbb{R})$.

C.2 Objectives, Rosenblatt transformation

We wish to transform a random vector $\mathbf{X}$ of pdf $f_{\mathbf{X}}$ and of copula C in a Gaussian vector $\mathbf{U}$ of same dimension but with independent, standard Gaussian as components.

If the variables are independent and that the marginals are known, the transformation is straightforward :

$$u_i = \phi^{-1}(F_i(x_i))$$

If there is a dependency structure in the variables, Rosenblatt and Nataf transformations are possibilities [?].

We present here the Rosenblatt [84] transformation. This transformation is not unique if the variables are correlated: it depends on the order in which the variables are transformed ¹.

Transformation is done as follows:

$$u_1 = \phi^{-1}(F_1(x_1))$$

$$u_2 = \phi^{-1}(F_2(x_2|X_1 = x_1))$$

...

$$u_d = \phi^{-1}(F_d(x_d|X_1 = x_1, \dots, X_{d-1} = x_{d-1}))$$

where $F_i(\cdot|X_1, \dots, X_{i-1})$ is the cdf of variable X_i given the realisations of the previous variables.

¹It has been shown in Lebrun and Dutfoy [57] that if the copula of $\mathbf{X}$ is Gaussian, the order in which the variables are transformed does neither impact the norm of the design point, nor the derivatives of the failure surface in this point. In other words, the following quantities use in FORM/SORM methods do not depend upon the order of transformation: $\beta_{HL}, \hat{P}_{FORM}, \hat{P}_{SORM}$.

Appendix D

Appendices for Chapter 3

D.1 Proofs of asymptotic properties

Proof of Lemma 3.2.1

Under assumption (i), we have

$$\int_{\text{Supp}(f_{i\delta})} \mathbf{1}_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}(x_i)}{f_i(x_i)} f(\mathbf{x}) \, d\mathbf{x} \leq \int_{\text{Supp}(f_{i\delta})} f_{i\delta}(x_i) \, dx_i = 1.$$

So that, the strong LLN may be applied to $\hat{P}_{i\delta N}$. Defining

$$\sigma_{i\delta}^2 = \text{Var} \left[\mathbf{1}_{\{G(\mathbf{X}) < 0\}} \frac{f_{i\delta}(X_i)}{f_i(X_i)} \right], \quad (\text{D.1})$$

one has

$$\sigma_{i\delta}^2 = \int_{\text{Supp}(f_i)} \mathbf{1}_{\{G(\mathbf{x}) < 0\}} \frac{f_{i\delta}^2(x_i)}{f_i(x_i)} \prod_{j \neq i}^d f_j(x_j) \, d\mathbf{x} - P_{i\delta}^2 < \infty \quad \text{under assumption (ii).}$$

Therefore the CLT applies:

$$\sqrt{N} \sigma_{i\delta}^{-1} \left(\hat{P}_{i\delta N} - P_{i\delta} \right) \xrightarrow{\mathcal{L}} \mathcal{N}(0, 1).$$

Under assumption (ii), the strong LLN applies to $\hat{\sigma}_{i\delta N}^2$. So that, the final result is straightforward using Slutsky's lemma.

Proof of Proposition 3.2.1

First, note that

$$\begin{aligned} \mathbb{E} \left[\widehat{P} \widehat{P}_{i\delta} \right] - PP_{i\delta} &= \mathbb{E} \left[\frac{1}{N^2} \left(\sum_{n=1}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \right) \left(\sum_{n=1}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)} \right) \right] - PP_{i\delta} \\ &= \frac{1}{N^2} \mathbb{E} \left[\sum_{n=1}^N \left[\mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \right]^2 \frac{f_{i\delta}(x_i^n)}{f_i(x_i^n)} + \sum_{n=1}^N \sum_{j \neq i}^N \mathbf{1}_{\{G(\mathbf{x}^n) < 0\}} \mathbf{1}_{\{G(\mathbf{x}^j) < 0\}} \frac{f_{i\delta}(x_i^j)}{f_i(x_i^j)} \right] \\ &\quad - PP_{i\delta} \\ &= \frac{1}{N^2} [NP_{i\delta} + N(N-1)PP_{i\delta}] - PP_{i\delta} \\ &= \frac{1}{N} (P_{i\delta} - PP_{i\delta}). \end{aligned}$$

Assuming the conditions under which Lemma 1 is true, the bivariate CLT follows with

$$\Sigma_{i\delta} = \begin{pmatrix} P(1-P) & P_{i\delta}(1-P) \\ P_{i\delta}(1-P) & \sigma_{i\delta}^2 \end{pmatrix}.$$

Each term of this matrix can be consistently estimated, using the results in Lemma 1 and Slutsky's lemma.

D.2 Computation of Lagrange multipliers

Let H be the Lagrange function:

$$H(\boldsymbol{\lambda}) = \psi_i(\boldsymbol{\lambda}) - \sum_{k=1}^K \lambda_k \delta_k.$$

Thus, using the results of Csizsar [26], one has

$$\boldsymbol{\lambda}^* = \arg \min H(\boldsymbol{\lambda}).$$

The expression of the gradient of H with respect to the j^{th} variable is

$$\nabla_j H(\boldsymbol{\lambda}) = \frac{\int g_j(x) f_i(x) \exp(\sum_{k=1}^K \lambda_k g_k(x)) dx}{\exp \psi_i(\boldsymbol{\lambda})} - \delta_j.$$

Similarly, the expression of the second derivative of H with respect to the h^{th} and the j^{th} variables is

$$\begin{aligned} D_{hj} H(\boldsymbol{\lambda}) &= \frac{\int g_h(x) g_j(x) f_i(x) \exp(\sum_{k=1}^K \lambda_k g_k(x)) dx}{\exp \psi_i(\boldsymbol{\lambda})} \\ &- \frac{\int g_j(x) f_i(x) \exp(\sum_{k=1}^K \lambda_k g_k(x)) dx}{\exp \psi_i(\boldsymbol{\lambda})} \frac{\int g_h(x) f_i(x) \exp(\sum_{k=1}^K \lambda_k g_k(x)) dx}{\exp \psi_i(\boldsymbol{\lambda})}. \end{aligned}$$

This method has been used in this paper for computing the optimal vector $\boldsymbol{\lambda}^*$ when a variance shifting was applied. The integrals were evaluated with Simpson's rule.

D.3 Proofs of the NEF properties

In this Appendix, the details of the calculus for the Proposition 3.3.4 are provided.

NEF specificities : If the original density $f_i(x)$ is a NEF, then under a set of K linear constraints on $f(x)$, one has :

$$f(x) = b(x) \exp [x\theta - \eta(\theta)],$$

thus :

$$f_\delta(x) = f(x) \exp \left[\sum_{k=1}^K \lambda_k g_k(x) - \psi(\boldsymbol{\lambda}) \right]$$

The regularization constant from (3.13) can be written as:

$$\psi(\boldsymbol{\lambda}) = \log \int b(x) \exp \left[x\theta + \sum_{k=1}^K \lambda_k g_k(x) - \eta(\theta) \right] dx \quad (\text{D.2})$$

If the integral on (D.2) is finite, f_δ exists and is a density.

Mean shifting With a single constraint formulated as in (3.15), (D.2) becomes :

$$\begin{aligned}\psi(\boldsymbol{\lambda}) &= \log \int b(x) \exp [x\theta + \lambda x - \eta(\theta)] dx \\ &= \log \int b(x) \exp [x(\theta + \lambda) - \eta(\theta) + \eta(\theta + \lambda) - \eta(\theta + \lambda)] dx\end{aligned}$$

if $\eta(\theta + \lambda)$ is well defined.

$$\begin{aligned}\psi(\boldsymbol{\lambda}) &= (\eta(\theta + \lambda) - \eta(\theta)) + \log \left[\int b(x) \exp [x(\theta + \lambda) - \eta(\theta + \lambda)] dx \right] \\ &= \eta(\theta + \lambda) - \phi(\theta)\end{aligned}$$

since

$$b(x) \exp [x(\theta + \lambda) - \eta(\theta + \lambda)] = f_{\theta+\lambda}(x)$$

with notation from (3.3.4), is a density of integral 1. Thus

$$\begin{aligned}f_{\delta}(x) &= b(x) \exp [x\theta - \phi(\theta)] \exp [\lambda x - \eta(\theta + \lambda) + \eta(\theta)] \\ &= b(x) \exp [x(\theta + \lambda) - \eta(\theta + \lambda)] = f_{\theta+\lambda}(x)\end{aligned}$$

Thus the mean shifting of a NEF of CDF $\eta(\cdot)$ results in another NEF with mean $\eta'(\theta + \lambda) = \delta$ (constraint) and variance $\eta''(\theta + \lambda)$.

Variance shifting With a single constraint formulated as in (3.19), using (D.2), the new distribution has for density:

$$f_{\delta}(x) = b(x) \exp [x\theta + x\lambda_1 + x^2\lambda_2 - \psi(\boldsymbol{\lambda}) - \eta(\theta)]$$

Since $\boldsymbol{\lambda}$ is known or computed, and θ is also known, consider the variable change $z = \sqrt{\lambda_2}x$ assuming λ_2 is strictly positive (the variable change is $z = \sqrt{-\lambda_2}x$ if λ_2 is strictly negative). Thus,

$$\begin{aligned}f_{\delta}(x) &= b\left(\frac{z}{\sqrt{\lambda_2}}\right) \exp [z^2] \exp \left[\frac{z}{\sqrt{\lambda_2}} (\theta + \lambda_1) - \psi(\boldsymbol{\lambda}) - \eta(\theta) \right] \\ &= \exp \left[\eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) - \eta(\theta) - \psi(\boldsymbol{\lambda}) \right] c(z) \exp \left[z \frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} - \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) \right]\end{aligned}$$

with

$$c(z) = b\left(\frac{z}{\sqrt{\lambda_2}}\right) \exp [z^2].$$

By (3.13),

$$\begin{aligned}\psi(\boldsymbol{\lambda}) &= \log \int b(x) \exp [x\theta + x\lambda_1 + x^2\lambda_2 - \eta(\theta)] dx \\ &= \log \int b\left(\frac{z}{\sqrt{\lambda_2}}\right) \exp [z^2] \exp \left[\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} z - \eta(\theta) + \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) - \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) \right] dx \\ &= \left(\eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) - \eta(\theta) \right) + \log \int c(z) \exp \left[\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} z - \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) \right] dz \\ &= \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) - \eta(\theta)\end{aligned}$$

Thus one has :

$$f_{\delta}(x) = c(z) \exp \left[z \frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} - \eta \left(\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}} \right) \right]$$

thus the variance shifting of a NEF results in another NEF parameterized by $\frac{(\theta + \lambda_1)}{\sqrt{\lambda_2}}$.

D.4 Numerical trick to work with truncated distribution

In the case where a mean shifting is considered on a left truncated distribution. We present a tip that can help to compute λ^* .

The studied truncated variable Y_T has distribution f_{Y_T} . Let us denote $Y \sim f_Y$ the corresponding non-truncated distribution. The truncation occurs for some real value a . This truncation may happen for some physical modelling reason. One has:

$$f_{Y_T}(y) = \frac{1}{1 - F(a)} 1_{[a, +\infty[}(y) f_Y(y).$$

The formal definition of $M_{Y_T}(\lambda)$ the mgf of Y_T for some λ is:

$$M_{Y_T}(\lambda) = \frac{1}{1 - F_Y(a)} \int_a^{+\infty} f_Y(y) \exp[\lambda y] dy.$$

Let us recall that we are looking for λ^* such as:

$$\delta = \frac{M'_{Y_T}(\lambda^*)}{M_{Y_T}(\lambda^*)} = \frac{\int_a^{+\infty} y f_Y(y) \exp[\lambda y] dy}{\int_a^{+\infty} f_Y(y) \exp[\lambda y] dy}. \quad (\text{D.3})$$

When the expression does not take a practical form, one can use numerical integration to estimate the integral terms. Unfortunately, for some heavy tailed distribution (for instance Gumbel distribution), this numerical integration might be complex or not possible. This is due to the multiplication by an exponential of y . The following tip helps to avoid such problems. Denoting $M_Y(\lambda)$ the mgf of the non-truncated distribution, one can remark that:

$$M_Y(\lambda) = \int_{-\infty}^{+\infty} f_Y(y) \exp[\lambda y] dy = \int_{-\infty}^a f_Y(y) \exp[\lambda y] dy + \int_a^{+\infty} f_Y(y) \exp[\lambda y] dy$$

Thus another expression for $M_{Y_T}(\lambda)$ is:

$$M_{Y_T}(\lambda) = \frac{1}{1 - F_Y(a)} \left[M_Y(\lambda) - \int_{-\infty}^a f_Y(y) \exp[\lambda y] dy \right].$$

The integral term is much smaller in the left heavy tailed distribution case. Therefore the numerical integration (for instance using Simpson's method) is much more precise or became possible.

The same goes for $M'_{Y_T}(\lambda)$ which has alternative expression:

$$M'_{Y_T}(\lambda) = \frac{1}{1 - F_Y(a)} \left[M'_Y(\lambda) - \int_{-\infty}^a y f_Y(y) \exp[\lambda y] dy \right].$$

Finally, another form of D.3 is:

$$\delta = \frac{M'_Y(\lambda) - \int_{-\infty}^a y f_Y(y) \exp[\lambda y] dy}{M_Y(\lambda) - \int_{-\infty}^a f_Y(y) \exp[\lambda y] dy}. \quad (\text{D.4})$$

This alternative expression may lead to more precise estimations of λ^* when $M_Y(\lambda)$ and $M'_Y(\lambda)$ are known (which is the case for most usual distribution) since the integral term are much smaller than in the first expression.

Résumé

Cette thèse porte sur l'analyse de sensibilité dans le contexte des études de fiabilité des structures. On considère un modèle numérique déterministe permettant de représenter des phénomènes physiques complexes. L'étude de fiabilité a pour objectif d'estimer la probabilité de défaillance du matériel à partir du modèle numérique et des incertitudes inhérentes aux variables d'entrée de ce modèle. Dans ce type d'étude, il est intéressant de hiérarchiser l'influence des variables d'entrée et de déterminer celles qui influencent le plus la sortie, ce qu'on appelle l'analyse de sensibilité. Ce sujet fait l'objet de nombreux travaux scientifiques mais dans des domaines d'application différents de celui de la fiabilité. Ce travail de thèse a pour but de tester la pertinence des méthodes existantes d'analyse de sensibilité et, le cas échéant, de proposer des solutions originales plus performantes. Plus précisément, une étape bibliographique sur l'analyse de sensibilité d'une part et sur l'estimation de faibles probabilités de défaillance d'autre part est proposée. Cette étape soulève le besoin de développer des techniques adaptées. Deux méthodes de hiérarchisation de sources d'incertitudes sont explorées. La première est basée sur la construction de modèle de type classifieurs binaires (forêts aléatoires). La seconde est basée sur la distance, à chaque étape d'une méthode de type subset, entre les fonctions de répartition originelle et modifiée. Une méthodologie originale plus globale, basée sur la quantification de l'impact de perturbations des lois d'entrée sur la probabilité de défaillance est ensuite explorée. Les méthodes proposées sont ensuite appliquées sur le cas industriel CWNR, qui motive cette thèse.

Mots-clés Analyse de sensibilité ; Fiabilité ; Incertitudes ; Expériences numériques ; Perturbation des lois

Abstract

This thesis' subject is sensitivity analysis in a structural reliability context. The general framework is the study of a deterministic numerical model that allows to reproduce a complex physical phenomenon. The aim of a reliability study is to estimate the failure probability of the system from the numerical model and the uncertainties of the inputs. In this context, the quantification of the impact of the uncertainty of each input parameter on the output might be of interest. This step is called sensitivity analysis. Many scientific works deal with this topic but not in the reliability scope. This thesis' aim is to test existing sensitivity analysis methods, and to propose more efficient original methods. A bibliographical step on sensitivity analysis on one hand and on the estimation of small failure probabilities on the other hand is first proposed. This step raises the need to develop appropriate techniques. Two variables ranking methods are then explored. The first one proposes to make use of binary classifiers (random forests). The second one measures the departure, at each step of a subset method, between each input original density and the density given the subset reached. A more general and original methodology reflecting the impact of the input density modification on the failure probability is then explored. The proposed methods are then applied on the CWNR case, which motivates this thesis.

Keywords Sensitivity Analysis; Reliability; Uncertainties; Computer experiments; Input perturbations