



Sensitivity analysis and graph-based methods for black-box functions with an application to sheet metal forming.

Jana Fruth

► To cite this version:

Jana Fruth. Sensitivity analysis and graph-based methods for black-box functions with an application to sheet metal forming.. General Mathematics [math.GM]. Ecole Nationale Supérieure des Mines de Saint-Etienne, 2015. English. NNT : 2015EMSE0779 . tel-01151170

HAL Id: tel-01151170

<https://theses.hal.science/tel-01151170>

Submitted on 12 May 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.

New methods for the sensitivity analysis of black-box functions with an application to sheet metal forming

A doctoral Thesis

by

Jana Fruth

in fulfillment of the requirements for the degrees

“Doktorin der Naturwissenschaften der Technischen Universität Dortmund”

and

“Docteur de l'École Nationale Supérieure des Mines de Saint-Étienne”

Submitted to

TU Dortmund University and Mines Saint-Etienne

Dortmund, 2015

Defended on March 12, 2015 with a jury composed of

Prof. Dr. Katja Ickstadt	Head of the jury
Prof. Dr. Sonja Kuhnt	Reviewer
Prof. Clémentine Prieur	Reviewer
Prof. Dr. Joachim Kunert	Reviewer
Prof. Olivier Roustant	Supervisor
Prof. Rodolphe Le Riche	Examiner

Acknowledgements

I deeply thank my supervisors Sonja Kuhnt in Dortmund and Olivier Roustant in Saint-Étienne for giving me the chance to research on such an interesting topic, to get insight into the research work of two different places, and to meet so many exciting people. I feel very lucky for that and as well for the great, constant support which I could not have imagined any better.

Furthermore, I want to thank some of the many people that helped me to write this thesis, e.g. by inspiring scientific conversations and inputs, amazing programming support, great cooperations, constant assistance, or kind encouragements: Thomas Mühlenstädt, Laurent Carraro, Clément Chevalier, David Ginsbourger, Bertrand Iooss, Roland Jegou, Art B. Owen, Clémentine Prieur, Marc Roelens, David Steinberg, Momchil Ivanov, Hamad ul Hassan, Leo Geppert, Max Wornowizki, André König, Uwe Ligges, Olaf Mersmann, Daniel Vogel, Alper Güner, Jörg Kolbe, Stefan Hess, Simon Neumärker, Malte Jastrow, Matthias Borowski, Thoralf Mildenerger, André Rehage, Nikolaus Rudak, Henrike Weinert, Ieva Zelo, Christine Exbrayat, Simon Ferber, Carina Hösl, Stefanie Krombacher, Natalina Külow, and Jakob Wieczorek.

Last but not least, I want to thank the Deutsche Forschungsgesellschaft for funding via SFB 708 project C3.

CONTENTS

1	Introduction	1
2	Sensitivity Analysis	5
2.1	Background	5
2.2	The role of metamodels	6
2.3	Variance-based indices	8
2.3.1	Definition of various variance-based indices	10
2.3.2	Monte Carlo estimation	12
2.3.3	Frequency-based estimation	17
2.4	Derivative-based methods	20
2.4.1	Morris screening by elementary effects	20
2.4.2	DGSM	21
2.5	Ongoing developments in sensitivity analysis	23
3	Sensitivity Analysis for Interaction Screening	25
3.1	Total interaction indices	28
3.2	TII estimation	31
3.3	Theoretical properties of TII estimators	34
3.4	Estimating the full set of TIIs	41
3.5	Simulation study	42

3.6	Threshold decision	48
3.7	Implementation	54
3.8	Crossed DGSM	56
4	Sensitivity Analysis for Functional Input	61
4.1	Sequential functional sensitivity analysis	64
4.2	Error-free test cases	70
4.3	Splitting and design	74
4.4	Implementation	77
4.5	Analytical example	77
5	Support Analysis	83
5.1	Support index functions	84
5.2	Analytical example	89
5.3	Expected values	89
6	Application to a Sheet Metal Forming Process	93
6.1	Thickness reduction	95
6.2	Functional input application	103
7	Conclusion and Outlook	109
A	Notations	111
B	Data	115
	Bibliography	125

CHAPTER 1

1. INTRODUCTION

This work deals with the topic of sensitivity analysis for computer experiments. In the situation that statistical experiments are too expensive, time-consuming, might harm lives, or are simply impossible to perform, computer simulations can sometimes be used as a replacement. These computer experiments are complex mathematical models constructed from known physical relations, e.g. in engineering or environmental problems, or chemical and medical experiments, where computer experiments are also known as *in silico* experiments. Specific examples for computer experiments are traffic simulation models (Punzo and Ciuffo, 2011), where the interest lies in observing those effects that lead to congestions, models of manufacturing processes like the deep drawing model which motivated this work, or highly complex weather models, used in the climate impact research (Katz, 2002). Computer experiments have become everyday procedures in those fields and are also rising in other areas like social sciences and economics. None the less, it has to be kept in mind that computer experiments are only approximations of the real process and should be validated by real experiments whenever possible.

From the statistical point of view, computer experiments have to be treated specially since they are usually deterministic, contain a high number of input variables and are highly complex in terms of nonlinear relations and interactions. The inherent mathematical equations are most often solved numerically which results in long-lasting

computations, so computer experiments can be very time-consuming. Methods have been developed for the design of computer experiments, above all space-filling designs like Latin hypercubes, maximin design, or entropy design (Santner et al., 2003). For deterministic modeling and prediction, advanced interpolation methods exist which respect the high complexity, with the most popular being the Gaussian process emulator, also called Kriging. Finally, there are also optimization procedures for computer experiments, which make use of the prediction models, one popular method being the Kriging-based EGO algorithm, which utilizes the fact that Kriging provides an estimation of the uncertainty of unobserved points.

In this work, the focus lies on a further broad field in the analysis of computer experiments, sensitivity analysis, which has importance in all steps beginning with the building of the computer experiment over its use and understanding to model building and model validation. It analyzes how sensitive the computer experiment responds to variations in the input.

The analysis of the response's sensitivity is of interest on its own, e.g. to answer the question what happens in a car crash simulation when the tire friction changes. Based on Saltelli et al. (2000), five further points of sensitivity analysis application can be named:

- (1) To check, if a model resembles the real process, i.e. if sensitivities reflect expectations,
- (2) to determine the most important input variables and optimal regions in the space of input variables for calibration studies,
- (3) to select and rank input variables by importance,
- (4) to detect regions in the space of input variables for which the model variation is maximum, and
- (5) to detect (groups of) input variables that interact with each other.

This work presents new sensitivity analysis methods, motivated by and framed around applications in the analysis of shape accuracy in sheet metal forming. Shape deformation after the removal of the forming punch is a serious problem in production processes, but can be simulated by finite elements methods (FEM).

With regard to point (5), the total interaction index (TII) for the analysis of computer experiments is presented, an extension to the usual Sobol indices. Defined for two input variables at a time, it gives the sum of Sobol indices of any order interaction containing both input variables. If we estimate the TII for all combinations of the two input variables, we can get an overview of the complete interaction structure of the computer experiment, a difficult task with Sobol indices due to the curse of dimensionality. The structure can be intuitively drawn in a so-called FANOVA graph (Mühlenstädt et al., 2012). Furthermore, similar to the input variable screening by Sobol indices, we can use TIIs for interaction screening. This leads to a block-additive decomposition of the computer experiment. An inactive interaction indicates that the two input variables come from two groups of variables that have no interaction in common and thus are additive. This knowledge about a block-additive structure in the computer experiment can be applied in subsequent procedures. Metamodels like Kriging can be adapted to follow the block-additive structure and optimization can be simplified and parallelized when being performed in each block separately. We present several estimation methods for the TII and analyze their statistical properties. Methods for interaction screening, including a suitable thresholding algorithm for inactive interactions, are developed. In addition, derivative-based indices, which provide a faster-to-compute upper bound for the TII, are presented.

Usually the process parameters are kept on a constant level during the computer experiment. For the deep drawing analysis, finite element simulations have been developed that allow the user to change parameters during the forming process in order to further improve the accuracy of the analysis. Here new methods are necessary for the sensitivity analysis of those temporal changeable parameters. This work presents an idea

for an effective sensitivity analysis of functional input including design and graphical representation of the functional influences of the inputs. To keep the number of function evaluations as low as possible, a sequential algorithm is introduced that increases the accuracy of the functional sensitivity analysis with every step. The points (1), (2) and (3) can be handled by the methods for functional input variables.

The ideas of the sequential algorithm are further developed for the analysis of the support of scalar input variables, the support analysis. Regarding point (4) the aim is to analyze the sensitivity of the different regions of the variables' distribution support. This information gives closer insight into the impact of the variables in the process and can furthermore help in finding optimal support settings for screening, modeling, and sensitivity analysis. In this work, support index functions as well as visualization methods are suggested and their relations to other indices are derived.

The thesis is structured as follows. Chapter 2 gives an introduction to sensitivity analysis together with a panorama of the current state of research in the field with special focus on variance-based global sensitivity indices. Chapter 3 introduces the TII along with the complete methodology for sensitivity analysis of the interaction structure. Chapter 4 presents the approach of sensitivity analysis for functional input and in Chapter 5 the developed technique for support analysis is presented. All methods are then applied in deep drawing simulations in Chapter 6. The work is concluded with a summary and outlook in Chapter 7.

2. SENSITIVITY ANALYSIS

2.1 Background

In the field of computer experiments, sensitivity analysis generally explores the relationships between information flowing in and out of the experiment. More specifically, it studies how the uncertainty in the output of the computer experiment can be apportioned to different sources of uncertainty in the inputs with the direct aim of process understanding and input variable screening, but also calibration, metamodel building and finally optimization. An elaborate overview can be found in the standard work of sensitivity analysis by Saltelli et al. (2000) and in the comparison of sensitivity analysis techniques by Confalonieri et al. (2010). This chapter starts with a brief outline of the background of sensitivity analysis.

In the beginning, sensitivity analysis was based on standard statistical methods like scatter plots. There, the influence of a variable can be read from plots of the output against each input variable. Further methods were regression analysis, where regression coefficients qualify the linear sensitivity of each input (e.g. Kleijnen (1997)) and correlation analysis. A next step were one-at-a-time methods (e.g. Daniel (1973)). Here, one input variable is changed while keeping all other at a constant level, thus information can be obtained only for that particular region in the input space. A similar approach came from chemical modeling (e.g. Chapter 5 of Saltelli et al. (2000)),

which uses partial derivatives around a nominal value as local sensitivities. In contrast to those local methods, Schaibly and Shuler (1973) developed a global sensitivity analysis method called FAST (Fourier amplitude sensitivity test), which allows for all input variables to be varied at the same time. Here, the input variables are sampled in such a way that the amplitudes obtained by a Fourier analysis of the output can be interpreted as sensitivity indices of the input variables. Later on, Sobol' (1993) introduced indices, which are now called Sobol indices, global sensitivity indices, which have been proved by Saltelli and Bolado (1998) to predict the same quantities as the FAST indices. Because of their clear interpretation, unbiased estimation and their access to interactions, Sobol indices have been widely used and are being further developed continuously. Parallel to that, alternative sensitivity analysis methods were introduced like derivative-based indices (Kucherenko et al., 2009) and, based on one-at-a-time methods, very effective screening designs have been developed with the most important being Morris screening (Morris, 1991).

2.2 The role of metamodels

The estimation of sensitivity indices usually requires a high number of model evaluations, as crude discrete integration methods are applied. Time-consuming computer experiments are therefore often replaced by a faster-to-evaluate second model, which is called a metamodel, surrogate model, emulator or response surface. Possible metamodels come from the large field of the modeling of computer experiments (e.g. Fang et al., 2006), a popular one being the Gaussian process model described below, but also polynomials, Polynomial Chaos Expansion, splines, or Neural Networks. As the metamodel is only an approximation of the real experiment, a metamodel error has to be taken into account and the accuracy of the model has to be assessed carefully, e.g. by leave-one-out cross validation.

Throughout this thesis, there will be no distinction between the true computer experiment and the metamodel. The term *underlying model*, symbolized by f , can either refer to the direct computer experiment or to a metamodel with negligible error. The underlying model f will be regarded as a black-box function, whose specific shape is inaccessible. However, there are some approaches that are able to utilize the specific properties of a metamodel. Blatman and Sudret (2010) compute Sobol indices analytically from Sparse Polynomial Chaos Expansion models and Marrel et al. (2009) define indices for the Gaussian process model that use the Gaussian process directly instead of the prediction function.

The Gaussian process model, also called Kriging as it was proposed by Krige (1951), is a standard tool in computer experiments for various reasons. It interpolates the data, which is important for deterministic computer codes. In addition, the Kriging model performs well in predicting, even for highly complex functions, and it gives a measure of uncertainty at unobserved points (Fang et al., 2006). The basic idea is to assume that the function is a realization of a Gaussian process with a linear trend as mean,

$$Y(\mathbf{x}) = \sum_{\ell=1}^p \beta_{\ell} f_{\ell}(\mathbf{x}) + Z(\mathbf{x}), \quad (2.1)$$

where $\mathbf{x} = (x_1, \dots, x_d)$ represents the vector of d input variables, $f_1(\mathbf{x}), \dots, f_p(\mathbf{x})$ are p known regression functions with $\beta_1, \dots, \beta_p$ the corresponding parameters, and $Z(\cdot)$ is a Gaussian process with zero mean and covariance function, or kernel k . A particular class for the covariance kernel is the stationary family, which implies that k depends only on the difference between two different locations, $k(\mathbf{x}^{(1)} - \mathbf{x}^{(2)})$ with $k(\cdot) = \sigma^2 R(\cdot; \boldsymbol{\theta})$, where σ^2 is the process variance, R the correlation function and $\boldsymbol{\theta}$ a vector of covariance parameters.

As the Gaussian process model in (2.1) returns a Gaussian process instead of a function, the predictor has to be defined separately,

$$\hat{y}(\mathbf{x}^*) = \sum_{\ell=1}^p \beta_{\ell} f_{\ell}(\mathbf{x}^*) + \mathbf{k}(\mathbf{x}^*)' \mathbf{K}^{-1}(\mathbf{y} - \mathbf{F}\boldsymbol{\beta}), \quad (2.2)$$

with $\mathbf{x}^*$ a new point at which to predict, $\mathbf{K}$ the covariance matrix at all data points, $\mathbf{k}(\mathbf{x}^*)$ the covariance vector between the data and $\mathbf{x}^*$, and $\mathbf{F}$ the experimental matrix containing the trend values at all data points. Usually, the parameters β_k , σ^2 and $\boldsymbol{\theta}$ are estimated by numerical maximization of the likelihood and inserted into Equation (2.2).

In practice, the kernels are often modeled as tensor products of one-dimensional kernels,

$$k(\mathbf{h}) = \sigma^2 \prod_{\ell=1}^d g_{\ell}(h_{\ell}; \theta_{\ell}), \quad (2.3)$$

where popular one-dimensional covariance functions are the Gaussian,

$$g(h; \theta) = \exp\left(-\frac{h^2}{2\theta^2}\right),$$

and the Matérn 5/2,

$$g(h; \theta) = \left(1 + \frac{\sqrt{5}|h|}{\theta} + \frac{5h^2}{3\theta^2}\right) \exp\left(-\frac{\sqrt{5}|h|}{\theta}\right).$$

An overview of possible covariance functions and their properties can be found in Rasmussen and Williams (2006). In the one-dimensional covariance functions, the parameter to be estimated, θ_i , controls the covariance in the direction of input variable x_i . If all input variables are scaled equally, the estimates of θ_i can be used as a first impression of the influence of the variables. A high covariance indicates a flat curve and thus a small influence of the variable whereas a low covariance indicates a higher fluctuation and a stronger influence.

2.3 Variance-based indices

The most popular global sensitivity analysis method, the so-called Sobol indices, is based on the variance of the function, decomposed into additive terms. One of the first

to mention this decomposition was Hoeffding (1948), who used it to obtain independent random variables to study properties of U-statistics. Efron and Stein (1981) showed the uniqueness of the decomposition and Sobol' (1993) revisited it in the context of sensitivity analysis. A generalization of the decomposition, the so-called high-dimensional model representation (HDMR), was introduced by Rabitz et al. (1999), who aimed at representing functions as a sum of low-dimensional components. Basing on so many (and more) contributors, the decomposition is known under a variety of names. In this work, it will be referred to as *functional ANOVA (FANOVA) decomposition* as it provides an ANOVA decomposition of the variance of the function.

Let us consider a black-box function $Y = f(\mathbf{X})$, where $\mathbf{X} = (X_1, \dots, X_d)'$ is a vector of independent random variables with distribution $\mu = \mu_1 \otimes \dots \otimes \mu_d$, and f is a d -dimensional function $f : \Delta \rightarrow \mathbb{R}$ with $f(\mathbf{X}) \in L^2(\mu)$, where $L^2(\mu)$ denotes the space of square-integrable functions with respect to the measure μ . The function can be decomposed into additive terms,

$$f(\mathbf{X}) = f_0 + \sum_{i=1}^d f_i(X_i) + \sum_{i < j} f_{i,j}(X_i, X_j) + \dots + f_{1,\dots,d}(X_1, \dots, X_d). \quad (2.4)$$

The terms represent first-order effects ($f_i(X_i)$), second-order interactions ($f_{i,j}(X_i, X_j)$), and all higher combinations of input variables. The decomposition is unique if all terms $f_I(\mathbf{X}_I)$, $I \subset \{1, \dots, d\}$ have zero mean,

$$\mathbb{E}(f_I(\mathbf{X}_I)) = 0, \quad I \subseteq \{1, \dots, d\} \quad (2.5)$$

and the conditional expectations fulfill the non-simplification conditions

$$\mathbb{E}(f_I(\mathbf{X}_I) \mid \mathbf{X}_J) = 0, \quad J \subset I \subseteq \{1, \dots, d\}. \quad (2.6)$$

From (2.5) and (2.6), it follows that the terms also have zero correlation,

$$\mathbb{E}(f_I(\mathbf{X}_I)f_{I'}(\mathbf{X}_{I'})) = 0, \quad I \neq I'.$$

The decomposition can be obtained by recursive integration,

$$f_0 = \mathbb{E}(f(\mathbf{X})),$$

$$f_i(X_i) = E(f(\mathbf{X})|X_i) - f_0,$$

$$f_{i,j}(X_i, X_j) = E(f(\mathbf{X})|X_i, X_j) - f_i(X_i) - f_j(X_j) - f_0,$$

and so on. By computing the variance of (2.4), an ANOVA-like variance decomposition is obtained, where each part quantifies the impact of the input variables on the response,

$$\begin{aligned} D = \text{Var}(f(\mathbf{X})) = \text{Var}(f_0) + \sum_{i=1}^d \text{Var}(f_i(X_i)) + \sum_{i < j} \text{Var}(f_{i,j}(X_i, X_j)) \\ + \cdots + \text{Var}(f_{1,\dots,d}(X_1, \dots, X_d)). \end{aligned} \quad (2.7)$$

2.3.1 Definition of various variance-based indices

The variances

$$D_I = \text{Var}(f_I(\mathbf{X}_I)) \quad (2.8)$$

are known as unscaled *Sobol indices* (Sobol', 1993). The first-order Sobol index (for $I \in \{1, \dots, d\}$) is widely used as a sensitivity measure for quantifying the influence of first-order effects. When I contains more than one input variable, the Sobol index quantifies the pure interaction influence of the variables indexed in I . As the estimation of the interactions of all possible variable combinations is usually laborious, other indices presented in the following have been developed for the assessment of interactions.

The Sobol index is often divided by the overall variance D , leading to the *scaled Sobol index*

$$S_I = \frac{\text{Var}(f_I(\mathbf{X}_I))}{D} = \frac{D_I}{D}. \quad (2.9)$$

By this division by the overall variance, the index is normalized to fall between 0 and 1 and thus is easier to assess. The same applies for all indices introduced in the following. For the sake of brevity, they will however be introduced in their unscaled versions since the scaling is equal for all methods considered.

An important extension of the Sobol index is the *total sensitivity index* D_I^T (Homma and Saltelli, 1996). Defined for a group of input variables $\mathbf{X}_I$ for any $I \subseteq \{1, \dots, d\}$, it describes the influence of the variables including all interactions of any order that contain at least one of them. Thus, the influence on all orders instead of only the first is measured, which makes it more useful for the screening of input variables. The total sensitivity index is defined as the sum of all partial variances that contain at least one of the variables,

$$D_I^T = \sum_{J \cap I \neq \emptyset} D_J. \quad (2.10)$$

Another way to describe the influence of a group of input variables is the *closed sensitivity index* D_I^C , see e.g. Sobol' (1993),

$$D_I^C = \sum_{J \subseteq I} D_J = \text{Var}(E[f(\mathbf{X})|\mathbf{X}_I]). \quad (2.11)$$

In contrast to total indices, interactions with variables not in $\mathbf{X}_I$ are not included in this index, but all effects, first-order effects as well as interactions, caused by subsets of it. It is equal to the variance of the conditional expectation and is therefore also known under this name (or shorter VCE) in the literature (McKay, 1995). In its first-order version, it matches the Sobol index.

Sobol indices of higher order can be obtained as a sum of closed sensitivity indices of this order and lower,

$$D_I = \sum_{M=1}^{|I|} (-1)^{|I|-M} \sum_{\substack{J \subseteq I, \\ |J|=M}} D_J^C. \quad (2.12)$$

The total and the closed sensitivity index are related by

$$D = D_{-I}^C + D_I^T, \quad (2.13)$$

where $-I$ denotes the complement of the set I . Through the *formula of total variance*, this leads to a further representation of the total sensitivity index,

$$D_I^T = E(\text{Var}[f(\mathbf{X})|\mathbf{X}_{-I}]).$$

2.3.2 Monte Carlo estimation

As the computer experiments considered here are treated as black-box functions, the different indices cannot be computed directly, but have to be estimated. The most forward way to estimate the indices, that is the expected values resulting from (2.4), is by crude Monte Carlo integration. A high number n – typically not less than $1\,000 \times d$ – of random or quasi-random sampled Monte Carlo runs $\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(n)}$ is drawn from μ , the distribution of $\mathbf{X}$, to approach the integral,

$$\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) \xrightarrow{n \rightarrow \infty} \int f(\mathbf{X}) d\mu(\mathbf{X}) = \mathbb{E}(f(\mathbf{X})). \quad (2.14)$$

The size of the Monte Carlo sample is denoted by n throughout this work, except when its specific role is of interest. The approximation is unbiased and, following the strong law of large numbers, convergent with probability one (Caflisch, 1998).

To estimate the closed sensitivity indices, and via them through (2.12) the Sobol indices, we need a representation of the index that can directly be adapted to Monte Carlo integration. A common way, sometimes called the *pick-and-freeze* (or just *pick-freeze formula*), is given by Sobol' (1993),

$$D_I^C = \mathbb{E}[f(\mathbf{X})f(\mathbf{X}_I, \mathbf{Z}_{-I})] - f_0^2, \quad (2.15)$$

where here as well as in the following, $\mathbf{Z}$ stands for an independent copy of $\mathbf{X}$. The constant term and overall variance can be obtained by

$$f_0 = \mathbb{E}(f(\mathbf{X})) \quad \text{and} \quad D = \text{Var}(f(\mathbf{X})).$$

The corresponding Monte Carlo estimate for two $n \times d$ -Monte Carlo samples $\mathbf{x}$ and $\mathbf{z}$ drawn from μ is straightforward.

$$\text{pf} \widehat{D}_I^C = \frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)}) - \widehat{f}_0^2, \quad (2.16)$$

$$\widehat{f}_0 = \frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}), \quad (2.17)$$

$$\widehat{D} = \widehat{\text{Var}}(f(\mathbf{x})). \quad (2.18)$$

Alternative expressions, which also take the evaluations $f(\mathbf{X}_I, \mathbf{Z}_{-I})$ into account in the estimation of f_0 and D and thus enable greater numerical stability, have been introduced by Monod et al. (2006, p. 86),

$$\begin{aligned} f_0 &= \mathbb{E} \left[\frac{f(\mathbf{X}) + f(\mathbf{X}_I, \mathbf{Z}_{-I})}{2} \right], \\ {}^* \hat{f}_0 &= \frac{1}{n} \sum_{k=1}^n \frac{f(\mathbf{x}^{(k)}) + f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)})}{2}, \end{aligned} \quad (2.19)$$

and

$$\begin{aligned} D &= \mathbb{E} \left[\frac{f(\mathbf{X})^2 + f(\mathbf{X}_I, \mathbf{Z}_{-I})^2}{2} \right] - \left(\mathbb{E} \left[\frac{f(\mathbf{X}) + f(\mathbf{X}_I, \mathbf{Z}_{-I})}{2} \right] \right)^2, \\ {}^* \hat{D} &= \frac{1}{n} \sum_{k=1}^n \frac{f(\mathbf{x}^{(k)})^2 + f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)})^2}{2} - {}^* \hat{f}_0^2. \end{aligned} \quad (2.20)$$

One problem of the pick-freeze estimator (2.16) is that its variance gets very large when f_0^2 is large in comparison to D_I^C . Sobol' (1993) therefore suggests to shift f by an amount close to f_0 . Another possibility is to avoid the subtraction of f_0^2 . In Owen (2013a), two such estimation strategies called *Correlation 1* (Mauntz, 2002, Formula 18) and *Correlation 2* are compared. They are based on the representations

$$\begin{aligned} D_I^C &= \mathbb{E} [f(\mathbf{X}) (f(\mathbf{X}_I, \mathbf{Z}_{-I}) - f(\mathbf{Z}))] \\ &= \mathbb{E} [(f(\mathbf{X}) - f(\mathbf{Z}_I, \mathbf{X}_{-I})) (f(\mathbf{X}_I, \mathbf{U}_{-I}) - f(\mathbf{U}))], \end{aligned}$$

with $\mathbf{U}$ a further independent copy of $\mathbf{X}$. The corresponding estimates, using Monte Carlo samples $\mathbf{x}$, $\mathbf{z}$, and $\mathbf{u}$ from μ , are much more accurate for small closed sensitivity indices.

$$\begin{aligned} {}^{\text{Cor1}} \hat{D}_I^C &= \frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) \left(f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)}) - f(\mathbf{z}^{(k)}) \right), \\ {}^{\text{Cor2}} \hat{D}_I^C &= \frac{1}{n} \sum_{k=1}^n \left(f(\mathbf{x}^{(k)}) - f(\mathbf{z}_I^{(k)}, \mathbf{x}_{-I}^{(k)}) \right) \left(f(\mathbf{x}_I^{(k)}, \mathbf{u}_{-I}^{(k)}) - f(\mathbf{u}^{(k)}) \right). \end{aligned} \quad (2.21)$$

Though they require more — 3 and 4 — vectors of function evaluations in the integral, no additional evaluations are necessary when applying the strategy of simultaneous estimation of all indices by Saltelli (2002), described later in this Section.

index	by Sobol indices	by variances	computation
D (overall variance)	$\sum_J D_J$	$\text{Var}(f(\mathbf{X}))$	$\text{E} \left[f(\mathbf{X})^2 \right] - f_0^2$ $\text{E} \left[\frac{f(\mathbf{X}_I)^2 + f(\mathbf{X}_I, \mathbf{Z}_{-I})^2}{2} \right] - f_0^2$
D_I (Sobol index)	D_I	$\text{Var}(f_I(\mathbf{X}_I))$	$\sum_{M=1}^{ I } (-1)^{ I -M} \sum_{\substack{J \subseteq I, \\ J =M}} D_I^C$
D_I^C (closed sensitivity index)	$\sum_{J \subseteq I} D_J$	$\text{Var}(\text{E}[f(\mathbf{X}) \mathbf{X}_I])$	$\text{E}[f(\mathbf{X})f(\mathbf{X}_I, \mathbf{Z}_{-I})] - f_0^2,$ $\text{E}[f(\mathbf{X})(f(\mathbf{X}_I, \mathbf{Z}_{-I}) - f(\mathbf{Z}))],$ $\text{E}[(f(\mathbf{X}) - f(\mathbf{Z}_I, \mathbf{X}_{-I}))(f(\mathbf{X}_I, \mathbf{U}_{-I}) - f(\mathbf{U}))]$
D_I^T (total sensitivity index)	$\sum_{J \cap I \neq \emptyset} D_J$	$\text{E}(\text{Var}[f(\mathbf{X}) \mathbf{X}_{-I}])$	$D - D_I^C, \frac{1}{2} \text{E}[(f(\mathbf{X}) - f(\mathbf{X}_{-I}, \mathbf{Z}_I))^2]$

Table 2.1: Overview of (unscaled) variance-based indices and their computations.

The total sensitivity index can on the one hand be obtained via the pick-freeze formula using the relation between closed and total sensitivity indices (2.13),

$$\begin{aligned}
 D_I^T &= D - D_{-I}^C \\
 &= D - \text{E}[f(\mathbf{X})f(\mathbf{Z}_I, \mathbf{X}_{-I})] + f_0^2.
 \end{aligned} \tag{2.22}$$

Another way, sometimes called the *Jansen formula*, which avoids the outer sum was first mentioned in Sobol' (1993) and later improved in Jansen (1999),

$$D_I^T = \frac{1}{2} \text{E}[(f(\mathbf{X}) - f(\mathbf{Z}_I, \mathbf{X}_{-I}))^2]. \tag{2.23}$$

The corresponding estimator

$$\text{Jan} \widehat{D}_I^T = \frac{1}{2n} \sum_{k=1}^n \left(f(\mathbf{x}^{(k)}) - f(\mathbf{z}_I^{(k)}, \mathbf{x}_{-I}^{(k)}) \right)^2.$$

is proved to be more efficient than (2.22) in Sobol' (2001, Theorem 4). A proof of its asymptotical efficiency will follow in Chapter 3.

An overview of all presented variance-based indices together with their estimation methods can be found in Tab. 2.1.

Estimation of a full set of indices

Saltelli (2002) presents a strategy to recycle runs when the full set of all d indices shall be estimated simultaneously, which is usually the case in applications. They exploit the fact that the two input data sets $\mathbf{x}$ and $\mathbf{z}$ can be used in more than one way in the estimation of an index. For instance the two pick-freeze formulas

$$\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)})f(\mathbf{x}_{i,j}^{(k)}, \mathbf{z}_{-\{i,j\}}^{(k)}) - \hat{f}_0^2 \quad \text{and} \quad \frac{1}{n} \sum_{k=1}^n f(z_i^{(k)}, \mathbf{x}_{-i}^{(k)})f(x_j^{(k)}, \mathbf{z}_{-j}^{(k)}) - \hat{f}_0^2$$

are both estimators for $D_{i,j}^C$, as they both freeze the variables corresponding to $\{i, j\}$. Only the names are changed from x_i and x_j in the first formula to z_i and x_j in the second one. Saltelli (2002) shows in his Theorem 1 that, using $n(d+2)$ evaluations $f(\mathbf{x}), f(\mathbf{z}), f(x_1, \mathbf{z}_{-1}), \dots, f(x_d, \mathbf{z}_{-d})$, it is possible to compute

- a) each first-order Sobol index $\hat{D}_i = \frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)})f(x_i^{(k)}, \mathbf{z}_{-i}^{(k)}) - \hat{f}_0^2$,
- b) each total sensitivity index $\hat{D}_i^T = \hat{D} - \frac{1}{n} \sum_{k=1}^n f(\mathbf{z}^{(k)})f(x_i^{(k)}, \mathbf{z}_{-i}^{(k)}) + \hat{f}_0^2$,
- c) each $(d-2)$ -order closed effect $\hat{D}_{-\{i,j\}}^C = \frac{1}{n} \sum_{k=1}^n f(x_i^{(k)}, \mathbf{z}_{-i}^{(k)})f(x_j^{(k)}, \mathbf{z}_{-j}^{(k)}) - \hat{f}_0^2$.

D and f_0^2 can be estimated from the evaluations as well, e.g. by (2.19) and (2.20). The method further allows for the estimation of all second-order total sensitivity indices and all the corresponding scaled indices.

If in addition $n \times d$ more evaluations $f(z_1, \mathbf{x}_{-1}), \dots, f(z_d, \mathbf{x}_{-d})$ are available, it is possible to obtain second-order closed sensitivity indices as well as double estimates, which allow for more efficient estimation by taking the mean of both estimates. More precisely, we can compute

- a) an additional estimate of each first-order Sobol index
 $\hat{D}_i = \frac{1}{n} \sum_{k=1}^n f(\mathbf{z}^{(k)})f(z_i^{(k)}, \mathbf{x}_{-i}^{(k)}) - \hat{f}_0^2$,
- b) an additional estimate of each total sensitivity index
 $\hat{D}_i^T = \hat{D} - \frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)})f(z_i^{(k)}, \mathbf{x}_{-i}^{(k)}) + \hat{f}_0^2$,
- c) an additional estimate of each $(d-2)$ -order closed effect
 $\hat{D}_{-\{i,j\}}^C = \frac{1}{n} \sum_{k=1}^n f(z_i^{(k)}, \mathbf{x}_{-i}^{(k)})f(z_j^{(k)}, \mathbf{x}_{-j}^{(k)}) - \hat{f}_0^2$,

d) double estimates of each second-order closed effect

$$\begin{aligned}\widehat{D}_{i,j}^C &= \frac{1}{n} \sum_{k=1}^n f(x_i^{(k)}, \mathbf{z}_{-i}^{(k)}) f(z_j^{(k)}, \mathbf{x}_{-j}^{(k)}) - \widehat{f}_0^2 \text{ and} \\ \widehat{D}_{i,j}^{*C} &= \frac{1}{n} \sum_{k=1}^n f(z_i^{(k)}, \mathbf{x}_{-i}^{(k)}) f(x_j^{(k)}, \mathbf{z}_{-j}^{(k)}) - \widehat{f}_0^2,\end{aligned}$$

which further allows again for the estimation of all scaled indices, second-order total sensitivity indices and second-order Sobol indices. Although not mentioned there, it is obvious that the same applies for the total indices estimated by the Jansen formula instead of pick-freeze, which needs the same function evaluations.

Properties of Monte Carlo estimators

To assess the error coming from the Monte Carlo estimation, a straightforward idea is to use bootstrap (Archer et al., 1997). To do this, in the Monte Carlo estimation of an index of interest, the n summands in formula (2.14) are resampled with replacement for a high number of times ($\approx 10\,000$). Each time, the corresponding index is computed and confidence intervals can be constructed from the resulting bootstrap sample. No new evaluations of f are necessary, as only the already present evaluations are used.

On the other hand, some theoretical results are sometimes available. Janon et al. (2013) discover asymptotic properties of two estimators for the closed sensitivity index (2.11) scaled by the overall variance, $\frac{D^C}{D}$. Both estimators use formula (2.16) for the estimation, but f_0 and the overall variance D are estimated differently. The first estimator, $\widehat{S}_n^1$, uses the straightforward expressions $\widehat{f}_0$ (2.17) and $\widehat{D}$ (2.18),

$$\widehat{S}_n^1 = \frac{\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)}) - \left(\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) \right)^2}{\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)})^2 - \left(\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) \right)^2}, \quad (2.24)$$

whereas in the second, $\widehat{S}_n^2$, the more stable estimators $\widehat{f}_0^*$ (2.19) and $\widehat{D}^*$ (2.20) are applied,

$$\widehat{S}_n^2 = \frac{\frac{1}{n} \sum_{k=1}^n f(\mathbf{x}^{(k)}) f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)}) - \left(\frac{1}{n} \sum_{k=1}^n \frac{f(\mathbf{x}^{(k)}) + f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)})}{2} \right)^2}{\frac{1}{n} \sum_{k=1}^n \frac{f(\mathbf{x}^{(k)})^2 + f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)})^2}{2} - \left(\frac{1}{n} \sum_{k=1}^n \frac{f(\mathbf{x}^{(k)}) + f(\mathbf{x}_I^{(k)}, \mathbf{z}_{-I}^{(k)})}{2} \right)^2}. \quad (2.25)$$

Janon et al. (2013) show that both estimators are consistent,

$$\widehat{S}_n^1 \xrightarrow[n \rightarrow \infty]{a.s.} \frac{D_I^C}{D}, \quad \widehat{S}_n^2 \xrightarrow[n \rightarrow \infty]{a.s.} \frac{D_I^C}{D},$$

and that they are asymptotically normal in the form of

$$\sqrt{n} \left(\widehat{S}_n^1 - \frac{D_I^C}{D} \right) \xrightarrow[n \rightarrow \infty]{d} \mathcal{N}(0, \sigma_1^2), \quad \sqrt{n} \left(\widehat{S}_n^2 - \frac{D_I^C}{D} \right) \xrightarrow[n \rightarrow \infty]{d} \mathcal{N}(0, \sigma_2^2).$$

See Janon et al. (2013, Prop. 3.3) for the specification of the variances σ_1^2 and σ_2^2 .

In their main result, they show that $\widehat{S}_n^2$ is asymptotically efficient for the estimation of D_I^C/D in the notion of van der Vaart (1998, Chapter 25) in a framework with exchangeable variables.

2.3.3 Frequency-based estimation

Beside direct Monte Carlo estimation, there is a further class of estimation methods, the frequency-based estimation, which is computationally cheaper than Monte Carlo estimation, but also slightly biased.

The first method introduced is the Fourier amplitude sensitivity test (FAST) by Cukier et al. (1978), which allows for the estimation of first-order Sobol indices. Sample points of $\mathbf{X}$ are chosen such that the indices can be interpreted as amplitudes obtained by Fourier analysis of the function. More precisely, the design of n_{FAST} points is such that

$$x_i^{(k)} = G_i(\sin(\omega_i s_k)), \quad i = 1, \dots, d, \quad k = 1, \dots, n_{\text{FAST}}, \quad s_k = \frac{2\pi(k-1)}{n_{\text{FAST}}},$$

with G_i functions to ensure that the sample points follow the distribution of $\mathbf{X}$. The set of integer frequencies $\{\omega_1, \dots, \omega_d\}$, associated to the input variables, is chosen to be as free of interferences as possible, where *free of interferences up to the order M* means that $\sum_{i=1}^p a_i \omega_i \neq 0$ for $\sum_{i=1}^p |a_i| \leq M+1$ (Tissot and Prieur, 2012). In practice $M = 4$ or $M = 6$ are used.

For a weight ω , the Fourier coefficients for each input variable can then be numerically estimated by

$$\begin{aligned}\hat{A}_\omega &= \frac{1}{n_{\text{FAST}}} \sum_{j=1}^{n_{\text{FAST}}} f(x(s_j)) \cos(\omega s_j), \\ \hat{B}_\omega &= \frac{1}{n_{\text{FAST}}} \sum_{j=1}^{n_{\text{FAST}}} f(x(s_j)) \sin(\omega s_j),\end{aligned}\tag{2.26}$$

and the first-order Sobol indices can be estimated by the sum of the corresponding amplitudes up to the order M ,

$$^{\text{FAST}}\hat{D}_i = 2 \sum_{p=1}^M (\hat{A}_{p\omega_i}^2 + \hat{B}_{p\omega_i}^2).\tag{2.27}$$

An estimate of the overall variance is given by the sum of all amplitudes,

$$^{\text{FAST}}\hat{D} = 2 \sum_{n=1}^{n_{\text{FAST}}/2} (\hat{A}_n^2 + \hat{B}_n^2).\tag{2.28}$$

Extended FAST (EFAST), first presented in Saltelli et al. (1999), is a method to also compute total sensitivity indices for single input variables using FAST. Here, a high frequency ω_i is assigned to the variable in question X_i and low frequencies, e.g. $\omega_{-i} = 1$, are assigned to all other variables. Then the total sensitivity index of X_i is estimated over the amplitude for the low frequencies,

$$^{\text{EFAST}}\hat{D}_i^T = D - 2 \sum_{p=1}^M (\hat{A}_{p\omega_{-i}}^2 + \hat{B}_{p\omega_{-i}}^2).$$

A similar method, which avoids the problem of interferences, is RBD-FAST, where the RBD stands for random balanced design. RBD-FAST is a group of modifications of FAST, which use random permutations of design points to avoid interferences. The original idea by Tarantola et al. (2006) is to assign the same frequency to all input variables, but to randomize their values independently before evaluating f . The first-order Sobol index of X_i can then be estimated by reordering the evaluations in the way X_i was permuted, so that the amplitude at the frequency returns the sensitivity of X_i only. This works because the order of the evaluations is important in (2.26). Thus, all first-order Sobol indices can be estimated simultaneously with one set of evaluations. In the same work a further RBD-FAST method called Hybrid is mentioned by which second-order and second-order closed sensitivity indices can be estimated.

Mara (2009) suggests a RBD-FAST method to compute the total sensitivity index of a group of variables $\widehat{D}_I^T$. Here, simple frequencies like $\omega = \{1, \dots, d\}$ are assigned to the variables. Then $n_{\text{RBD}} = 2(Md + L)$ design points are generated over a periodic curve, with $L > 100$ a tunable integer number, which regulates the sample size. The values of all variables in I are then randomly permuted, either differently per variable or identically, and the experiment is evaluated at the points. The total sensitivity index is estimated by

$$^{\text{RBD}}\widehat{D}_I^T = \frac{n_{\text{RBD}}}{L} \sum_{p=dM+1}^{n_{\text{RBD}}/2} (\widehat{A}_p^2 + \widehat{B}_p^2). \quad (2.29)$$

The EASI algorithm by Plischke (2010) is a further frequency-based method for the estimation of first-order Sobol indices. It has the advantage that no specific sampling scheme is necessary, which makes it possible to use already available data. It can be seen as a reverse RBD-FAST approach. The evaluations are reordered so that the x_i are approximately triangular-shaped, which corresponds to the frequency $\omega = 1$. Then the first-order Sobol index can be estimated as for RBD-FAST.

2.4 Derivative-based methods

This section shows different sensitivity analysis methods that, in contrast to variance-based methods, use the partial derivative of an input variable as an indicator of its influence. As the derivative quantifies the change in the output when the variable is changed, it can be used as a sensitivity measure.

2.4.1 Morris screening by elementary effects

In the field of input variable screening in computer experiments, that is finding the important variables out of a large number of variables using a minimal number of runs, Morris screening is one of the most popular methods (Morris, 1991). Among the advantages of the method over common screening methods are that it covers the entire input space and that it is not restricted to linear influences. Partial differences at different points in the input space, so-called elementary effects, are computed

$$d_i(\mathbf{x}) = \frac{|y(x_1, \dots, x_i + \Delta_p, \dots, x_k) - y(\mathbf{x})|}{\Delta_p},$$

where the values of $\mathbf{x}$ are sampled on a grid with p levels for each input and Δ_p is a predetermined multiple of $1/(p - 1)$. The distribution of the d_i provides information on the behavior of X_i . The mean, the so-called Morris Importance Measure, indicates the overall influence. The standard deviation indicates the linearity of the influence. A value close to 0 suggests linear behavior, a high value nonlinear or interaction behavior. To obtain the distribution of elementary effects, the values of $\mathbf{x}$ are sampled randomly or by a more elaborate design that uses trajectories in the space of the input variables, starting from random points on the grid. This reduces the number of necessary runs almost by half compared to random sampling, resulting in a total number of $r(d + 1)$ runs, with r the number of elementary effects per variable.

2.4.2 DGSM

A more accurate method based on Monte Carlo integration was introduced by Kucherenko et al. (2009). They use indices they call *derivate-based global sensitivity measures* (DGSM), first presented by Sobol' and Gershman (1995) as an interpretation of variance-based indices. In contrast to Morris screening, the square of the partial derivative is used,

$$\nu_i = \int \left(\frac{\partial f(\mathbf{X})}{\partial X_i} \right)^2 d\mu(\mathbf{X}). \quad (2.30)$$

Like variance-based indices, the DGSM quantify the influence of the input variables and they share with the total sensitivity indices the ability that, (under the mild assumption that μ is continuous and its support is equal to the domain Δ of X), if $\nu_i = 0$, then $f(\mathbf{x})$ does not depend on x_i . Indeed, Sobol' and Kucherenko (2009) could show a connection between the two sensitivity methods for the uniform and normal distributions, which was extended by Lamboni et al. (2013) for general continuous distributions. It states that

$$D_i^T \leq C(\mu_i) \nu_i, \quad (2.31)$$

provided that μ belongs to a class of distributions that satisfy a Poincaré inequality,

$$\int g(\mathbf{X})^2 d\mu(\mathbf{X}) \leq C(\mu) \int \|\nabla g(\mathbf{X})\|^2 d\mu(\mathbf{X}), \quad (2.32)$$

for all functions g in $L^2(\mu)$ such that $\int g(\mathbf{X}) d\mu(\mathbf{X}) = 0$, and $\|\nabla g\| \in L^2(\mu)$. A *Poincaré constant* $C(\cdot)$ is a characteristic of the corresponding distribution μ . The best possible constant for a given distribution, the optimal Poincaré constant, is denoted by $C_{opt}(\mu)$. When $C(\mu) = C_{opt}(\mu)$, there can exist certain functions f_{opt} for which the Poincaré inequality is an equality. See Lamboni et al. (2013) or Roustant et al. (2014) for a comprehensive summary.

Formula (2.31) implies that the DGSM can be used as upper boundaries for the total sensitivity indices and thus can serve as a cheaper method for the screening of non-influential input variables since, after Lamboni et al. (2013), the DGSM seems to be

computationally more tractable than variance-based measures, especially for problems with higher numbers of inputs. This applies especially when the computer experiment additionally provides the partial derivatives. Then, the Monte Carlo integration can be performed directly. If not, the derivatives can be approximated numerically by finite differences. The DGSM for a $n \times d$ data set $\mathbf{x}$ and a small real number δ^* can be estimated by

$$\hat{\nu}_i = \frac{1}{n} \sum_{k=1}^n \left(\frac{f(x_i^{(k)} + \delta^*, \mathbf{x}_{-\{i\}}^{(k)}) - f(\mathbf{x}^{(k)})}{\delta^*} \right)^2.$$

To explore the difference between the DGSM and variance-based indices, let us compare two one-dimensional case functions,

$$\begin{aligned} f_1(x) &= \frac{1}{1 + 49 \exp(-2 \log(49)x)}, \\ f_2(x) &= \frac{1}{2} \sin(50x) + \frac{1}{2}. \end{aligned} \tag{2.33}$$

Plots are shown in Fig. 2.1. In this one-dimensional situation, the total index, the main index, and the overall variance coincide, the variance-based index reduces to

$$D_1 = D_1^T = D = \int_0^1 f(x)^2 dx - \left(\int_0^1 f(x) dx \right)^2,$$

and for the DGSM we have

$$\nu_1 = \int_0^1 (f'(x))^2 dx.$$

When computing these values for the two functions, we see that the DGSM is much higher for function f_2 (≈ 310.9) than for f_1 (≈ 1.3), as it visibly varies much more. Nevertheless, the overall variance is for both functions almost the same (≈ 0.126), as the function evaluations over the space of x share a similar variance. The variance-based index ignores the local variance. This highlights the difference of the notion *sensitivity* for variance-based indices and DGSM. In an hypothetical underlying model $y = f_1(x_1) + f_2(x_2)$, the DGSM would rank the second variable much higher whereas the

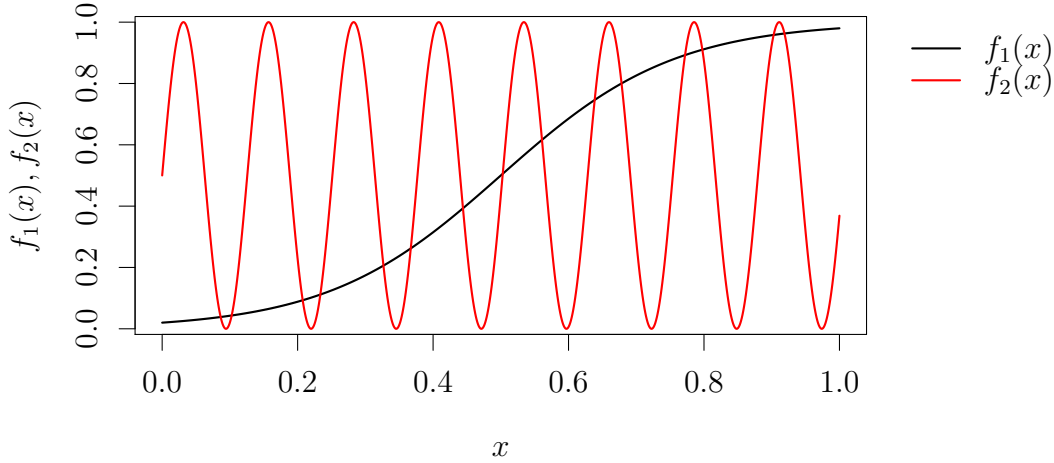


Figure 2.1: Plot of f_1 and f_2 of Equation (2.33).

variance-based indices would see both variables as equally important. The variance-based indices return the variance of the output evoked by the input, the DGSM summarize the local variation.

2.5 Ongoing developments in sensitivity analysis

To complete the overview of methods in sensitivity analysis, we want to show some interesting topics of current research in order to give a little insight into the progression of the field.

Owen (2013b) creates a general framework for variance-based indices by using linear combinations of Sobol indices. The various possible estimators for the different indices are summarized and categorized, which leads to some new estimators. Among others, it supplies a bias-corrected pick-freeze estimator and estimators for Sobol indices with a reduced number of necessary function evaluations.

For the case of multivariate output, several ideas exist that incorporate correlations between output variables to global sensitivity analysis. Indices were introduced by Lamboni et al. (2011) and studied further by Gamboa et al. (2013) that extend sensitivity indices to multivariate output, which – relating to the FANOVA decomposition – base on the decomposition of the covariance of the output. An overview of different approaches to the analysis of multivariate output can be found in Garcia-Cabrejo and Valocchi (2014).

Control variates is a technique to reduce the variance in Monte Carlo integration. Applying control variates to the estimation of total sensitivity indices via the Jansen formula, Kucherenko et al. (2014) find the formula

$$D_i^T = \frac{1}{2} \mathbb{E} [f(\mathbf{X}) - f_i(X_i) - (f(Z_i, \mathbf{X}_{-i}) - f_i(Z_i))]^2 + D_i,$$

which can indeed improve the estimators efficiency. It requires the FANOVA decomposition term f_i as well as the first-order Sobol index D_i . If not known analytically, these terms have to be extracted from metamodels.

A further interesting idea is an extension to Sobol indices that takes the goal of further analysis into consideration (Fort et al., 2014). The idea is that the estimation of a mean or a median could involve different input variables than the estimation of extreme quantiles. They set up a framework they call *Goal Oriented Sensitivity Analysis*, where new sensitivity indices for each statistical purpose are defined by applying contrast functions. When considering the estimation of α -quantiles, a function of indices over α is returned.

This chapter presented the basic sensitivity analysis techniques used for the exploration of input variables in computer experiments. The following chapter will present and examine methods that focus on the exploration of interactions between variables.

3. SENSITIVITY ANALYSIS FOR INTERACTION SCREENING

The *total interaction index* (TII) is a sensitivity index for interaction screening in computer experiments, which can further be used for visualization and block-additive decomposition of the computer model. As the total sensitivity index, introduced in Chapter 2, is ideal for input screening, since it sums up the total influence of the input variables, the screening of interactions can be done by the TII. It contains, for any pair of input variables, the influence of the interaction of the pair plus all interactions containing the pair. Thus, corresponding to the total sensitivity index where a value close to zero is a strong indicator that a variable can be excluded (Saltelli et al., 2006), a TII close to zero indicates that the two variables do not interact. It has been used but not investigated closer in Mühlenstädt et al. (2012) in order to set up so-called FANOVA graphs and Kriging models with block-additive kernels. Liu and Owen (2006) introduced a more general index for uniform distributions as a measure of importance of interactions, which was also used in Hooker (2004) in the data mining framework.

This chapter starts with an example that gives a direct motivation to the TII. Then, several variance-based estimation methods for the TII are introduced. Their properties are analyzed theoretically as well as on simulations. The problem of finding a threshold is addressed as well as the implementation within the statistical software programming

environment R (R Core Team, 2014). In addition, *crossed DGSM*, extensions of the DGSM to second-order interactions, are presented, which provide upper bounds for the TII.

Some results from this chapter are published in the contributions “Total interaction index: a variance-based sensitivity index for second-order interaction screening” (Fruth et al., 2014b) and “Crossed-derivative based sensitivity measures for interaction screening” (Roustant et al., 2014).

Motivational example

With the TII, it is possible to discover the block-additive structure of the underlying function, that is we can identify a decomposition into groups of input variables such that variables between groups do not interact. As an illustration, we analyze the following function,

$$f(X_1, \dots, X_6) = \cos([1, X_1, X_3, X_5] \alpha) + \sin([1, X_2, X_4, X_6] \gamma),$$

with $X_i \stackrel{\text{i.i.d.}}{\sim} U[-1, 1]$, $i = 1, \dots, 6$, $\alpha = [-0.8, -1.1, 1, 1.1]'$ and $\gamma = [-0.51, 0.9, -1.1]'$, where the prime $'$ stands for the transpose.

Clearly, the variables in the group $\{X_1, X_3, X_5\}$ do not interact with the variables in the group $\{X_2, X_4, X_6\}$, which induces a block-additive structure of the form $f(x) = f_{135}(x_1, x_3, x_5) + f_{246}(x_2, x_4, x_6)$. This, however, cannot be seen from the common first-order and total sensitivity indices (Fig. 3.1, left). We only see that all variables have first-order and total effects, but do not know how variables interact with each other.

If we now estimate the TII for each combination of input variables, we obtain such information, which can be conveniently plotted in a so-called FANOVA graph in Fig. 3.1 on the right (Mühlenstädt et al., 2012). The TII of the two variables is represented by

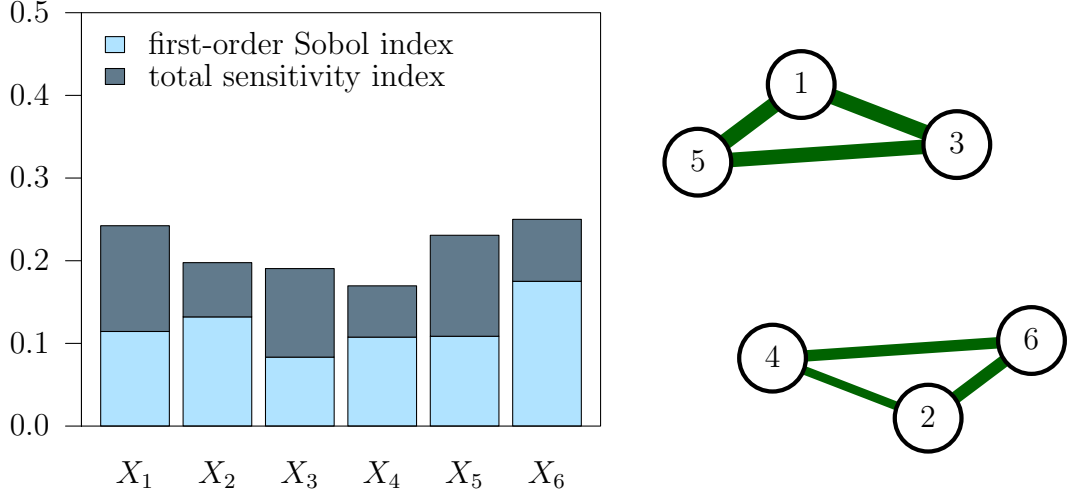


Figure 3.1: Sensitivity analysis of the motivational example. First-order Sobol and total sensitivity indices (left), FANOVA graph of total interaction indices (right).

the thickness of the edge between the two corresponding vertices. With the graph, the partition into the two additive groups is clearly visible.

This information about the block-additive interaction structure of a function can be exploited in several fields of computer experiment analysis. Mühlenstädt et al. (2012) show that it can be used to improve Kriging model predictions by adapting the Kriging kernel to the block-additive structure. For the example above that would mean to modify the kernel k from Equation (2.3) as follows,

$$k(\mathbf{h}) = k_1(h_2, h_4, h_6) + k_2(h_1, h_3, h_5).$$

Furthermore, the block-additive structure can be exploited to simplify and parallelize optimization, as described in Ivanov and Kuhnt (2014). There, the idea is to optimize the separate groups independently, which reduces the optimization dimensions and enables parallelization, an important property as optimization techniques are usually sequential. For the example above, the six-dimensional optimization problem of minimizing f simplifies into two three-dimensional ones, where the $c_1, \dots, c_6$ are constants,

which are not varied in the optimization,

$$\min_{x_1, \dots, x_6} f(x_1, \dots, x_6) = \min_{x_1, x_3, x_5} f(x_1, c_2, x_3, c_4, x_5, c_6) + \min_{x_2, x_4, x_6} f(c_1, x_2, c_3, x_4, c_5, x_6).$$

3.1 Total interaction indices

The TII measures the total influence of a second-order interaction between two input variables. It contains, for any pair of input variables, the influence of their interaction plus all interactions containing both indices. This is different from the second-order Sobol index (2.8), which does not contain higher interactions, and from the second-order total sensitivity index (2.10), which additionally contains first-order effects as well as higher-order interactions of only one of the two variables.

Definition 3.1. *The total interaction index (TII) of two input variables X_i and X_j is defined by*

$$\mathfrak{D}_{i,j} = \text{Var} \left(\sum_{I \supseteq \{i,j\}} f_I(\mathbf{X}_I) \right) = \sum_{I \supseteq \{i,j\}} D_I. \quad (3.1)$$

The TII is a special second-order version of the more general superset importance, which was introduced by Liu and Owen (2006) for uniform distributions (Υ_u^2) as a measure of importance of interactions and their supersets. It was also investigated by Hooker (2004) in the data mining framework ($\bar{\sigma}_u^2$). The superset importance Υ_u^2 is defined for any subset $u \subseteq \{1, \dots, d\}$ as

$$\Upsilon_u^2 = \bar{\sigma}_u^2 = \sum_{I \supseteq u} D_I. \quad (3.2)$$

One important property of the TII is its link to total sensitivity indices and closed indices respectively.

Proposition 3.1.

$$\mathfrak{D}_{i,j} = D_i^T + D_j^T - D_{i,j}^T, \quad (3.3)$$

$$\mathfrak{D}_{i,j} = D + D_{-\{i,j\}}^C - D_{-i}^C - D_{-j}^C. \quad (3.4)$$

Proof. The proof of (3.3) is obtained by simple set transformation. Since

$$\sum_{I \supseteq \{i\} \vee I \supseteq \{j\}} D_I = \sum_{I \supseteq \{i\}} D_I + \sum_{I \supseteq \{j\}} D_I - \sum_{I \supseteq \{i,j\}} D_I,$$

it holds that

$$\sum_{I \supseteq \{i,j\}} D_I = \sum_{I \supseteq \{i\}} D_I + \sum_{I \supseteq \{j\}} D_I - \sum_{I \supseteq \{i\} \vee I \supseteq \{j\}} D_I.$$

Equation (3.4) then can be deduced from (3.3) using the connection between total and closed indices (2.13),

$$\begin{aligned} \mathfrak{D}_{i,j} &= D_i^T + D_j^T - D_{i,j}^T \\ &= (D - D_{-i}^C) + (D - D_{-j}^C) - (D - D_{-\{i,j\}}^C) \\ &= D + D_{-\{i,j\}}^C - D_{-i}^C - D_{-j}^C. \end{aligned}$$

□

From the connection to the superset importance, another representation of the TII is given by Liu and Owen (2006).

Proposition 3.2. *Liu and Owen's formula*

$$\mathfrak{D}_{i,j} = \frac{1}{4} E \left[\left(f(X_i, X_j, \mathbf{X}_{-\{i,j\}}) - f(X_i, Z_j, \mathbf{X}_{-\{i,j\}}) \right. \right. \\ \left. \left. - f(Z_i, X_j, \mathbf{X}_{-\{i,j\}}) + f(Z_i, Z_j, \mathbf{X}_{-\{i,j\}}) \right)^2 \right], \quad (3.5)$$

where Z_i (resp. Z_j) is an independent copy of X_i (resp. X_j).

Proof. A proof is given in Liu and Owen (2006). In addition, (3.5) can be connected to (3.4). When expanding the squared sum in (3.5), of the 10 resulting terms the four squared terms are simply equal to $E(f(\mathbf{X})^2) = D + f_0^2$. The six double products gather two by two, and result in the pick-freeze formula (2.15)

- $E[f(Z_i, X_j, \mathbf{X}_{-\{i,j\}})f(Z_i, Z_j, \mathbf{X}_{-\{i,j\}})]$
 $= E[f(X_i, X_j, \mathbf{X}_{-\{i,j\}})f(X_i, Z_j, \mathbf{X}_{-\{i,j\}})] = D_{-j}^C + f_0^2$
- $E[f(X_i, Z_j, \mathbf{X}_{-\{i,j\}})f(Z_i, Z_j, \mathbf{X}_{-\{i,j\}})]$
 $= E[f(X_i, X_j, \mathbf{X}_{-\{i,j\}})f(Z_i, X_j, \mathbf{X}_{-\{i,j\}})] = D_{-i}^C + f_0^2$
- $E[f(Z_i, X_j, \mathbf{X}_{-\{i,j\}})f(X_i, Z_j, \mathbf{X}_{-\{i,j\}})]$
 $= E[f(X_i, X_j, \mathbf{X}_{-\{i,j\}})f(Z_i, Z_j, \mathbf{X}_{-\{i,j\}})] = D_{-\{i,j\}}^C + f_0^2$

Finally, combining all the terms gives (3.4). □

Remark: Connection to ANOVA Liu and Owen (2006) state that formula (3.5) “can also be obtained through the classical formulas for expected mean squares in the discrete ANOVA as established by Cornfield and Tukey (1956)”. This connection becomes evident when looking at the formula for the interaction Sum of Squares in an ANOVA with two input variables A and B (Cornfield and Tukey, 1956, p. 916),

$$SS_{AB} = n \sum_{i=1}^a \sum_{j=1}^b (\bar{y}_{ij\cdot} - \bar{y}_{i\cdot\cdot} - \bar{y}_{\cdot j\cdot} + \bar{y}_{\cdot\cdot\cdot})^2,$$

with a and b the number of levels and y_{ijk} , $i = 1, \dots, a$, $j = 1, \dots, b$, $k = 1, \dots, n$ the realized output values with n the number of repetitions.

For a further TII computation notion, say we keep all variables other than X_i and X_j fixed and look at the Liu and Owen formula for the resulting two-dimensional function $f_{\text{fixed}}(X_i, X_j)$,

$$\mathfrak{D}_{i,j} = \frac{1}{4} E \left[(f_{\text{fixed}}(X_i, X_j) - f_{\text{fixed}}(X_i, Z_j) - f_{\text{fixed}}(Z_i, X_j) + f_{\text{fixed}}(Z_i, Z_j))^2 \right].$$

Then we see that the TII of the original function can be computed as the expected value over this second-order interaction index of the two-dimensional function obtained by fixing all variables except X_i and X_j . Thus, the Liu and Owen formula is a so-called *fixing method*, introduced by Mühlenstädt et al. (2012).

Proposition 3.3. (*Fixing method*). *For any $x_{-\{i,j\}}$, define f_{fixed} as the two-dimensional function $f_{\text{fixed}} : (x_i, x_j) \rightarrow f(\mathbf{x})$ obtained from f by fixing all input variables except x_i and x_j . Let $D_{i,j|\mathbf{X}_{-\{i,j\}}}$ denote the second-order Sobol index of $f_{\text{fixed}}(X_i, X_j)$, which depends on the fixed variables $\mathbf{X}_{-\{i,j\}}$. Then the TII of X_i and X_j is obtained by integrating $D_{i,j|\mathbf{X}_{-\{i,j\}}}$ with respect to $\mathbf{X}_{-\{i,j\}}$,*

$$\mathfrak{D}_{i,j} = E\left(D_{i,j|\mathbf{X}_{-\{i,j\}}}\right). \quad (3.6)$$

Proof. Since the function f_{fixed} is two-dimensional, it has only one interaction, which is a second-order one, and coincides with its TII. Hence, this interaction can be computed by applying (3.5) to f_{fixed} ,

$$D_{i,j|\mathbf{X}_{-\{i,j\}}} = \frac{1}{4} E[f_{\text{fixed}}(X_i, X_j) - f_{\text{fixed}}(X_i, Z_j) - f_{\text{fixed}}(Z_i, X_j) + f_{\text{fixed}}(Z_i, Z_j)]^2.$$

Now we can rewrite the right hand side by using conditional expectations,

$$D_{i,j|\mathbf{X}_{-\{i,j\}}} = \frac{1}{4} E \left[[f(X_i, X_j, \mathbf{X}_{-\{i,j\}}) - f(X_i, Z_j, \mathbf{X}_{-\{i,j\}}) - f(Z_i, X_j, \mathbf{X}_{-\{i,j\}}) + f(Z_i, Z_j, \mathbf{X}_{-\{i,j\}})]^2 | \mathbf{X}_{-\{i,j\}} \right].$$

Taking the expectation with respect to $\mathbf{X}_{-\{i,j\}}$ gives the result. $\square$

3.2 TII estimation

The different representations of the TII derived in the last section and the estimation methods of the common indices from Section 2.3 are now combined to construct estimation methods for the TII, followed by remarks on their properties.

Estimation via total sensitivity indices

To apply the link between the TII and total sensitivity indices from Prop. 3.1, formula (3.3), estimators for the total sensitivity index of first- and second-order are required. Thus, either the Monte Carlo estimator via the Jansen formula (2.23) or the RBD-FAST method (2.29) can be used, resulting in the TII estimators *Jansen TII estimator* and *RBD-FAST TII estimator*,

$${}^{\text{Jan}}\widehat{\mathfrak{D}}_{i,j} = {}^{\text{Jan}}\widehat{D}_i^T + {}^{\text{Jan}}\widehat{D}_j^T - {}^{\text{Jan}}\widehat{D}_{\{i,j\}}^T \quad (3.7)$$

$${}^{\text{RBD}}\widehat{\mathfrak{D}}_{i,j} = {}^{\text{RBD}}\widehat{D}_i^T + {}^{\text{RBD}}\widehat{D}_j^T - {}^{\text{RBD}}\widehat{D}_{\{i,j\}}^T \quad (3.8)$$

For the *Jansen TII estimator*, evaluations can be saved by computing the last term ${}^{\text{Jan}}\widehat{D}_{\{i,j\}}^T$ via $f(\mathbf{x}_{-i}, \mathbf{z}_i)$ and $f(\mathbf{x}_{-j}, \mathbf{z}_j)$, which have been evaluated in the computation of the previous terms, rather than via $f(\mathbf{x})$ and $f(\mathbf{x}_{-\{i,j\}}, \mathbf{z}_{\{i,j\}})$. The approach is discussed further in Section 3.4.

Estimation via closed sensitivity indices

The second link from Prop. 3.1 between the TII and the closed sensitivity indices (3.4) requires the estimation of high-order closed indices D_{-i}^C and $D_{i,j}^C$. To this point, no reliable FAST or RBD-FAST method exists, so Monte Carlo integration is considered. As the required closed indices are of high order, they are expected to be large. Thus, the estimator based on the pick-freeze formula (2.15, 2.16) is more suitable here than the one based on the strategies from *Correlation 1* and *Correlation 2* (2.21) since they are not recommended when the index to be estimated is rather large (Owen, 2013a). The *pick-freeze TII estimator* of the TII is constructed as

$${}^{\text{pf}}\widehat{\mathfrak{D}}_{i,j} = \widehat{D} + {}^{\text{pf}}\widehat{D}_{-\{i,j\}}^C - {}^{\text{pf}}\widehat{D}_{-i}^C - {}^{\text{pf}}\widehat{D}_{-j}^C. \quad (3.9)$$

As for the *Jansen TII estimator*, evaluations can be reused in the second-order term by using $f(\mathbf{x}_{-i}, z_i)$ and $f(\mathbf{x}_{-j}, z_j)$ rather than $f(\mathbf{x})$ and $f(\mathbf{x}_{-\{i,j\}}, \mathbf{z}_{\{i,j\}})$.

Estimation via Liu and Owen's formula

Another TII estimation method is directly obtained by Liu and Owen's formula in Prop. 3.2, the *Liu and Owen TII estimator*

$$\begin{aligned} {}^{\text{LO}}\widehat{\mathfrak{D}}_{i,j} = \frac{1}{4} \times \frac{1}{n} \sum_{k=1}^n [f(x_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) - f(x_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k) \\ - f(z_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) + f(z_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k)]^2. \end{aligned} \quad (3.10)$$

Estimation via fixing method

Following Prop. 3.3, the TII can be computed by averaging the second-order interaction of two-dimensional functions in the way it has been done in Mühlenstädt et al. (2012) by the following scheme:

Let us consider a couple of integers (i, j) , with $i < j$. For $k = 1, \dots, n_{\text{MC}}$ carry out the following steps.

1. Simulate $\mathbf{x}_{-\{i,j\}}^k$ from the distribution of $\mathbf{X}_{-\{i,j\}}$, that is take a single sample of all input variables except X_i and X_j .
2. Create the two-dimensional function f_{fixed} by fixing f on $\mathbf{x}_{-\{i,j\}}^k$

$$f_{\text{fixed}}(X_i, X_j) = f(x_1^k, \dots, X_i, \dots, X_j, \dots, x_d^k).$$

3. Using the FAST estimator (Section 2.3.3), compute the second-order Sobol index of f_{fixed} , denoted $\widehat{D}_{i,j|\mathbf{X}_{-\{i,j\}}}^k$, by

$$\widehat{D}_{i,j|\mathbf{X}_{-\{i,j\}}}^k = {}^{\text{FAST}}\widehat{D}_{|\mathbf{X}_{-\{i,j\}}}^k - {}^{\text{FAST}}\widehat{D}_{i|\mathbf{X}_{-\{i,j\}}}^k - {}^{\text{FAST}}\widehat{D}_{j|\mathbf{X}_{-\{i,j\}}}^k,$$

TII estimator	derived from	formula
<i>Jansen</i>	Prop. 3.1 (3.3)	${}^{\text{Jan}}\widehat{\mathfrak{D}}_{i,j} = {}^{\text{Jan}}\widehat{D}_i^T + {}^{\text{Jan}}\widehat{D}_j^T - {}^{\text{Jan}}\widehat{D}_{\{i,j\}}^T$
<i>RBD-FAST</i>	Prop. 3.1 (3.3)	${}^{\text{RBD}}\widehat{\mathfrak{D}}_{i,j} = {}^{\text{RBD}}\widehat{D}_i^T + {}^{\text{RBD}}\widehat{D}_j^T - {}^{\text{RBD}}\widehat{D}_{\{i,j\}}^T$
<i>pick-freeze</i>	Prop. 3.1 (3.4)	${}^{\text{pf}}\widehat{\mathfrak{D}}_{i,j} = {}^{\text{pf}}\widehat{D} + {}^{\text{pf}}\widehat{D}_{-\{i,j\}}^C - {}^{\text{pf}}\widehat{D}_{-i}^C - {}^{\text{pf}}\widehat{D}_{-j}^C$
<i>Liu and Owen</i>	Prop. 3.2	${}^{\text{LO}}\widehat{\mathfrak{D}}_{i,j} = \frac{1}{4} \times \frac{1}{n} \sum_{k=1}^n \left[f(x_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) - f(x_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k) \right. \\ \left. - f(z_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) + f(z_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k) \right]^2$
<i>fixing method</i>	Prop. 3.3	${}^{\text{fix}}\widehat{\mathfrak{D}}_{i,j} = \frac{1}{n_{\text{MC}}} \sum_{k=1}^{n_{\text{MC}}} \widehat{D}_{i,j \mathbf{X}_{-\{i,j\}}}^k$

Table 3.1: Overview of considered TII estimators.

where $\widehat{D}_{|\mathbf{X}_{-\{i,j\}}}^k$ denotes the overall variance of f_{fixed} and $\widehat{D}_{i|\mathbf{X}_{-\{i,j\}}}^k$ and $\widehat{D}_{j|\mathbf{X}_{-\{i,j\}}}^k$ the first-order indices of X_i and X_j , respectively.

After carrying out the three points for all $k = 1, \dots, n_{\text{MC}}$, compute the *fixing method TII estimator* by

$${}^{\text{fix}}\widehat{\mathfrak{D}}_{i,j} = \frac{1}{n_{\text{MC}}} \sum_{k=1}^{n_{\text{MC}}} \widehat{D}_{i,j|\mathbf{X}_{-\{i,j\}}}^k. \quad (3.11)$$

A summary of all considered TII estimators can be found in Tab. 3.1.

3.3 Theoretical properties of TII estimators

In the following, the introduced TII estimators are compared according to different statistical properties.

Nonnegativity

As the indices measure the sum of variances, the estimates should be nonnegative. This holds for the *Liu and Owen TII estimator* (3.10) which is a sum of squares. However, negative estimates can occur for the *Jansen* (3.7), the *pick-freeze* (3.9) and the *RBD-FAST TII estimator* (3.8). For the *fixing method TII estimator* (3.11), there is a sufficient condition, which results from the following proposition:

Proposition 3.4. *Let f be a two-dimensional function, and consider its second-order interaction $D_{12} = D - D_1 - D_2$. Denote by $\hat{D}_{12} = {}^{FAST}\hat{D} - {}^{FAST}\hat{D}_1 - {}^{FAST}\hat{D}_2$ its FAST estimate (see Section 2.3.3) and assume that:*

- (i) ω_1 and ω_2 are free of interference up to order $2M$,
- (ii) $N \geq 2M \times \max(\omega_1, \omega_2)$.

Then $\hat{D}_{12} \geq 0$.

Proof. Denote the sets $W_{\omega_i, M} = \{p\omega_i, p = 1, \dots, M\}$ for $i = 1, 2$, and $W_N = \{1, \dots, N/2\}$. With (2.27) and (2.28), we have

$$\hat{D}_{12}/2 = \sum_{n \in W_N} (\hat{A}_n^2 + \hat{B}_n^2) - \sum_{n \in W_{\omega_1, M}} (\hat{A}_n^2 + \hat{B}_n^2) - \sum_{n \in W_{\omega_2, M}} (\hat{A}_n^2 + \hat{B}_n^2).$$

Now, the condition (i) ensures that $W_{\omega_1, M} \cap W_{\omega_2, M} = \emptyset$ while (ii) implies that $W_{\omega_i, M} \subseteq W_N$, for $i = 1, 2$. Hence,

$$\hat{D}_{12}/2 = \sum_{\substack{n \in W_N \\ -(W_{\omega_1, M} \cup W_{\omega_2, M})}} (\hat{A}_n^2 + \hat{B}_n^2) \geq 0.$$

□

It is a direct consequence of Prop. 3.4 that if (i) and (ii) are satisfied, then (3.11) returns positive values. In practice, one can use for instance $\omega_1 = 11, \omega_2 = 35$ (Mara, 2009), which are free of interferences up to $2M$ for the usual orders $M = 4, 6$. Then the minimal value of N is $2 \times 6 \times \max\{11, 35\} = 420$.

Bias

The methods based on direct Monte Carlo estimation, *Jansen* and *Liu and Owen TII estimator*, are unbiased since only direct mean estimators are used for the conditional expectations. This is especially remarkable in combination with the positivity of the *Liu and Owen TII estimator*, as it implies that when the true value is zero, the estimator is identical to zero as well.

Proposition 3.5. *If $\mathfrak{D}_{i,j} = 0$, then the Liu and Owen TII estimator is equal to zero:*

$$^{LO}\widehat{\mathfrak{D}}_{i,j} \equiv 0.$$

Proof. Starting with the FANOVA decomposition (2.4), if $\mathfrak{D}_{i,j} = 0$, then all of the terms containing both x_i and x_j vanish. So the decomposition reduces to

$$\begin{aligned} f(\mathbf{x}) &= \sum_{i \notin I \wedge j \notin I} f_I(\mathbf{x}_I) + \sum_{i \in I \wedge j \notin I} f_I(\mathbf{x}_I) + \sum_{i \notin I \wedge j \in I} f_I(\mathbf{x}_I) \\ &= a_0(\mathbf{x}_{-\{i,j\}}) + a_i(x_i, \mathbf{x}_{-\{i,j\}}) + a_j(x_j, \mathbf{x}_{-\{i,j\}}). \end{aligned}$$

Inserting this representation of f into the squared term of the *Liu and Owen TII estimator* (3.10) returns zero,

$$\begin{aligned} &f(x_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) - f(x_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k) - f(z_i^k, x_j^k, \mathbf{x}_{-\{i,j\}}^k) + f(z_i^k, z_j^k, \mathbf{x}_{-\{i,j\}}^k) \\ &= a_i(x_i^k, \mathbf{x}_{-\{i,j\}}^k) - a_i(x_i^k, \mathbf{x}_{-\{i,j\}}^k) - a_i(z_i^k, \mathbf{x}_{-\{i,j\}}^k) + a_i(z_i^k, \mathbf{x}_{-\{i,j\}}^k) \\ &\quad + a_j(x_j^k, \mathbf{x}_{-\{i,j\}}^k) - a_j(z_j^k, \mathbf{x}_{-\{i,j\}}^k) - a_j(x_j^k, \mathbf{x}_{-\{i,j\}}^k) + a_j(z_j^k, \mathbf{x}_{-\{i,j\}}^k) = 0. \end{aligned}$$

□

In the *pick-freeze TII estimator*, the estimation of f_0^2 is necessary in the estimation of the closed sensitivity indices (2.16). If we estimate f_0^2 by taking the square of $\widehat{f}_0$ (2.17) or $^*\widehat{f}_0$ (2.19), we introduce a bias of $-\text{Var}(f_0)$, as was noted by Owen (2013b),

$$[\text{E}(f_0)]^2 = \text{E}[(f_0)^2] - \text{Var}(f_0).$$

However, as long as $E(f(X)^4) < \infty$, the bias is asymptotically negligible. Owen (2013b) remarks that the bias may be important in the case of small closed sensitivity indices, but those are not expected in the estimation of ${}^{\text{pf}}\widehat{\mathfrak{D}}_{i,j}$, as mentioned in Section 3.2. To nonetheless ensure unbiasedness, unbiased estimators for the closed sensitivity index can be used such as ${}^{\text{Cor1}}\widehat{D}_I^C$ or ${}^{\text{Cor2}}\widehat{D}_I^C$ or (2.21), which do not need an estimate of f_0^2 , or the estimator $\widetilde{\tau}^2$ proposed in Section 7 of Owen (2013b).

Frequency-based estimators are generally prone to bias (Tissot and Prieur, 2012). This bias might even be enhanced here through the use of a combination of *FAST* estimators for the *fixing method TII estimator* and *RBD-FAST* estimators for the *RBD-FAST TII estimator*.

Asymptotic properties of the Liu and Owen TII estimator

Corresponding to the asymptotic properties of the pick-freeze estimator of closed sensitivity indices (end of Section 2.3.2), it is possible to derive asymptotic properties of the *Liu and Owen TII estimator* (3.10). To do this, the estimator for a pair of input variables $\{X_i, X_j\}$ is written as

$$T_n = {}^{\text{LO}}\widehat{\mathfrak{D}}_{i,j} = \frac{1}{n} \sum_{k=1}^n \frac{(\Delta_{i,j}^k)^2}{4},$$

with

$$\begin{aligned} \Delta_{i,j}^k = & f(X_i^k, X_j^k, \mathbf{X}_{-\{i,j\}}^k) - f(X_i^k, Z_j^k, \mathbf{X}_{-\{i,j\}}^k) - f(Z_i^k, X_j^k, \mathbf{X}_{-\{i,j\}}^k) \\ & + f(Z_i^k, Z_j^k, \mathbf{X}_{-\{i,j\}}^k). \end{aligned}$$

Proposition 3.6. *Let $\mathcal{P}$ be the set of all cumulative distribution functions of exchangeable random vectors in $L^2(\mathbb{R}^2)$, that is for a $P \in \mathcal{P}$ it holds for any random vectors X and X' that $P(X, X') = P(X', X)$. Then the following propositions hold for the Liu and Owen TII estimator T_n .*

a) T_n is consistent for $\mathfrak{D}_{i,j}$,

$$T_n \xrightarrow[n \rightarrow \infty]{a.s.} \mathfrak{D}_{i,j}.$$

b) T_n is asymptotically normally distributed,

$$\sqrt{n}(T_n - \mathfrak{D}_{i,j}) \xrightarrow[n \rightarrow \infty]{d} \mathcal{N}\left(0, \frac{\text{Var}[(\Delta_{i,j}^1)^2]}{16}\right).$$

c) T_n is asymptotically efficient in the notion of van der Vaart (1998) for estimating $\mathfrak{D}_{i,j}$ for $P \in \mathcal{P}$.

Proof. The results a) and b) are a direct application of the law of large numbers and the central limit theorem, applied to the variables $(\Delta_{i,j}^k)^2$.

Result c) follows from the fact that estimators with symmetrical expressions are asymptotically efficient in the framework of exchangeable variables, as proved by Janon et al. (2013) in their Lemma 2.6 (2):

Let $\Phi_2 : \mathbb{R}^2 \rightarrow \mathbb{R}$ be a symmetric function in $L^2(P)$ and f a deterministic function on $\mathcal{R} \subset \mathbb{R}^{p_1+p_2}$ of independent random input variables $\mathbf{X} \in \mathbb{R}^{p_1}$ and $\mathbf{Z} \in \mathbb{R}^{p_2}$. $\mathbf{Z}'$ denotes an independent copy of $\mathbf{Z}$ and $\mathbf{X}_k, \mathbf{Z}_k, \mathbf{Z}'_k$, $k = 1, \dots, n$ denote independent samples of the corresponding variables. The sequence $\{\Phi_n^2\}_{n \in \mathbb{N}}$ given by

$$\Phi_n^2 = \frac{1}{n} \sum_{k=1}^n \Phi_2\left(f(\mathbf{X}_k, \mathbf{Z}_k), f(\mathbf{X}_k, \mathbf{Z}'_k)\right)$$

is asymptotically efficient for estimating $E(\Phi_2(f(X, Z), f(X, Z')))$ for $P \in \mathcal{P}$.

More precisely, denote $\mathbf{X}_k = (X_j^k, Z_j^k, \mathbf{X}_{-\{i,j\}}^k)$, $\mathbf{Z}_k = X_i^k$, $\mathbf{Z}'_k = Z_i^k$, and let g be the function defined over $\mathbb{R}^d \times \mathbb{R}$ by

$$g(\mathbf{a}, b) = f(b, a_1, a_3, \dots, a_d) - f(b, a_2, a_3, \dots, a_d).$$

Then

$$\Delta_{i,j}^k = g(\mathbf{X}_k, \mathbf{Z}_k) - g(\mathbf{X}_k, \mathbf{Z}'_k).$$

Therefore

$$T_n = \frac{1}{n} \sum_{k=1}^n \Phi_2(g(\mathbf{X}_k, \mathbf{Z}_k), g(\mathbf{X}_k, \mathbf{Z}'_k)),$$

and

$$\mathfrak{D}_{i,j} = \mathbb{E}(\Phi_2(g(\mathbf{X}_1, \mathbf{Z}_1), g(\mathbf{X}_1, \mathbf{Z}'_1))),$$

where Φ_2 is the two-dimensional function defined over $\mathbb{R}^2$

$$\Phi_2(u, v) = \frac{(u - v)^2}{4}.$$

Remark that $\mathbf{Z}_k$ and $\mathbf{Z}'_k$ are independent copies of each other, both independent of $\mathbf{X}_k$, and that Φ_2 is a symmetric function. With the following change in notation

$$i \leftarrow k, \mathbf{X} \leftarrow \mathbf{X}, \mathbf{Z} \leftarrow \mathbf{Z}, \mathbf{Z}' \leftarrow \mathbf{Z}', f \leftarrow g$$

the result follows from the Lemma of Janon et al. (2013) named above.

□

The two last propositions can be extended to the general superset importance (3.2), including the case of the total sensitivity index of one input variable.

Proposition 3.7. *Let $\Upsilon_I = \sum_{J \supseteq I} D_J$ be the superset importance for a set I . Define*

$$T_{I,n} = \frac{1}{n} \sum_{k=1}^n \frac{(\Delta_I^k)^2}{2^{|I|}},$$

with $\Delta_I^k = \sum_{J \subseteq I} (-1)^{|I-J|} f(\mathbf{Z}_J^k, \mathbf{X}_{-J}^k)$, where $|\cdot|$ stands for the number of elements in a set and $A - B = \{x | x \in A \text{ and } x \notin B\}$ denotes the difference of two sets A and B .

Then $T_{I,n}$ is asymptotically normal and asymptotically efficient for Υ_I .

Proof. Note that $T_{I,n}$ is the sample version of the formula (10) given by Liu and Owen (2006) for Υ_I (with suitable change of notations). The proof of asymptotic normality is thus a direct consequence of the central limit theorem. For asymptotic efficiency, the proof relies on arguments similar to Prop. 3.6.

- When $I = \{i\}$ is a single input variable, we have

$$\Delta_I^k = f(Z_i^k, \mathbf{X}_{-i}^k) - f(X_i^k, \mathbf{X}_{-i}^k),$$

which is of the form $g(\mathbf{x}_k, \mathcal{Z}_k) - g(\mathbf{x}_k, \mathcal{Z}'_k)$ with $\mathcal{Z}_k = Z_i^k$, $\mathcal{Z}'_k = X_i^k$, $\mathbf{x}_k = \mathbf{X}_{-i}^k$, and $g(\mathbf{a}, b) = f(b, \mathbf{a})$.

- When $|I| \geq 2$, let us choose $i \in I$. Then, by splitting the subsets of I into two parts, depending whether they contain $\{i\}$, we have

$$\begin{aligned} \Delta_I^k &= \sum_{J \subseteq I - \{i\}} (-1)^{|I-J|} f(\mathbf{Z}_J^k, \mathbf{X}_{-J}^k) \\ &\quad + \sum_{J \subseteq I - \{i\}} (-1)^{|I-(J \cup \{i\})|} f(\mathbf{Z}_{J \cup \{i\}}^k, \mathbf{X}_{-(J \cup \{i\})}^k) \\ &= \sum_{J \subseteq I - \{i\}} (-1)^{|I-J|} f(\mathbf{Z}_J^k, X_i^k, \mathbf{X}_{-(J \cup \{i\})}^k) \\ &\quad - \sum_{J \subseteq I - \{i\}} (-1)^{|I-J|} f(\mathbf{Z}_J^k, Z_i^k, \mathbf{X}_{-(J \cup \{i\})}^k), \end{aligned}$$

which is also of the form $g(\mathbf{x}_k, \mathcal{Z}_k) - g(\mathbf{x}_k, \mathcal{Z}'_k)$ with $\mathcal{Z}_k = X_i^k$, $\mathcal{Z}'_k = Z_i^k$, $\mathbf{x}_k = (\mathbf{X}_{I-\{i\}}^k, \mathbf{Z}_{I-\{i\}}^k, \mathbf{X}_{-I}^k)$, and a suitable g since the second term in the difference is obtained from the first one by exchanging X_i^k and Z_i^k .

The result is then obtained by applying Lemma 2.6 in Janon et al. (2013) to the symmetric function $\Phi_2(u, v) = \frac{(u-v)^2}{2|I|}$, remarking that $\mathcal{Z}_k$ and $\mathcal{Z}'_k$ are independent copies of each other, both independent of $\mathbf{x}_k$. $\square$

Corollary 3.1. *It follows from Prop. 3.7 that the Jansen estimator of the total sensitivity index of a single variable D_i^T is asymptotically efficient.*

3.4 Estimating the full set of TIIs

In most of the practical applications, the full set of $\frac{d(d-1)}{2}$ TIIs, that is the TII of each possible pair of the d input variables, is required simultaneously. The function evaluations, especially in the case without a metamodel, where f represents the actual computer experiment, are usually the most time-consuming part of the estimation. Table 3.2 therefore gives an overview of the number of evaluations N that are required by each method for the estimation of the full set of TIIs, depending on the different parameter settings for each method.

In each of the n Monte Carlo runs, the *Liu and Owen TII estimator* requires one estimate corresponding to $f(\mathbf{x})$, d estimates corresponding to $f(z_i, \mathbf{x}_{-i})$ — and at the same time to $f(z_j, \mathbf{x}_{-j})$ — and $\binom{d}{2}$ estimates corresponding to $f(z_i, z_j, \mathbf{x}_{-\{i,j\}})$. Similarly, the *RBD-FAST TII estimator* requires d estimates of $^{\text{RBD}}\widehat{D}_i^T$ and $\binom{d}{2}$ estimates of $^{\text{RBD}}\widehat{D}_{\{i,j\}}^T$, and each estimate requires a number of $2(Md + L)$ runs. The *fixing method TII estimator* requires for each of the $\binom{d}{2}$ indices a number of n_{MC} FAST estimates.

For the *Jansen* and the *pick-freeze TII estimator*, strategies that reuse computations corresponding to the strategy by Saltelli (2002) (Section 2.3.2) can be applied, as addressed in their definitions (3.7) and (3.9). In the estimation of the first- and second-order total indices for the *Jansen TII estimator* as well as in the estimation of the $(d - 1)$ - and $(d - 2)$ -order closed indices for the *pick-freeze TII estimator*, only the $(d + 1)n$ samples $f(\mathbf{x}^k)$ and the $f(z_i^k, \mathbf{x}_{-i}^k)$, $1 \leq k \leq n$, $1 \leq i \leq d$ are required for the estimation of all TII indices. In addition, those function evaluations can be used to estimate the full set of first-order total sensitivity indices using either (2.22) or (2.23). By adding only the n samples $f(\mathbf{z}^k)$, $1 \leq k \leq n$, all first-order Sobol indices (2.16) can be estimated as well.

Table 3.2 clearly shows that in contrast to the other estimators, the *Jansen* as well as the *pick-freeze TII estimator* are scalable in the sense that the number of required runs increases linearly with the number of input variables d .

TII estimator	number of function evaluations
<i>Jansen</i>	$N = (d + 1) \times n$
<i>RBD-FAST</i>	$N = \left(\binom{d}{2} + d\right) \times 2(Md + L)$
<i>pick-freeze</i>	$N = (d + 1) \times n$
<i>Liu and Owen</i>	$N = \left(\binom{d}{2} + d + 1\right) \times n$
<i>fixing method</i>	$N = \binom{d}{2} \times n_{\text{MC}} \times n_{\text{FAST}}$

Table 3.2: Number of function evaluations for the considered TII estimators.

3.5 Simulation study

In the following, the behavior of the different estimators is investigated in simulations on known analytical real-valued functions. In order to cover a wide range of functions, functions are chosen that fall in different categories according to different properties based on Kucherenko et al. (2009). The complexity of the function is considered in two different categories, the interaction order and the functional form. For the latter, simple forms like low-order polynomial functions are considered as of low and complex nonlinear forms as of high complexity. In detail, the categories are

- the number of input variables: low ($d \leq 5$) or high ($d \geq 10$),
- interaction behavior: low (only low-order terms are dominant) or high (important high-order interactions),
- functional form: simple (simple functional form like low-order polynomial) or complex (complex nonlinear behavior).

See Tab. 3.3 for an overview of the functions used in the simulation. The functions are listed along with an identification and the distribution of their inputs. Function *function 1* is taken from Fruth et al. (2014b), *Ishigami* from Ishigami and Homma (1990), and *Branin4* from Lehman et al. (2004). Table 3.4 shows how the functions are classified into the defined categories.

name	function	$\mu(\mathbf{X})$
<i>function 1</i>	$\sin(X_1 + X_2) + 0.4 \cos(X_3 + X_4)$	$U[-1, 1]^4$
<i>Ishigami</i>	$\sin(X_1) + 7 \sin^2(X_2) + 0.1 X_3^4 \sin(X_1)$	$U[-\pi, \pi]^3$
<i>pure 3rd order</i>	$\sqrt{3^3} \prod_{i=1}^3 X_i$	$U[-1, 1]^3$
<i>Branin4</i>	$\frac{1}{30} f^{\text{Branin}}(X_1, X_2) f^{\text{Branin}}(X_3, X_4) + (X_1 - \pi)^2$ with $f^{\text{Branin}}(X_1, X_2) = \left(X_2 - \frac{5.1}{4\pi^2} X_1^2 + \frac{5}{\pi} X_1 - 6\right)^2$ $+10 \left(1 - \frac{1}{8\pi}\right) \cos(X_1) + 10$	$X_1, X_3 \sim U[-5, 10],$ $X_2, X_4 \sim U[0, 15]$
<i>high 2nd order</i>	$3 \left(\sum_{i=1}^7 X_i\right) \left(\sum_{j=8}^{14} X_j\right)$	$U[-1, 1]^{14}$
<i>high 14th order</i>	$\sqrt{3^{14}} \prod_{i=1}^{14} X_i$	$U[-1, 1]^{14}$

Table 3.3: Overview of test functions.

$d \leq 5$		$d \geq 10$	
form = simple		form = complex	
interaction = low	<i>function 1</i>	<i>Ishigami</i>	<i>high 2nd order</i>
interaction = high	<i>pure 3rd order</i>	<i>Branin4</i>	<i>high 14th order</i>

Table 3.4: Classification of test functions.

For each of the functions, the TIIs are estimated a hundred times by all 5 methods. For each estimation, 2 000 Monte Carlo draws per index are used, $2\,000 \times \binom{d}{2}$ draws in total. The parameters of each method are determined using the connections in Tab. 3.2.

Boxplots of the respective 100 estimates give an impression of the performance in terms of bias and variance for the different classes of functions and are presented in the figures 3.2 - 3.4. The detailed numerical results can be found in the Appendix in Tab. B.1 - B.6.

Figure 3.2 shows the boxplots of the simulation results of the two test functions with low dimension and only second-order interactions. As expected in Section 3.3, negative results are observed for the *RBD-FAST*, *Jansen*, and *pick-freeze* estimates, but not for the *Liu and Owen* and the *fixing method*. Overall, the *Liu and Owen* and the *fixing method TII estimator* outperform the other estimators clearly, especially for influences close to zero. That means that the *Liu and Owen TII estimator* enables a precise detection of inactive interactions, an important task for interaction screening. The *RBD-FAST TII estimator* shows a rather large variance for the case with low complex functional form *function 1* and a severe bias for the more complex *Ishigami* function. One reason for this might be the bias for RBD-FAST methods, mentioned in Section 3.3. This problematic behavior of the *RBD-FAST TII estimator* continues in the other four simulations. The other four estimators perform well in terms of bias, confirming the results of Section 3.3. The described bias of the *pick-freeze TII estimator* is not noticeable, the deviations from the true mean are comparable to the *Jansen* and the *Liu-Owen TII estimator*.

Simulation boxplots for functions with high order interactions but a rather small number of input variables are shown in Fig. 3.3. Here, the *Liu and Owen* and the *fixing method TII estimator* show a stronger variance than before. The reason for this lies in the variance of the estimates $\hat{D}_{i,j|\mathbf{X}_{-\{i,j\}}}^k$ in the basic fixing method formula (3.6), which applies for both methods. For second-order interactions, those estimates do not depend on the fixed values $\mathbf{X}_{-\{i,j\}}$ and thus vary only slightly while for higher interactions the estimates should differ with the fixed variables included in the interaction. The much larger variance of the *fixing method TII estimator* can be additionally explained by the fact that the two steps, second-order interaction estimation and overall expectation, are performed separately, and the values for $\mathbf{X}_{-\{i,j\}}$ are kept constant in each second-order index estimation, whereas in the *Liu and Owen method* both steps are performed simultaneously and thus are more efficient. Generally, for complex

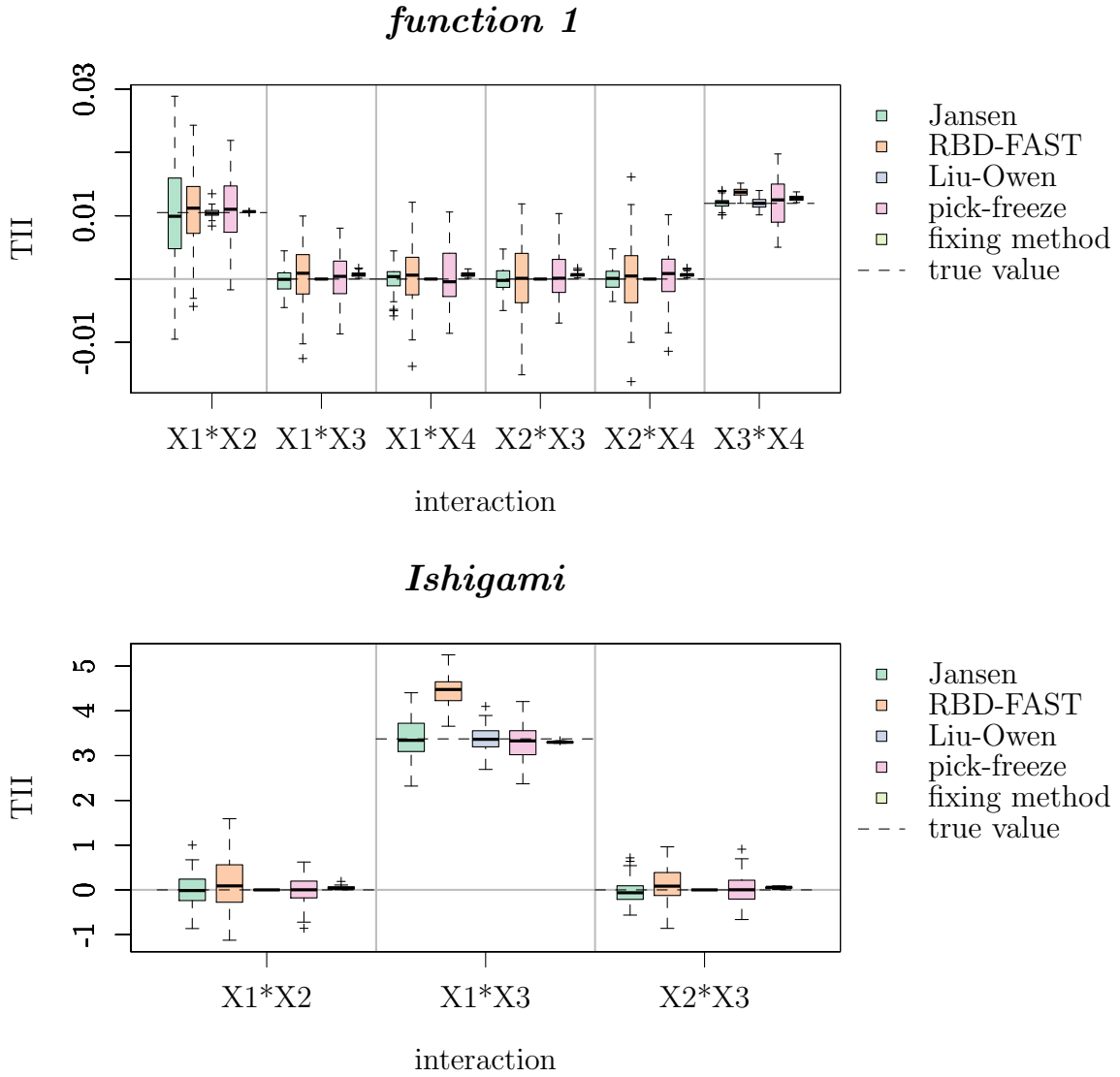


Figure 3.2: Simulation study on TII estimators, boxplots of 100 estimations for functions with low interactions, low number of input variables.

functions, the methods *Jansen*, *Liu* and *Owen* and *pick-freeze* TII estimator seem to perform equally well.

In the simulation study of the two high-dimensional test functions, a 10 times higher number, 20 000, of Monte Carlo draws is necessary due to numerical problems for the extremely complex case of a 14th order interaction. As the *RBD-FAST* TII estimator

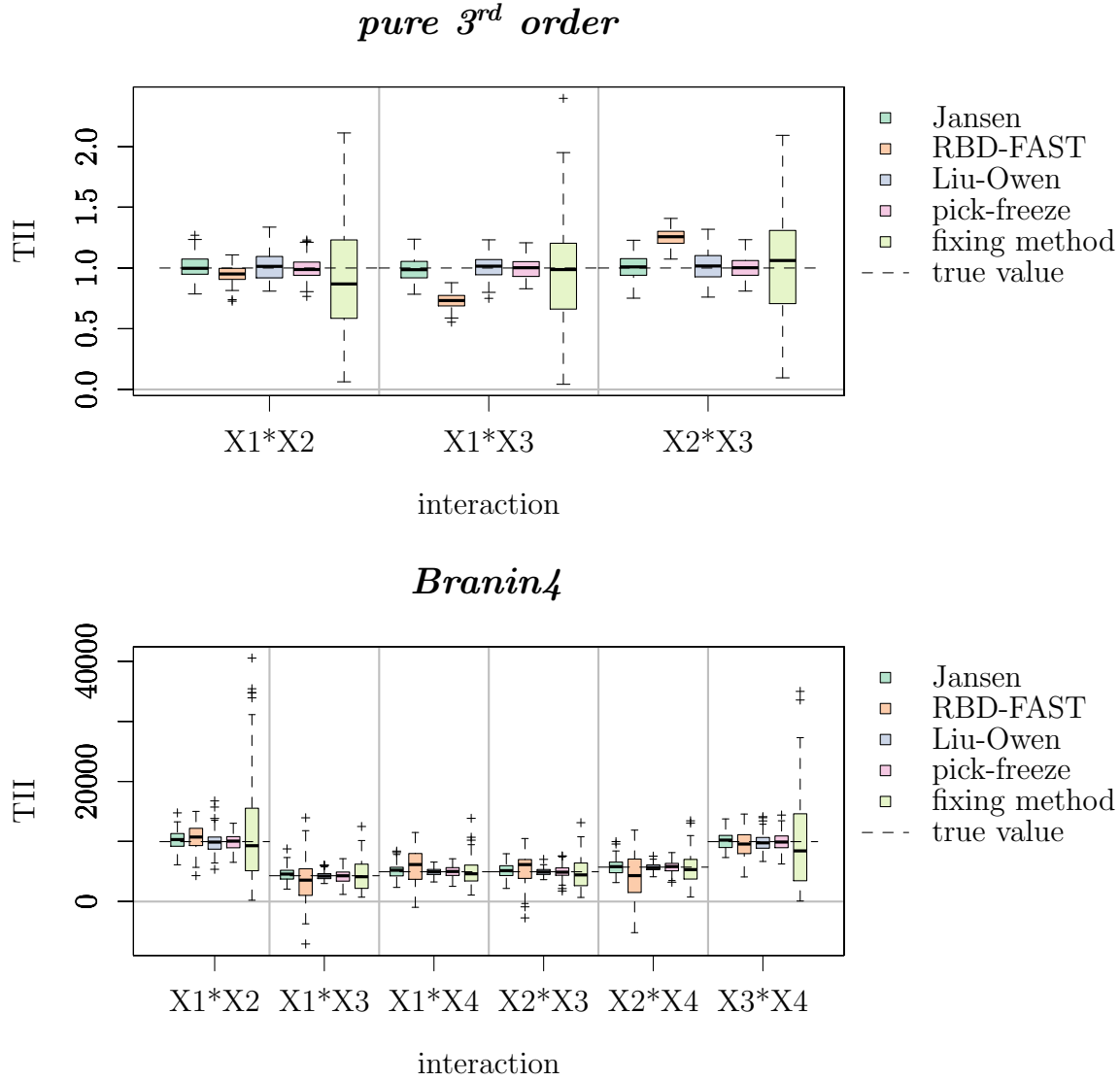


Figure 3.3: Simulation study on TII estimators, boxplots of 100 estimations for functions with high interactions, low number of input variables.

showed a severe bias as well as a large variance, it is not studied any further. The same observations as from the low-interaction cases can be made for the *high 2nd order* function, the *Liu and Owen* and the *fixing method* TII estimator outperform the other estimators. Thus, for low interactions, those estimators still work well in high dimensions. Only for the complex *high 14th order* function, the *Jansen* and the *pick-*

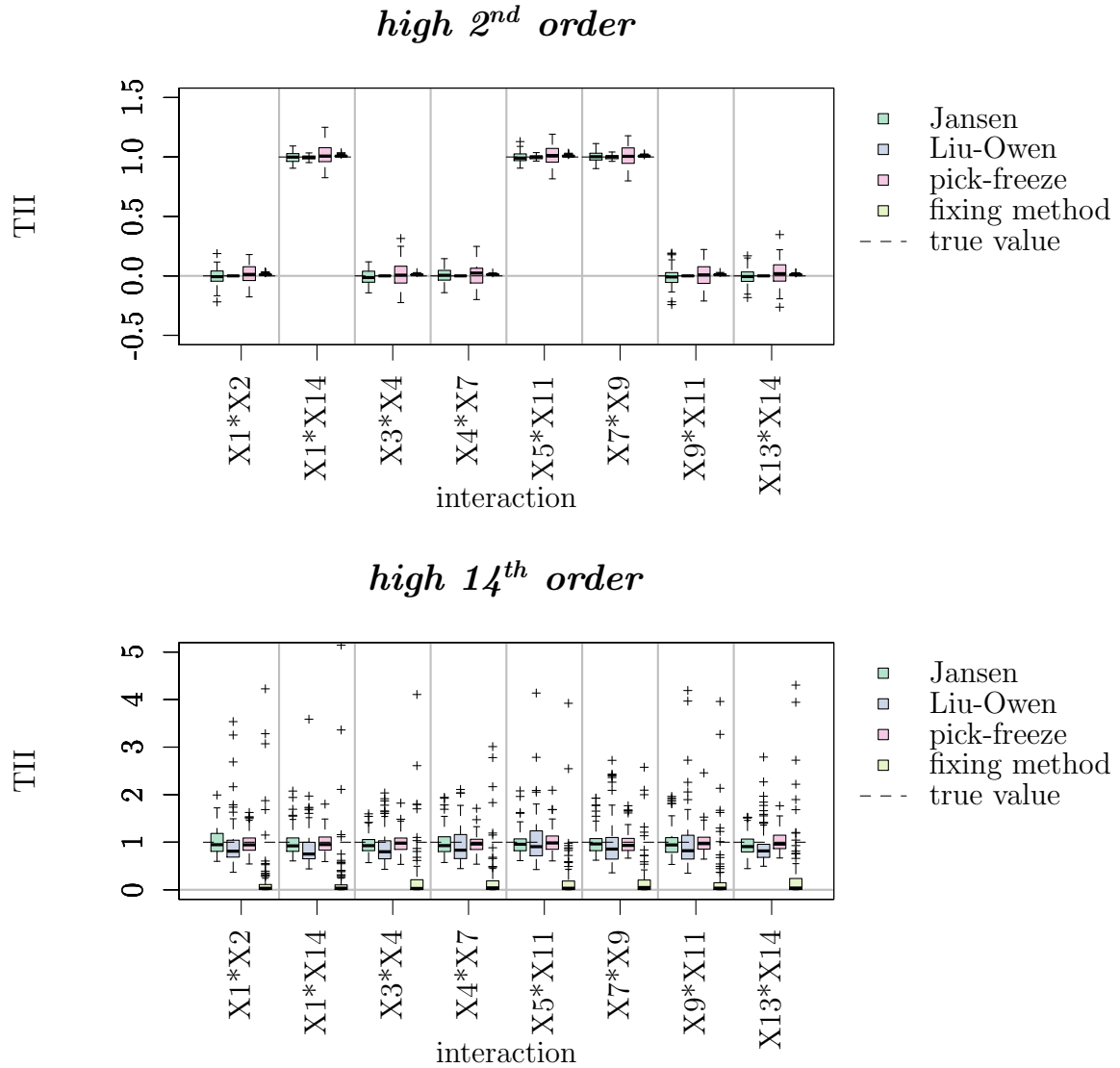


Figure 3.4: Simulation study on TII estimators, boxplots of 100 estimations of the TIIs of selected interactions for functions with high number of input variables.

freeze TII estimator outperform the *Liu and Owen TII estimator* in terms of variance. This is due to the strategy to reuse computations, which makes those two methods scalable (see Section 3.4). The *fixing method TII estimator* shows a bias, perhaps caused by the bias in FAST described e.g. in Tissot and Prieur (2012).

To get an impression of the asymptotic behavior, Fig. 3.5 shows exemplarily the estimates of the four reliable estimators *Jansen*, *Liu and Owen*, *pick-freeze* and *fixing method TII estimator* for index $\mathfrak{D}_{13}$ of the *pure 3rd order* function. For each method, the index is estimated starting with the lowest number of evaluations possible and then evaluations are added step by step. The estimate of the asymptotically efficient estimator by *Liu and Owen* is the fastest to come close the true value followed by the *pick-freeze* and the *Jansen TII estimator* while the estimate of the *fixing method TII estimator* does not reach it within the 50 000 evaluations.

The simulation study shows that the *Liu and Owen TII estimator* indeed provides good results for almost all situations and can thus be recommended as method of choice. The only exception are situations with a high number of variables and strong interactions. There, the *pick-freeze TII estimator* is rather preferable.

3.6 Threshold decision

Errors in the initial metamodel and in the index estimation usually perturb the TII estimation so that, when applying them to interaction screening, inactive indices are not necessarily zero. Therefore, a threshold needs to be introduced to separate active from inactive interactions. The decision rule is then to include the interaction if

$$\frac{\hat{\mathfrak{D}}_{i,j}}{\hat{D}} > \delta. \quad (3.12)$$

If the TII estimation is further used for adapted modeling (Mühlenstädt et al., 2012) or parallelized optimization (Ivanov and Kuhnt, 2014), it is more problematic to cut off an active interaction than to falsely accept an inactive interaction as active. For modeling, Mühlenstädt et al. (2012, p. 735) mention that the first case can lead to substantial modeling error whereas the second leads to a higher number of parameters, but still the true structure is contained in the model. In the optimization application, Ivanov and Kuhnt (2014, pp. 8-9) report that the first case of cutting an active interaction can

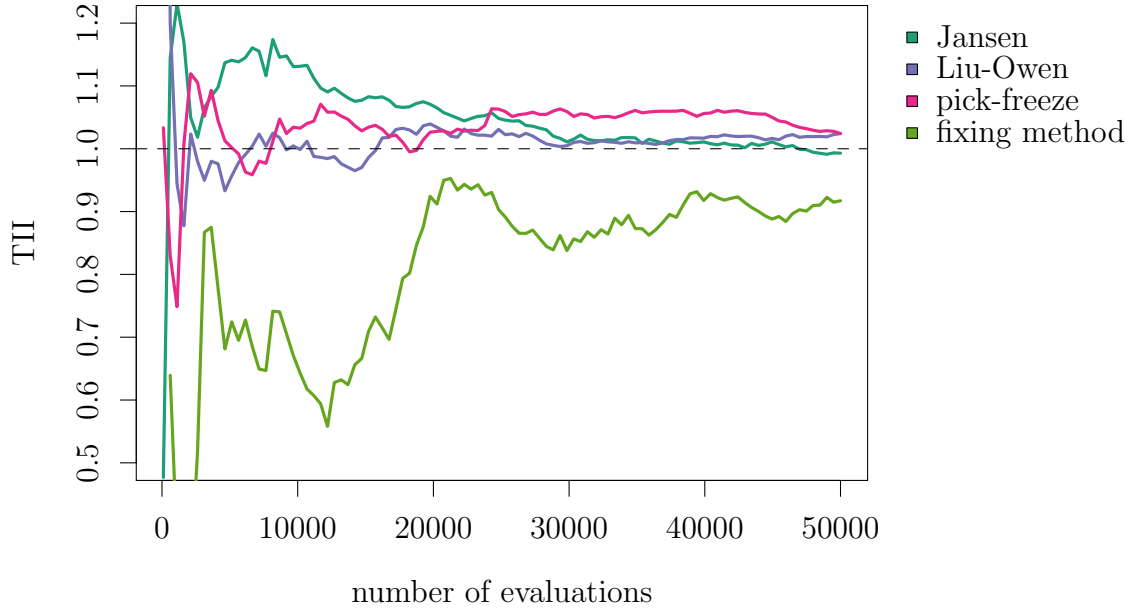


Figure 3.5: Simulation study on TII estimators, asymptotical exploration on the *pure 3rd order* function.

lead to a wrong optimum as the assumption of additivity is crucial in the parallelization whereas the keeping of an inactive interaction makes the procedure less efficient but does not effect the optimum. Finding at least all active interactions therefore is an important property of a good threshold.

In the following, a simple graphical method to support the threshold decision is introduced as well as a data-driven approach, which uses block-additive Kriging kernels.

Delta jump plot

As decision plots are a simple and clear way to help the user to find data dependent values, we suggest a plot, the *delta jump plot*, which helps to find suitable candidates

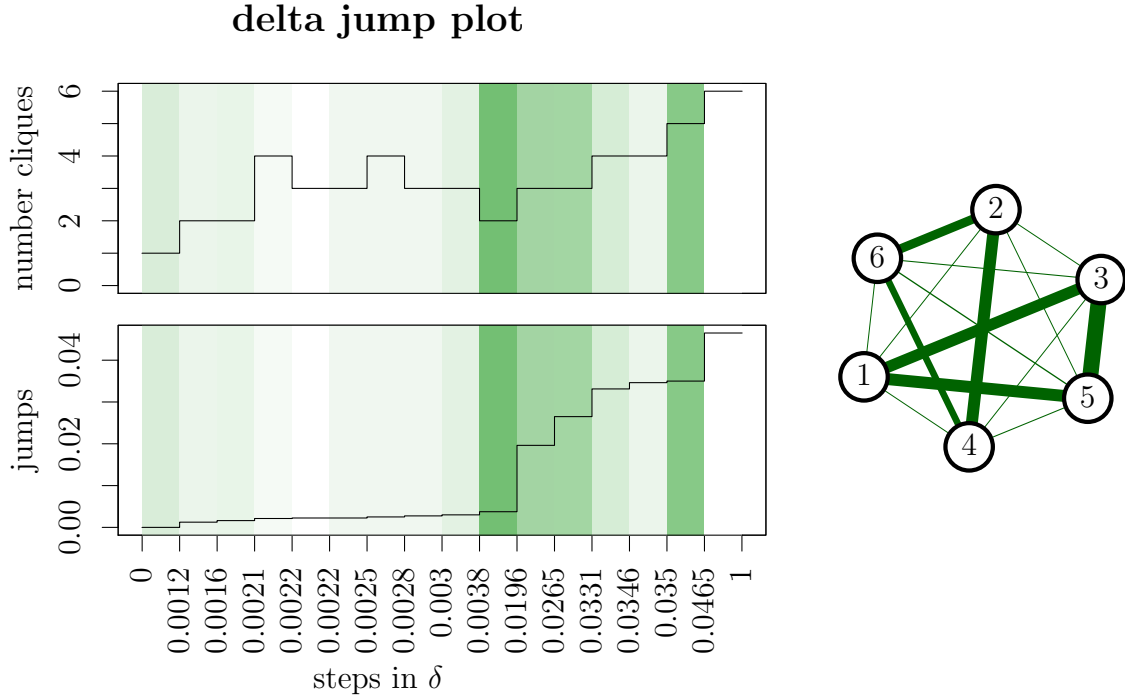


Figure 3.6: Delta jump plot of an artificial example (left) and corresponding FANOVA graph (right).

for the threshold value δ . An artificial example of such a plot can be seen in Fig. 3.6 on the left. The corresponding full FANOVA graph, where the TIIs are represented by graph edges, is shown on the right.

The delta jump plot consists of two parts. In the bottom part, it shows the ordered values of all $\binom{d}{2}$ estimated TIIs. High jumps point to big differences between successive indices. In a situation with moderate perturbation, the difference between inactive and active indices should be high and thus be revealed by a jump. High jumps are additionally highlighted by shading of the relevant interval, the darker the shading the higher the jump, so that values of delta that lie inside dark intervals indicate good choices for the threshold. Simultaneously, the upper part of the plot shows the number of cliques of the corresponding FANOVA graph a threshold would lead to. A small number of cliques might be preferable in order to obtain a clearer structure. For the

artificial case in Fig. 3.6, we see that a suitable δ lies between 0.0038 and 0.0196, for example 0.01, because the jump to the right of this interval is the highest and the number of cliques is very low. Another possible candidate would be 0.04.

Kriging kernel comparison

One method to compare different candidates for the threshold, possibly found by looking at the *delta jump plot*, is to compare the performance of Kriging models with corresponding block-additive structures.

As described in Mühlenstädt et al. (2012) and already mentioned in Section 3, Kriging kernels can be adapted to mimic the interaction structure given from the TII interaction screening. Instead of a product kernel (compare (2.3)), a *block-additive kernel* is applied. Basing on the cliques $C_1, \dots, C_L$ of the FANOVA graph, it is constructed as the sum of kernels $k_{C_1}, \dots, k_{C_L}$, where each k_{C_l} is chosen to be a tensor-product kernel in dimension $|C_l|$,

$$k(\mathbf{h}) = \sum_{l=1}^L k_{C_l}(\mathbf{h}_{C_l}).$$

Kernels adapted to the true interaction structure of the underlying model are supposed to be more suited to represent the underlying model and thus lead to better predictions. The idea for the threshold decision is to set up Kriging models according to some candidate thresholds and compare their leave-one-out prediction performance. See below for the precise approach.

- Choose a (low) number c of possible candidate thresholds and add 0 and 1 for the extreme cases of all input variables interacting and no interaction: $\boldsymbol{\delta} = (0, \delta_1, \dots, \delta_c, 1)$.
- For each threshold $\delta_i \in \boldsymbol{\delta}$:
 - Determine the according active interactions following rule (3.12).
 - Set up the FANOVA graph and determine the cliques $C_1, \dots, C_l$.

- Construct the adapted Kriging model with block-additive kernel,

$$k(\mathbf{h}) = \sum_{l=1}^L k_{C_l}(\mathbf{h}_{C_l}).$$

- Predict available original evaluations of the underlying model via leave-one-out cross-validations and compute the resulting RMSE value.
- Compare the leave-one-out RMSE values for each threshold candidate and choose the threshold with the smallest RMSE. It is also possible to plot scatterplots of predicted versus true values to get an impression of the precision of each prediction.

Simulation study

The performance of the threshold decision methods shall be presented in a small simulation study with four input variables ($d = 4$). Five different interaction structures (*str1* to *str5*) are studied, here represented as sets of cliques of the corresponding graph structure. The goal is to find the threshold that rediscovers the corresponding graph.

- *str1*: $\{1\}, \{2\}, \{3\}, \{4\}$ (pure additivity),
- *str2*: $\{1, 2\}, \{3\}, \{4\}$ (two variables interacting),
- *str3*: $\{1, 2\}, \{3, 4\}$ (two pairs of variables interacting),
- *str4*: $\{1, 2, 3\}, \{4\}$ (three variables interacting, one additive),
- *str5*: $\{1, 2, 3, 4\}$ (full graph).

For each structure, a number of 100 simulations are performed via the following steps, using mainly the R package **fanovaGraph**, described in the next chapter. A maximin distance design within the class of Latin hypercubes (Morris and Mitchell, 1995) is constructed for a number of 50 runs using the R package **DiceDesign** (Franco et al., 2014). As underlying model f , a Gaussian process with block-additive kernel corresponding to one of the interaction structures *str1*, ..., *str5* is simulated at the design points. See Fruth et al. (2013) for the simulation of block-additive processes. On the

interaction structure	<i>str1</i>	<i>str2</i>	<i>str3</i>	<i>str4</i>	<i>str5</i>
complete graph identified	0.99	0.68	0.50	0.85	0.60
all active interactions identified	1.00	0.98	0.81	0.95	0.60

Table 3.5: Proportions of correctly identified graphs in threshold simulation.

simulated data, an initial metamodel is fit. We use a Gaussian process model with a special ANOVA kernel (Durrande et al., 2013) since it is assumed to be more adapted to block-additive situations than the standard product kernel (2.3), which a priori assumes a full graph. For each interaction, the TII is estimated via the *Liu and Owen TII estimator* on a number of 20 000 evaluations for the whole set. The Kriging kernel comparison procedure described above is then executed. The candidate thresholds for the procedure are obtained from the three biggest jumps in the delta jump plot.

The results can be seen in Tab. 3.5. In the first row, the proportion of all correctly identified graphs is shown. The second row shows the proportion of those situations where all active interactions are identified correctly but not necessary all inactive ones. As described before, this property is crucial for subsequent methods. A threshold that identifies all active interactions does at least not disturb subsequent analysis. For *str1* and *str5*, the results of the second row are obvious due to their interaction structure. The simulation results show that the performance depends on the interaction structure. Situations with few interactions are better discovered than complex ones. The interaction structure *str3* with two similar interactions seems to be more difficult to predict than the unbalanced one *str4*. The results show that the Kriging kernel comparison leads to the right decision in the majority of the simulations and, except for the most complex case of a full graph, the active interactions are identified in more than 4 out of 5 simulations.

The simulations are repeated using the standard product kernel instead of the ANOVA kernel in the initial model in order to study the impact of the initial kernel on the

interaction structure	<i>str1</i>	<i>str2</i>	<i>str3</i>	<i>str4</i>	<i>str5</i>
complete graph identified	0.99	0.40	0.03	0.18	0.87
all active interactions identified	1.00	0.97	0.79	0.87	0.87

Table 3.6: Proportions of correctly identified graphs in threshold simulation using the standard product kernel as initial kernel.

performance. The results can be seen in Tab. 3.6. As assumed, the performance stays behind the performance of the ANOVA kernel, especially for the cases with two cliques, *str3* and *str4*, for which only a very small proportion of the simulation is identified correctly. An exception is the full graph *str5*, which was identified well. Despite the worse performance in the graph identification, the kernel lead to satisfactory results in the identification of all active interactions. This shows that with the standard kernel the procedure tends to identify too much interactions as active, explainable by the product structure of the kernel, which assumes a full graph. Thus, in this simulation study on simulated block-additive kernels, the ANOVA kernel performed generally better than the standard product kernel.

3.7 Implementation

As part of this work, the methods presented in this chapter are implemented in the R package `fanovaGraph` (Fruth et al., 2013), which is online available on the repository CRAN (Comprehensive Archive Network). It includes the TII estimators of Section 3.2, the threshold decision methods of Section 3.6, and the FANOVA graph visualization mentioned in the motivational example together with Kriging model adaption. The package is explained in detail in the manual by Fruth et al. (2013).

Figure 3.7 shows a diagram of the structural setup of the package. The TII estimation is done via the function `estimateGraph`. If the computer experiment can be incorporated into the R environment, it can be used directly. Otherwise a metamodel can be inserted as described in Section 2.2. The function allows for the four TII estimation methods *RBD-FAST*, *pick-freeze*, *Liu and Owen*, and *fixing method TII estimator* (compare Tab. 3.1). Its input and output behavior matches with the corresponding functions in the comprehensive package for sensitivity analysis, `sensitivity` (Pujol et al., 2014). The package can thus easily be included into a broader study. The function `estimateGraph` returns an S3 object named `graphlist`, which contains the TII values and further information necessary to set up the FANOVA graph. Applying the generic functions `print` or `plot` to the object results in a printed output of the estimated TII values or a plot of the FANOVA graph via package `igraph` (Csárdi and Nepusz, 2006), respectively. Corresponding to Section 3.6, a threshold cut can be performed by function `threshold` that takes the `graphlist` object as input and returns a new `graphlist` object, whose TII values below a given threshold are set to zero and thus not printed in the FANOVA graph plot. The two methods that support the threshold decision are implemented via the functions `plotDeltaJumps` for the *delta jump plot* and `thresholdIdentification` for the Kriging kernel comparison. Finally, the `graphlist` object can be used to adapt Kriging models as described in Mühlenstädt et al. (2012) and mentioned in the motivational example. The function `kmAdditive` takes a data set coming from the computer experiment under study as well as the given block-additive structure in the `graphlist` object and returns a block-additive Kriging model by use of the popular Kriging package `DiceKriging` (Roustant et al., 2012). This block-additive Kriging model can then again be used to predict new data via the function `predictAdditive`.

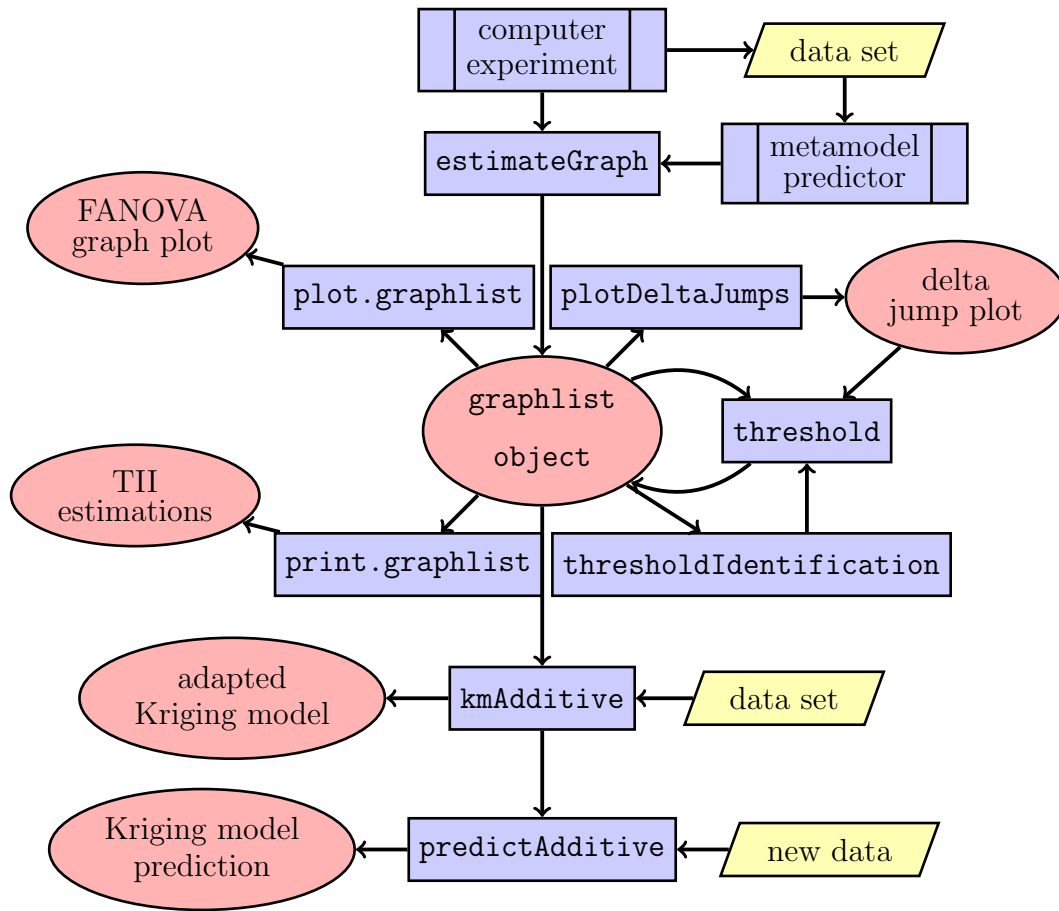


Figure 3.7: Structural setup of R package `fanovaGraph`. Functions are represented by rectangles (with double vertical lines if defined externally), data by parallelograms, and output by circles. Arrows represent information transfer.

3.8 Crossed DGSM

Similar to the total sensitivity index and the DGSM, the notion of the TII can as well be formulated as a derivative-based index, which provides a faster-to-compute index, which serves as an upper bound for the TII. In the context of statistical learning, Friedman and Popescu (2008) introduced a quantity that generalizes the DGSM to

interactions by using cross-partial derivatives and that we propose to denote by *crossed DGSM*,

$$\nu_{i,j} = \int \left(\frac{\partial^2 f(\mathbf{x})}{\partial x_i \partial x_j} \right)^2 d\mu(\mathbf{x}). \quad (3.13)$$

It is assumed that f fulfills $\frac{\partial f(\mathbf{x})}{\partial x_i \partial x_j} \in L^2(\mu)$. From the definition, it is apparent that if $\nu_{i,j} = 0$, no term of the FANOVA decomposition (2.4) that contains x_i and x_j can be active at the same time, i.e. $\nu_{i,j} = 0 \Rightarrow \mathfrak{D}_{i,j} = 0$. The crossed DGSM can thus be used for interaction screening.

Link to the TII

In Roustant, Fruth, Iooss, and Kuhnt (2014) it is shown that the crossed DGSM provide upper bounds for the TII like the DGSM for the total sensitivity index.

Proposition 3.8. *Let us consider n distributions $\mu_1, \dots, \mu_n$ on the real line $\mathbb{R}$, and $\mu = \mu_1 \otimes \dots \otimes \mu_n$. Assume that all μ_i ($i = 1, \dots, n$) satisfy the Poincaré inequality (2.32). Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ be a function in $L^2(\mu)$, such that all first-order and crossed second-order partial derivatives are in $L^2(\mu)$. Then for all pairs $\{i, j\}$ ($1 \leq i, j \leq n$),*

$$\mathfrak{D}_{i,j} \leq C(\mu_i)C(\mu_j)\nu_{i,j}. \quad (3.14)$$

Furthermore, $C_{opt}(\mu_i)C_{opt}(\mu_j)$ is the best constant: If equalities can be achieved in the Poincaré inequalities satisfied by each distribution, then the inequality (3.14) reaches equality as well.

Two different proofs are given in Roustant et al. (2014), one via sequential application of Poincaré inequalities and one using the crossed finite difference in the Liu and Owen formula (3.5).

Estimation

As the crossed DGSM serve as upper boundaries for the TII, they can be used as a cheaper method for the screening of interactions, especially if the Hessian of the underlying model is available. As this is rarely the case, the necessary second-order derivatives usually have to be approximated, e.g. by finite differences. With $\mathbf{x}$ a $n \times d$ -data set and δ^* a small real number, an estimator is given by

$$\hat{\nu}_{i,j} = \frac{1}{n} \sum_{k=1}^n \left(\frac{\Delta_\nu}{(\delta^*)^2} \right)^2,$$

with $\Delta_\nu =$

$$f(x_i^{(k)} + \delta^*, x_j^{(k)} + \delta^*, \mathbf{x}_{-\{i,j\}}^{(k)}) - f(x_i^{(k)} + \delta^*, x_j^{(k)}, \mathbf{x}_{-\{i,j\}}^{(k)}) - f(x_i^{(k)}, x_j^{(k)} + \delta^*, \mathbf{x}_{-\{i,j\}}^{(k)}) + f(\mathbf{x}^{(k)}).$$

Analytical examples as well as practical applications can again be found in Roustant et al. (2014).

Conclusion

This chapter introduced a complete methodology for sensitivity analysis of the interaction structure of a black-box function as extension to a classical sensitivity analysis of first-order and total effects. The total interaction index (TII) was introduced, which, in the same way as the total sensitivity index screens out the most influential input variables, provides a screening of interactions. It reveals the way the input variables interact, which can be visualized in the FANOVA graph. The resulting information on the block-additive interaction structure can further be employed in adapted modeling or parallelized optimization.

Five possible estimators were presented together with their properties. The *Liu and Owen TII estimator* was proven to have good properties to estimate a single TII, as it is unbiased, nonnegative, and asymptotically efficient. Its superiority was confirmed

in simulations on analytical functions. Only for complex functions with a high number of inputs, the *pick-freeze TII estimator* performed better. There, the increase of costs when the dimension increases is linear instead of quadratic as for the *Liu and Owen TII estimator*. Furthermore, an approach to find a threshold for interaction screening was presented, a combination of a graphical decision and Kriging model comparison. The corresponding R package **fanovaGraph** was presented. Finally, crossed DGSM were introduced as extensions of the DGSM to interaction analysis, and the inequality between total effects and DGSM was generalized in the context of interaction screening. So far, all input variables were assumed to be continuous scalar variables. The next chapter presents a method to explore the sensitivity of additional functional input variables.

4. SENSITIVITY ANALYSIS FOR FUNCTIONAL INPUT

The complexity of the input settings of computer experiments is usually not limited to scalar numbers. Examples are industrial processes, where parameters can be varied over the process time or parameters in geological simulations, which have a spatial distribution. The inputs in these examples can be considered as functional input variables, dependend on time or space, respectively. They introduce an additional dimension to computer experiments and make new methods for the analysis necessary. For sensitivity analysis, it is especially interesting to explore the influence over the functional domain in order to find out about the inputs influence at different regions of the domain.

Usually, data with functional input is studied by functional linear regression (see Ramsay and Silverman (1997)), a framework that allows for approximation, modeling and prediction of data with functional input. Applied as a technique for sensitivity analysis, it has some drawbacks. In the functional linear regression framework, data is usually assumed as already provided so that no statistical design techniques are required. In addition, the interpretation of the influence of different regions of the functional domain is not easy (see e.g. James et al. (2009)) and it is restricted to linear behavior. Non-linear modeling, however, is possible by functional Gaussian process frameworks as for instance presented by Morris (2012) and, specifically for scalar output, by Mühlenstädt et al. (2014). These methods also allow for functional space-filling design, but are not constructed for sensitivity analysis. Iooss and Ribatet (2009), Lilburne and Taran-

tola (2009), and Fort et al. (2013) eventually present several ideas for the sensitivity analysis of computer experiments with functional inputs. The methods result in one uncertainty index for each functional input as a whole and thus do not give insight into the sensitivity with respect to changes at specific intervals of the functional domain.

This chapter presents a general method for the performance of sensitivity analysis in computer experiments that can take functions as input variables, and return a scalar value as output. The aim is not only to discover the sensitivity of those functional variables as a whole but to identify relevant regions in the functional domain. A very economical sequential approach is presented, which reduces the functional space to a scalar one and makes use of methods from the so-called group factor screening. It results in a descriptive graphical representation of the functional sensitivities. The chapter starts with a motivational example. Then the methodology is presented, followed by properties and graphical presentation. Afterwards, the sequential design strategy for functional input is described. The chapter is finished with some advice on the practical implementation, including the R implementation as well as a detailed presentation of the procedure on an analytical example.

The results of the chapter are based on the contribution “Sequential designs for sensitivity analysis of functional inputs in computer experiments” by Fruth, Roustant, and Kuhnt (2014a).

Motivational example

To illustrate the problem, an application in sheet metal forming is introduced, which originally motivated this research as part of the PhD thesis. The complete analysis of the problem along with a detailed description follows later in Chapter 6.2. Here, it suffices to say that the aim is to analyze the sensitivity of a scalar output, springback, for two inputs that can be changed during the forming process: friction and blankholder

force. The corresponding computer experiment is very time-consuming, which is even enhanced through the functional variation of the two parameters.

A few questions arise at this point. Does the effort of varying blankholder force and friction during the process indeed change anything in the outcome? For example, might it be that only the area under each input curve, and thus the mean of the settings over the time, has an influence? In this case, the functional approach would be useless as it would not give more information than the scalar analysis, where each input variable is kept constant during the run.

To check this issue, a small preview study on a few simple runs is executed (Fig. 4.1). Each row shows the setting of the parameter friction of one run of the computer experiment. No other parameter is changed over the runs, so that the functional behavior of the friction can be analyzed exclusively. In the first two runs, the friction is constant over the 15 seconds of the punch travel, leading to a springback of 7.1 for a low and 2.2 for a high value. In the following 4 runs, varying curves are chosen such that the area below each curve is the same in each run. It can be seen that varying the friction can dramatically reduce springback. In addition, although the four last runs have equal mean friction over time, very clear differences in the springback results are obtained. This indicates the potential that lies in the functional approach and motivates an exploration of the function influence of the parameters.

Notation and situation

For the handling of functional inputs, the notation is slightly modified. An underlying model f is considered with d_{scal} scalar input variables $x_i \in [-1, 1]$, $i = 1, \dots, d_{\text{scal}}$ and d_{fun} functional input variables $g_j : D_j \mapsto [-1, 1]$, $j = 1, \dots, d_{\text{fun}}$. The output stays scalar, $Y \in \mathbb{R}$. The input parameters, scalar as well as functional, are bounded to fall into $[-1, 1]$. The input functions are defined on the domain $D_j = [0, 1]$ for each $j = 1, \dots, d_{\text{fun}}$ and depend on a single argument, not necessarily the same for all

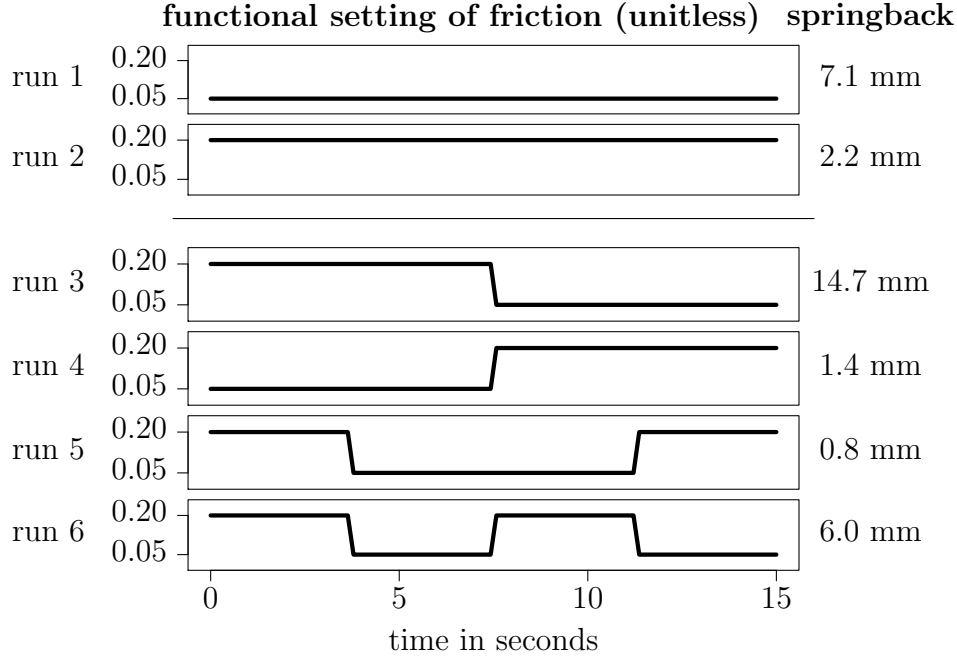


Figure 4.1: Preview runs. Comparison of different functional friction settings.

functional inputs. There are no further conditions on the functions, that is they can be designed freely with no assumptions on the shape or the smoothness of the functions.

The input parameters are connected to Y by f via $f : [-1, 1]^{d_{\text{scal}}} \times \mathcal{F}_{[0,1]}^{d_{\text{fun}}} \mapsto \mathbb{R}$,

$$Y = f(x_1, \dots, x_{d_{\text{scal}}}, g_1, \dots, g_{d_{\text{fun}}}),$$

where $\mathcal{F}_{[0,1]}$ denotes the space of all functions on $[0, 1]$.

4.1 Sequential functional sensitivity analysis

Modeling of input functions

With functional inputs, the input dimension becomes infinite dimensional while at the same time the computer experiment becomes more time-consuming. The functions

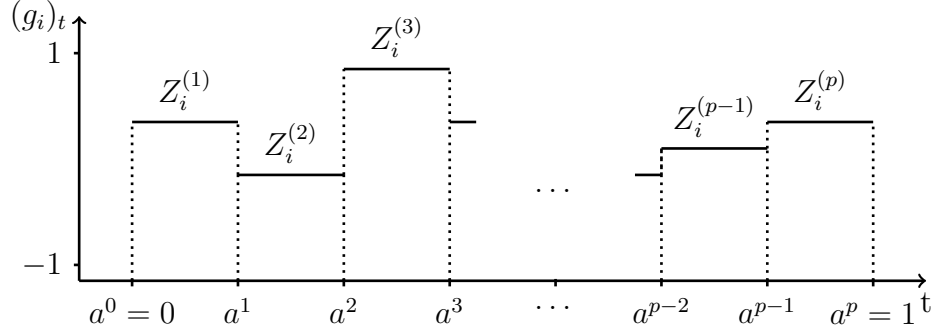


Figure 4.2: Visualization of the piecewise constant representation of functional inputs.

shall therefore be constructed within a flexible framework that reduces the input space considerably, and that eventually allows for the identification of important regions of the functional domain. The idea is to use piecewise constant functions as shown in Fig. 4.2, that is to restrict the input functions to the space of piecewise constant functions. Say we have a decomposition of the domain of each functional input g_j into $p_j \in \mathbb{N}^+$ subintervals at splitting points $\mathbf{a}_j = (a_j^0, \dots, a_j^{p_j})$ with $0 = a_j^0 < a_j^1 < \dots < a_j^{p_j-1} < a_j^{p_j} = 1$,

$$D_j = [a_j^0, a_j^1[\cup [a_j^1, a_j^2[\cup \dots \cup [a_j^{p_j-1}, a_j^{p_j}].$$

We restrict each g_j to belong to $V_{\mathbf{a}_j}$, the space of piecewise constant functions over the subintervals defined by $\mathbf{a}_j$,

$$V_{\mathbf{a}_j} = \left\{ Z_j^{(1)} \mathbb{1}_{[0, a_j^1[}(t) + \dots + Z_j^{(p_j)} \mathbb{1}_{[a_j^{p_j-1}, 1]}(t), \text{ with } Z_j^{(k)} \in [-1, 1], 1 \leq k \leq p_j \right\}. \quad (4.1)$$

The specific procedure of how to choose the split points is provided later in Section 4.3.

With this representation, every functional input $g_j \in V_{\mathbf{a}_j}$ is identified by the random variables $Z_j^{(1)}, \dots, Z_j^{(p_j)}$. Hence, the input dimension is transformed from functional to scalar space,

$$\begin{aligned} Y &= f(x_1, \dots, x_{d_{\text{scal}}}, g_1, \dots, g_{d_{\text{fun}}}) \\ &= \tilde{f}_{a^1, \dots, a^{p_{d_{\text{fun}}}}} \left(x_1, \dots, x_{d_{\text{scal}}}, Z_1^{(1)}, \dots, Z_1^{(p_1)}, \dots, Z_{d_{\text{fun}}}^{(1)}, \dots, Z_{d_{\text{fun}}}^{(p_{d_{\text{fun}}})} \right). \end{aligned}$$

The piecewise constant representation can be viewed in the context of functional representation, see e.g. Ramsay and Silverman (1997), Chapter 3. B-Splines, piecewise polynomial functions with certain knots as breakpoints, are a popular basis to model functions (de Boor, 2001). They cover various types of functions and are always bounded. Our functional framework (4.1) can be regarded as a B-spline of order 1, a linear combination of piecewise constant functions. In addition, the transformation is connected to wavelet theory, more precisely to Haar Wavelets (see e.g. Walker (1999)), where again the basis functions are piecewise constant functions but scaled and shifted to the desired functional form. Wavelets and our framework resemble each other also in the sequential splitting procedure described later in Section 4.3, where sequences of embedded spaces $V_{\mathbf{a}_j} \subseteq V_{\mathbf{a}'_j} \subseteq V_{\mathbf{a}''_j} \subseteq \dots$ are considered. However, the objectives of this sequential splitting are different, as in our approach, the interest lies in the time localization, the identification of influential time points, not in the time-frequency localization.

Remark: Scalar inputs The design of scalar inputs does not have to be mentioned separately, as the space of piecewise constant functions also includes scalar inputs. A scalar input can be considered as a functional input over only one interval $[0, 1]$ which is never split. Thus, from now on, $d = d_{\text{scal}} + d_{\text{fun}}$ is simply used to denote the number of inputs.

Sensitivity Indices

The defined framework allows to perform sensitivity analysis on the values $Z_j^{(k)}$ in

$$Y = \tilde{f}_{a^1, \dots, a^{p_d}} \left(Z_1^{(1)}, \dots, Z_1^{(p_1)}, \dots, Z_d^{(1)}, \dots, Z_d^{(p_d)} \right),$$

leading to a clear and easily understandable sensitivity analysis of intervals. The specific sensitivity method can be chosen according to the evaluation costs and complexity

of the problem from the methods described in Chapter 2. However, since the functional computer experiment is assumed to be expensive, regression analysis as a very economical method for sensitivity analysis is recommended and explored further in this work. The following linear model is considered,

$$Y = \alpha + \sum_{j=1}^d \sum_{k=1}^{p_j} \beta_j^{(k)} Z_j^{(k)} + \varepsilon, \quad (4.2)$$

with α and $\beta_j^{(k)}$ the coefficients to be estimated using a design matrix $\mathbf{Z}$ on the values of $Z_j^{(k)}$ sufficient for estimation, i.e. such that $\mathbf{Z}'\mathbf{Z}$ is invertible (see e.g. Saltelli et al. (2000), p. 124). The estimates $\hat{\alpha}$ and $\hat{\beta}_j^{(k)}$ are obtained by the Least Squares formula $(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{y}$, with $\mathbf{y}$ the vector of the output realizations. As the underlying computer experiment is deterministic, no assumptions are made about ε , the difference between response and regression model and the Least Squares approach is interpreted as simple curve fitting. The coefficients $\hat{\beta}_j^{(k)}$ can then be used as sensitivity indices as they indicate the linear influence of the corresponding functional interval on the output. Additionally, higher order effects of interest can be estimated. One might be especially interested in the second-order interactions between intervals of the same functional input, which leads to the model

$$Y = \alpha + \sum_{j=1}^d \sum_{k=1}^{p_j} \beta_j^{(k)} Z_j^{(k)} + \sum_{j=1}^d \sum_{1 \leq k < k' \leq p_j} \beta_j^{(k,k')} Z_j^{(k)} Z_j^{(k')} + \varepsilon. \quad (4.3)$$

Graphical representation and normalization

The functional indices can be represented by plotting them over the functional domain of each functional input, which visualizes the functional behavior of each input in a compact way. An artificial example of such a visualization is depicted in Fig. 4.3 (top left). It shows the regression coefficients of a single functional input, the bar colors highlight direction and size of the index. A positive bar size indicates that a raise of the input at this interval causes an increase in the mean output, a negative size causes

a decrease. In the specific plot, a small negative influence is visible for the first half of the functional domain and a larger positive influence in the second half. In a study with more than one functional input, one such plot is drawn for each input. Scalar inputs can be depicted as single bars.

The drawback of using $\widehat{\beta}$ directly as functional sensitivity index is that the index values are not independent of the current splitting $\mathbf{a}_j = (a_j^0, \dots, a_j^{p_j})$. If we change the splitting so that the last two intervals are gathered, this larger interval now has a much stronger influence than the two single intervals before, since a larger part of the functional input is varied, which leads to a larger change in the output. Figure 4.3 (top right) shows this situation. In contrast to the previous plot, the plots indicates an increasing influence towards the end of the functional domain.

Independence of the chosen decomposition can be achieved by normalization through the interval size.

Definition 4.1. Consider a set of splitting points $\mathbf{a}_1, \dots, \mathbf{a}_{p_d}$ and assume that $g_j \in V_{\mathbf{a}_j}, j = 1, \dots, d : g_j(t) = \sum Z_j^{(k)} \mathbb{1}_{[a_j^{k-1}, a_j^k]}(t)$. Denote by $\widehat{\beta}_j^k$ and $\widehat{\beta}_j^{(k,k')}$ the estimated first-order and second-order regression coefficients, then the normalized regression index of $Z_j^{(k)}$ is defined by

$$\widehat{H}_j^k = \frac{\widehat{\beta}_j^{(k)}}{a_j^k - a_j^{k-1}},$$

and the normalized interaction regression index of $Z_j^{(k)}$ and $Z_j^{(k')}$ by

$$\widehat{H}_j^{k,k'} = \frac{\widehat{\beta}_j^{(k,k')}}{(a_j^k - a_j^{k-1})(a_j^{k'} - a_j^{k'-1})}$$

for $j \in \{1, \dots, d\}, 1 \leq k < k' \leq p_j$.

In Fig. 4.3, the plot at the bottom shows the corresponding normalized regression indices to the artificial example. The behavior at the end now corresponds to the

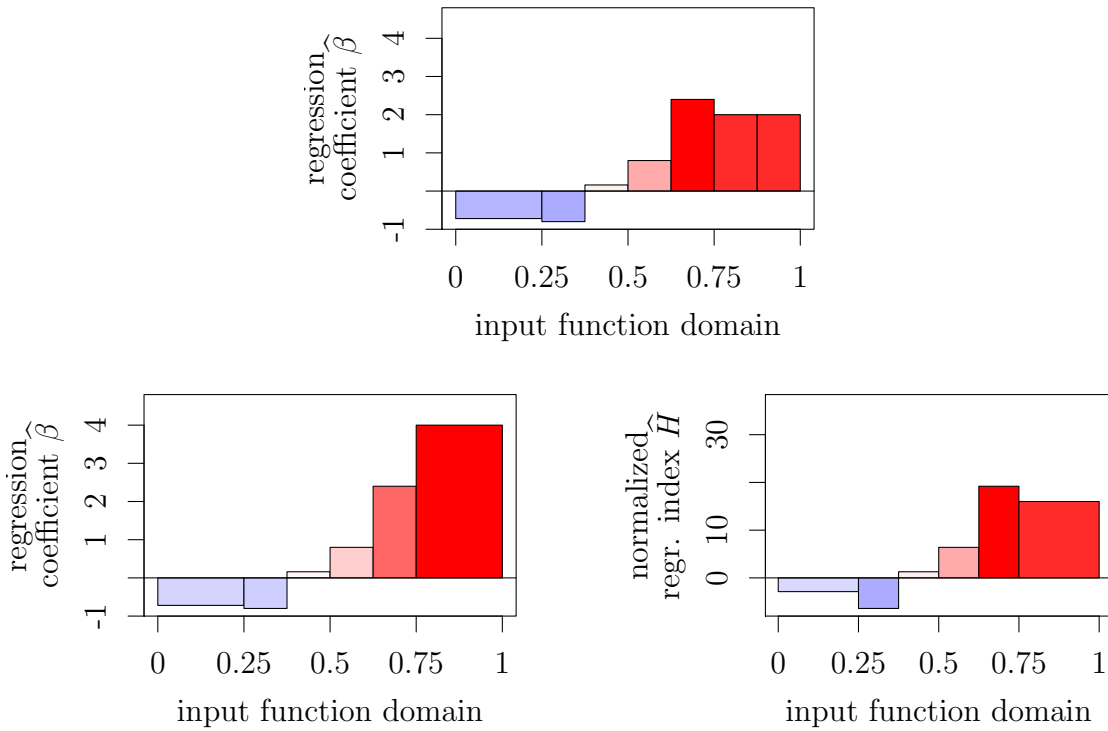


Figure 4.3: Graphical representation of the functional sensitivity of an artificial example, via regression coefficients $\hat{\beta}$ (top and bottom left, for different splittings) and via normalized regression coefficients $\hat{H}$ (bottom right). The bar colors highlight size and direction of the indices. Unnormalized, the regression coefficients are sensitive to the splitting.

picture on the top. In addition, the bar of the first interval is now smaller, as its size is larger and thus was not comparable to the other intervals in the upper plots.

As a regression coefficient $\hat{\beta}$ can be interpreted as the change in the output when the functional input is increased over the corresponding interval, $\hat{H}$ can be interpreted as the change in the output when increasing the functional input over one unit of the functional domain, e.g. one time unit for temporal inputs. Yet, care has to be taken in the interpretation. In an extreme case, the joint interval of two intervals with influences

of different signs could show zero influence when the intervals cancel each other out. Any interpretation of the indices can only relate to the mean influence of the interval, as can be seen in the following test cases.

4.2 Error-free test cases

In order to show that the normalized indices fulfill desirable properties, their performance on some error-free test cases for the underlying model f shall be examined. Each test case depends on only one functional input g and follows a purely linear behavior.

Proposition 4.1. *Let f be the integral over the functional input g weighted by an integrable function $w : [0, 1] \mapsto \mathbb{R}$*

$$f(g) = \alpha + \int_0^1 w(t)g(t) dt,$$

with $\alpha \in \mathbb{R}$. Then, for any splitting $\mathbf{a} = (a^0, \dots, a^p)$ of the functional domain of g , the normalized regression indices return the mean values of the weight function over the intervals, and each interaction index is zero,

$$\widehat{H}^k = \frac{\int_{a^{k-1}}^{a^k} w(t) dt}{a^k - a^{k-1}} \quad \text{and} \quad \widehat{H}^{k,k'} = 0, \quad k, k' = 1, \dots, p.$$

Asymptotically, if the interval size goes to zero, the normalized regression indices return the exact value of the weight function,

$$\lim_{\Delta \rightarrow 0} \frac{\int_t^{t+\Delta} w(u) du}{\Delta} = w(t).$$

Proof. As g is bounded and w is integrable, wg is integrable. When we insert the piecewise constant function representation

$$g(t) = \sum_{k=1}^p Z^{(k)} \mathbb{1}_{[a^{k-1}, a^k[}(t)$$

into $f(g)$, we obtain

$$\begin{aligned} f(g) &= \alpha + \int_0^1 w(t) \sum_{k=1}^p Z^{(k)} \mathbb{1}_{[a^{k-1}, a^k[}(t) dt \\ &= \alpha + Z^{(1)} \int_{a^0}^{a^1} w(t) dt + \cdots + Z^{(p)} \int_{a^{p-1}}^{a^p} w(t) dt \\ &= (\mathbf{1}, Z^{(1)}, \dots, Z^{(p)}) \eta, \end{aligned}$$

with $\eta = \left(1, \int_{a^0}^{a^1} w(t) dt, \dots, \int_{a^{p-1}}^{a^p} w(t) dt \right)'$. Now in the experiments, recalling that the design matrix $\mathbf{Z}$ on the values of $Z^{(1)}, \dots, Z^{(p)}$ is such that $(\mathbf{Z}'\mathbf{Z})$ is invertible, and denoting by $\mathbf{y}$ the corresponding output, we get $y = \mathbf{Z}\eta$. The Least Squares coefficients are then given by

$$\hat{\beta} = (\mathbf{Z}'\mathbf{Z})^{-1} \mathbf{Z}'\mathbf{y} = \eta,$$

which leads to

$$\hat{H}^k = \frac{\int_{a^{k-1}}^{a^k} w(t) dt}{a^k - a^{k-1}}.$$

As there are no interaction terms in $\mathbf{Z}$, we get $\hat{H}^{k,k'} = 0$, $k, k' = 1, \dots, p$. □

In the same way it can be shown that interactions are recovered.

Proposition 4.2. *Let f be the product of two integrals of g over two different intervals $[a^{i^*-1}, a^{i^*}]$ and $[a^{j^*-1}, a^{j^*}]$ with $c, \alpha \in \mathbb{R}$*

$$f(g) = \alpha + c \int_{a^{i^*-1}}^{a^{i^*}} g(t) dt \times \int_{a^{j^*-1}}^{a^{j^*}} g(t) dt.$$

Then for any splitting $\mathbf{a}$ of the functional domain of g with $\{a^{i^-1}, a^{i^*}, a^{j^*-1}, a^{j^*}\} \in \mathbf{a}$, the estimates of the indices and the interaction indices recapture the interactions as they hold*

$$\hat{H}^k = 0, k = 1, \dots, p \quad \text{and} \quad \hat{H}^{k,k'} = \begin{cases} c, & k = i^*, k' = j^*, \\ 0, & \text{otherwise.} \end{cases}$$

Proof. Similarly to the proof of Prop. 4.1 we get

$$\begin{aligned}
 f(g) &= \alpha + c \int_{a^{i^*-1}}^{a^{i^*}} g(t) dt \times \int_{a^{j^*-1}}^{a^{j^*}} g(t) dt \\
 &= \alpha + c \int_{a^{i^*-1}}^{a^{i^*}} \sum_{k=1}^p Z^{(k)} \mathbb{1}_{[a^{k-1}, a^k[}(t) dt \times \int_{a^{j^*-1}}^{a^{j^*}} \sum_{k=1}^p Z^{(k)} \mathbb{1}_{[a^{k-1}, a^k[}(t) dt \\
 &= \alpha + c \times Z^{(i^*)}(a^{i^*} - a^{i^*-1}) \times Z^{(j^*)}(a^{j^*} - a^{j^*-1}).
 \end{aligned}$$

Then the Least Squares estimation returns $\widehat{\beta}^k = 0$, $k = 1, \dots, p$ and

$$\widehat{\beta}^{k,k'} = \begin{cases} c(a^{i^*} - a^{i^*-1})(a^{j^*} - a^{j^*-1}), & k = i^*, k' = j^*, \\ 0, & \text{otherwise} \end{cases}$$

leading to

$$\widehat{H}^k = 0, k = 1, \dots, p \text{ and } \widehat{H}^{k,k'} = \begin{cases} c, & k = i^*, k' = j^*, \\ 0, & \text{otherwise.} \end{cases}$$

□

Furthermore it can be shown that the normalized regression index is robust against nonlinear transformations of the input g in the computer model.

Proposition 4.3. *If g is transformed by a not necessarily linear, but strictly monotonically increasing function $\zeta : [0, 1] \mapsto \mathbb{R}$*

$$f(g) = \alpha + \int_0^1 w(t) \zeta(g(t)) dt$$

then for a two-level design the values $\widehat{H}^k$ are linear in the mean values of the weight function, and the weighting function still determines the importance of the intervals compared to each other. Thus, there exists a fix $\lambda > 0$ such that

$$\widehat{H}^k = \lambda \frac{\int_{a^{k-1}}^{a^k} w(t) dt}{a^k - a^{k-1}}, \quad k = 1, \dots, p.$$

Proof.

$$\begin{aligned} f(g) &= \alpha + \int_0^1 w(t) \zeta(g(t)) dt, \\ &= \alpha + \zeta(Z^{(1)}) \int_{a^0}^{a^1} w(t) dt + \cdots + \zeta(Z^{(p)}) \int_{a^{p-1}}^{a^p} w(t) dt. \end{aligned}$$

Now for a design on two levels $\{-1, 1\}$, we can define λ and $\kappa \in \mathbb{R}$, such that

$$\zeta(z) = \lambda z + \kappa \quad \text{for the two values } z = -1, z = 1.$$

Since ζ is increasing, we have $\lambda > 0$. It follows for all $Z^{(k)} \in \{-1, 1\}$, $k = 1, \dots, p$ that

$$\begin{aligned} f(g) &= \alpha + (\lambda Z^{(1)} + \kappa) \int_{a^0}^{a^1} w(t) dt + \cdots + (\lambda Z^{(p)} + \kappa) \int_{a^{p-1}}^{a^p} w(t) dt \\ &= \left[\alpha + \kappa \sum_{k=1}^p \int_{a^{k-1}}^{a^k} w(t) dt \right] + \sum_{k=1}^p \left[\lambda \int_{a^{k-1}}^{a^k} w(t) dt \right] Z^{(k)}. \end{aligned}$$

Then in the computations of the coefficients, we obtain

$$\begin{aligned} \widehat{\beta}^k &= \lambda \int_{a^{k-1}}^{a^k} w(t) dt, \quad k = 1, \dots, p \\ \Rightarrow \widehat{H}^k &= \lambda \frac{\int_{a^{k-1}}^{a^k} w(t) dt}{a^k - a^{k-1}}, \quad k = 1, \dots, p. \end{aligned}$$

□

We have now described the functional shape for the design of the input functions, the way to perform sensitivity analysis, and its visualization. In the following, the missing part on how to choose the splitting into the intervals and appropriate design techniques are presented.

4.3 Splitting and design

If a desired splitting is available, the method can be performed right away, using a suitable design matrix $\mathbf{Z}$, e.g. full or fractional factorial design. If not, one aims at discovering the functional domain while investing as few as possible model evaluations. Therefore an economical splitting approach is proposed which is performed in consecutive steps. In each step preceding information is used for the splitting.

For the sake of readability, only one functional input $g : D \mapsto [-1, 1]$ as sole input is considered in this section, i.e. $d = 1$. The approach is easily extended to more inputs by considering them as additional groups. For a specific iteration step r , we slightly change the notation of the splitting points to

$$\mathbf{a}^r = (a^{0,r}, \dots, a^{p^r,r}), \quad 0 = a^{0,r} < a^{1,r} < \dots < a^{p^r-1,r} < a^{p^r,r} = 1$$

the space of piecewise constant functions to

$$V\mathbf{a}^r = \{Z^{(1,r)}1_{[0,a^{1,r}]}(t) + \dots + Z^{(p^r,r)}1_{[a^{p^r-1,r},1]}(t), Z^{(k,r)} \in [0, 1]\},$$

and the normalized indices to

$$\hat{H}^{k-r}, \quad k = 1, \dots, p^k \quad \text{and} \quad \hat{H}_j^{k,k'-r}, \quad 1 \leq k \leq k' \leq p^k.$$

Sequential splitting approach

The approach is based on a family of methods called *group factor screening* (see Watson, G. S. (1961) or for an overview Morris (2006)), a very economical method to screen influential input variables in experiments with a high number of input variables. The basic idea is to group variables, explore the influence of the groups as a whole, and sequentially divide only those groups that are influential. In the literature on group factor screening, there are various ways on how to design the groups at the different steps,

which differ e.g. in their assumptions, orthogonality, the treatment of interactions, or the way of reusing evaluations.

The idea can be transferred to functional sensitivity analysis by interpreting the (infinite) points in the functional domain as individual input variables and the splitting intervals as groups of them. The approach is then to start with a very low number p^1 of intervals $\mathbf{a}^1 = (a^{0,1}, \dots, a^{p^1,1})$, perform sensitivity analysis by choosing a suitable design for the corresponding variables $Z^{(1,1)}, Z^{(2,1)}, \dots, Z^{(p^1,1)}$, perform the experiments, and compute the corresponding indices $\hat{H}^{1-1}, \dots, \hat{H}^{p^1-1}$. Then, for the second step, only those intervals that show a considerable influence on the output are examined further. So the vector of splitting points of the second step $\mathbf{a}^2$ includes all points of $\mathbf{a}^1$ plus the point $\frac{a^{k-1,1} + a^{k,1}}{2}$ for each interval $[a^{k-1,1}, a^{k,1}]$ to be examined. This procedure of estimation and splitting is repeated until the functional domain is sufficiently explored or the maximum budget is reached. The split decisions should be taken in close cooperation with experts. Generally, an interval should be seen as important if its index is bigger than an assumed approximation error.

Remark: Interpretation care It has to be kept in mind that the indices are only mean values over a given interval, as it could be seen in Prop. 4.1. Intervals, even if the sensitivity index is small, should be explored further if a change in sign is suspected.

Design by sequential bifurcation

The application of a specific group factor screening method, the sequential bifurcation (Bettonvil, 1995), shall be presented in detail here. In this method, the evaluations of the different steps are effectively reused in the subsequent steps, resulting in a very economical procedure.

Adapted to functional sensitivity analysis, the procedure is the following. The variables are designed on two extreme levels, encoded with -1 and 1 . It starts with only one

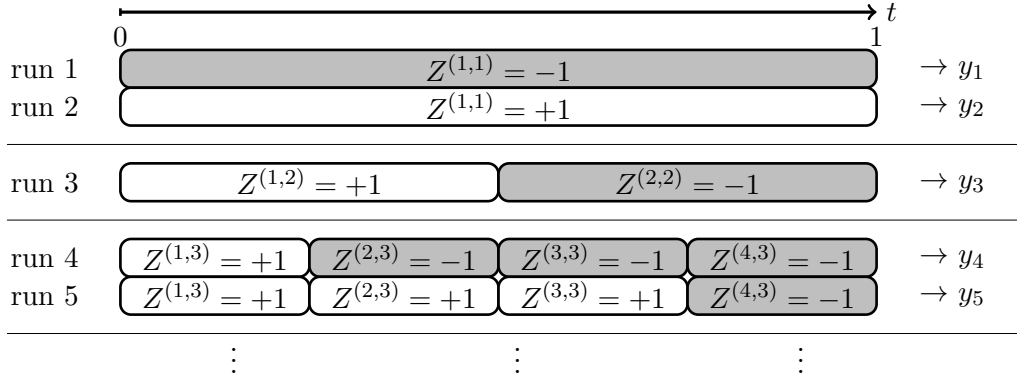


Figure 4.4: Scheme for the design based on sequential bifurcation. White and gray shading indicate the setting to $+1$ and -1 .

interval, corresponding to a constant curve, which is set to -1 and, in a second run, to $+1$. Using the results of these two runs, y_1 and y_2 , the normalized regression index of the whole function can be estimated by $\hat{H}^{1-1} = \frac{0.5(y_2 - y_1)}{1}$. The domain is then split in the middle, resulting in two intervals in the second step. Only one additional run is required to estimate the coefficients of both intervals in which the first interval is set to $+1$ and the second interval to -1 , resulting in y_3 . Values for $(+1, +1)$, $(-1, -1)$ are already known from the previous step, so that the indices of both intervals can be estimated by $\hat{H}^{1-2} = \frac{0.5(y_3 - y_1)}{0.5}$ and $\hat{H}^{2-2} = \frac{0.5(y_2 - y_3)}{0.5}$. Generally, for each split, only one additional run is required to estimate the influence of both new intervals, the one with $+1$ up to the cut and -1 from there. The approach is depicted in Fig. 4.4. It shows three sequential steps requiring a total of 5 runs.

When the presence of interactions has to be considered, additional mirror runs are suggested by Bettonvil (1995), i.e. adding a new run for each run in the design that contains the same settings but with opposite signs. By this, unbiased estimates of the first-order effect coefficients are obtained.

Different designs may be necessary if orthogonality and/or the possibility to estimate interactions are required. Then factorial or fractional factorial designs are good design options. A new such design is then constructed in each step r on all variables of interest

$Z^{(1,r)}, \dots, Z^{(p^r,r)}$, all noninteresting intervals are set to a constant value, e.g. 0, in the design and thus are not regarded in the sensitivity analysis any more. Here again, runs from former steps can be reused. It is easy to show that for full factorial designs in a step where all intervals have been split half of the required runs can be reused. For fractional factorial designs, the reuse possibilities depend strongly on the confounding.

4.4 Implementation

The presented methodology is implemented in the R package **seqSAFI** (Fruth and Jastrow, 2014). The package allows for all described steps of sequential design, modeling and plotting of the functional sensitivities. See Fig. 4.5 for an overview of the package structure. Its core is an object, **safidesign**. It contains the current design, which can be enhanced to get to the next step of the sequential procedure. The splitting can be performed following the sequential bifurcation algorithm or by any desired design given manually. Mirror runs can be added. At any step, the object can be accessed to obtain a transferable design matrix, and also be plotted. A model can be fit to corresponding output from the computer experiment. The resulting **safimodel** object contains the normalized regression indices computed according to the chosen design. Applying the function **plot** to this object leads to the graphical representation by barplots as shown in Fig. 4.3.

4.5 Analytical example

To show how the method can practically be applied, a small analytical test example is analyzed in this section. For the unknown and expensive to evaluate underlying model, we choose the following function which depends on two functional inputs, g_1 and g_2 ,

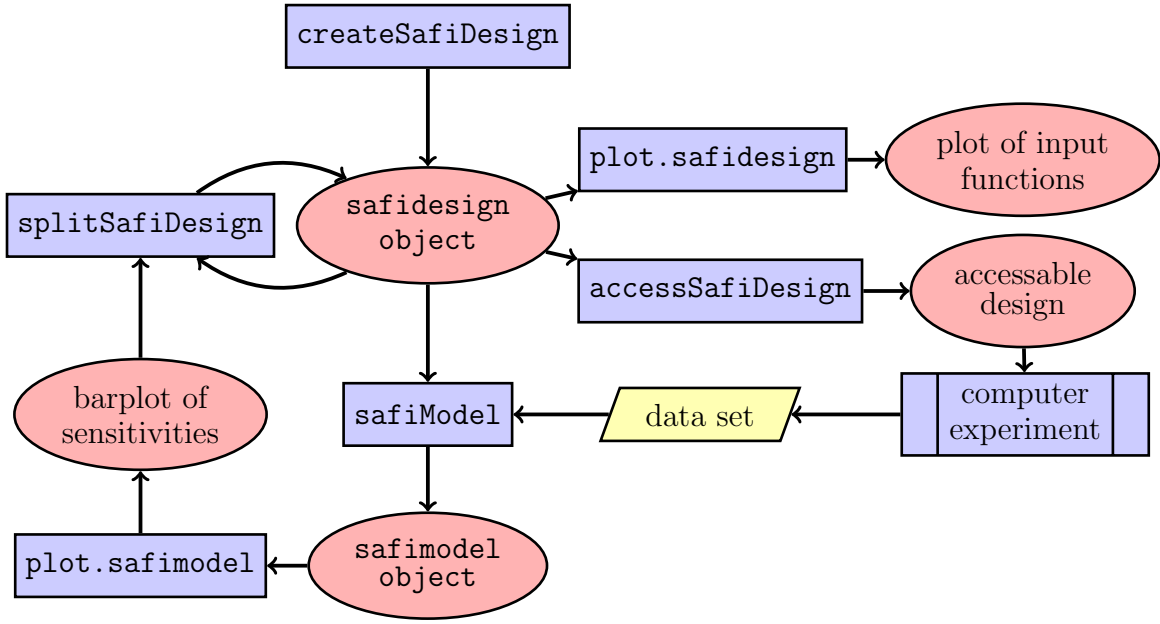


Figure 4.5: Structural setup of R package `seqSAFI`. Functions are represented by rectangles (with double vertical lines if defined externally), data by parallelograms, and output by circles. Arrows represent information transfer.

and one scalar input x .

$$\begin{aligned}
 f(g_1, g_2, x) = & \int_0^1 \left[(3 - 9t) \mathbb{1}_{[0, \frac{1}{3}]}(t) + Z(t) \right] g_1(t) dt \\
 & - \int_0^1 \left[\frac{1}{30} \mathbb{1}_{[\frac{3}{10}, 1]}(t) g_2(t) + 3 \right]^3 dt \\
 & + \frac{8}{10} \sin(x),
 \end{aligned} \tag{4.4}$$

with $Z(\cdot)$ a simulation of a Gaussian process with zero mean, a small variance of 0.05, and Matérn 5/2 covariance function with parameter $\theta = 0.1$. The process $Z(\cdot)$ as well as the third power and the sine function are disturbances meant to achieve a more realistic example. The aim is to discover the linear effects of each input, that is

- the decreasing positive influence of g_1 over $[0, \frac{1}{3}]$,
- the constant negative influence of g_2 over $[\frac{3}{10}, 1]$,
- the positive influence of x .

The computations are performed using the R package `seqSAFI`, presented before. At first, a starting design has to be created. As we want to use as little evaluations as possible and we say that we do not assume active interactions, we use a sequential bifurcation design without mirror runs. Not having any previous knowledge about the functional influence, we start with two equidistant intervals per functional input, i.e with the decompositions

$$\mathbf{a}^1 = \mathbf{a}^2 = \left(0, \frac{1}{2}, 1\right)$$

for g_1 and g_2 respectively. The scalar input x is treated as a functional input on the sole interval $\mathbf{a}^3 = (0, 1)$ in all steps. Thus, a starting design is set up that contains six runs, the initial bifurcation runs (all $+1$ and all -1) plus one run for each splitting which makes two for the splitting into three variables and two for the splitting of each of the two functional domains. The design can be seen on the top in Fig. 4.6. The six runs are executed via Equation (4.4) and the normalized regression indices of Def. 4.1 are computed as described in Section 4.3. They are visualized in the bottom of Fig. 4.6. The first functional input g_1 shows only negligible influence in its second interval. Thus, only the first interval shall be explored further and is split into three intervals. For g_2 , no intervals can be screened out, so that both intervals are split up. This results in four new runs depicted in the top of Fig. 4.7. The resulting indices can be seen in the bottom of the same figure.

The intervals with visible influence of the second step are split again to gain a more detailed look in the functional influence, leading to two new runs for g_1 and three for g_2 . They are shown in Fig. 4.8 together with the resulting indices. We stop the procedure at this point, making Fig. 4.8 the final result. Though they are disturbed and do not show the exact influences, the main linear behavior of the three parameters is revealed in the 15 runs. We can see the decreasing influence of g_1 at the beginning of the domain, though not exactly at the first third, the constant influence of g_2 in the last part, and the positive influence of the scalar input x .

input interval	g_1		g_2		x
	$[0, \frac{1}{2}]$	$[\frac{1}{2}, 1]$	$[0, \frac{1}{2}]$	$[\frac{1}{2}, 1]$	$([0, 1])$
run 1	+1	+1	+1	+1	+1
run 2	-1	-1	-1	-1	-1
run 3	+1	+1	-1	-1	-1
run 4	+1	+1	+1	+1	-1
run 5	+1	-1	-1	-1	-1
run 6	+1	+1	+1	-1	-1

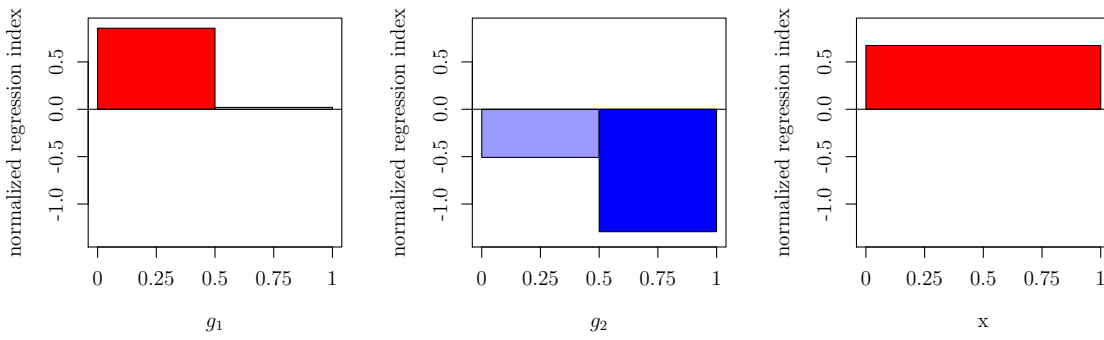


Figure 4.6: Analytical example, step 1, design and normalized regression indices.

Chapter discussion

In this chapter a methodology was presented to perform sensitivity analysis for time-consuming computer codes that take, additional to scalar inputs, functions as input variables. The focus was not only on discovering the sensitivity of the function as a whole but also on exploring the behaviour over the functional domain. The method comprises sequential group factor screening of piecewise constant functions and specially normalized regression indices.

Due to the infinite functional dimension and the small budget of evaluations, a very economical method has been developed that has several limitations. Independence as well as only linear behaviour over the intervals is assumed. Another limitation concerns the interpretation of the indices, as theoretically influences of opposite sign could get

input interval	g_1				g_2				x $([0, 1])$
	$[0, \frac{1}{6}]$	$[\frac{1}{6}, \frac{1}{3}]$	$[\frac{1}{3}, \frac{1}{2}]$	$[\frac{1}{2}, 1]$	$[0, \frac{1}{4}]$	$[\frac{1}{4}, \frac{1}{2}]$	$[\frac{1}{2}, \frac{3}{4}]$	$[\frac{3}{4}, 1]$	
run 7	+1	-1	-1	-1	-1	-1	-1	-1	-1
run 8	+1	+1	-1	-1	-1	-1	-1	-1	-1
run 9	+1	+1	+1	+1	+1	-1	-1	-1	-1
run 10	+1	+1	+1	+1	+1	+1	+1	-1	-1

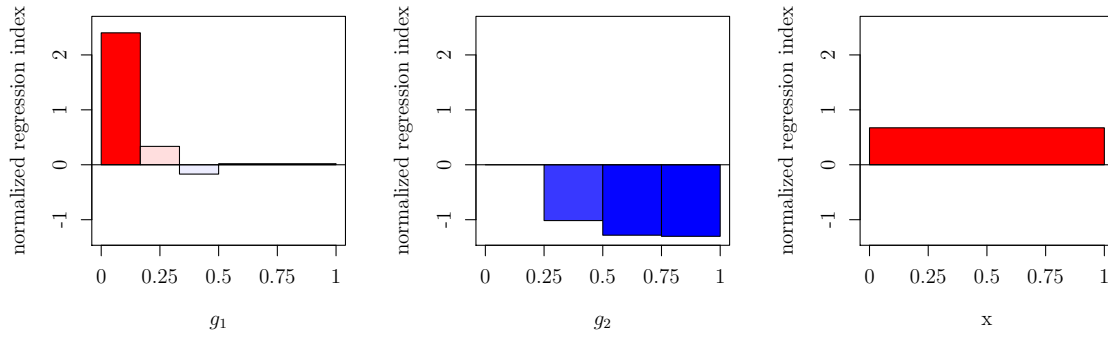


Figure 4.7: Analytical example, step 2, design and normalized regression indices.

	g_1						g_2							x $([0, 1])$
	$[0, \frac{1}{12}]$	$[\frac{1}{12}, \frac{2}{12}]$	$[\frac{2}{12}, \frac{3}{12}]$	$[\frac{3}{12}, \frac{4}{12}]$	$[\frac{4}{12}, \frac{6}{12}]$	$[\frac{6}{12}, 1]$	$[0, \frac{2}{8}]$	$[\frac{2}{8}, \frac{3}{8}]$	$[\frac{3}{8}, \frac{4}{8}]$	$[\frac{4}{8}, \frac{5}{8}]$	$[\frac{5}{8}, \frac{6}{8}]$	$[\frac{6}{8}, \frac{7}{8}]$	$[\frac{7}{8}, 1]$	
run 11	+1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1
run 12	+1	+1	+1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1
run 13	+1	+1	+1	+1	+1	+1	+1	+1	-1	-1	-1	-1	-1	-1
run 14	+1	+1	+1	+1	+1	+1	+1	+1	+1	+1	-1	-1	-1	-1
run 15	+1	+1	+1	+1	+1	+1	+1	+1	+1	+1	+1	+1	-1	-1

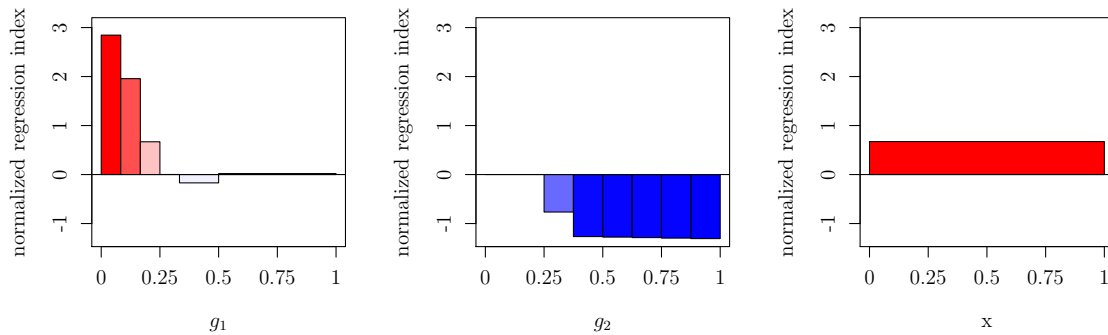


Figure 4.8: Analytical example, step 3, design and normalized regression indices.

canceled out. The method, however, was shown to be robust against some nonlinear transformations and proves to work well in the application in Chapter 6.

The basic idea of this chapter, to explore single intervals and reduce their size more and more, can be transferred to another application of sensitivity analysis, the analysis of the support of input variables, which is presented in the next chapter.

CHAPTER 5

5. SUPPORT ANALYSIS

Minimum and maximum of the input variable distribution are usually specified by the practitioner as extreme values of the parameter in question. Those values are often chosen vaguely, but can have an important impact. As an example, the total effects of the Ishigami function with the usual distributions $X_i \sim U[-\pi, \pi]$ are $D_1^T = 7.72$, $D_2^T = 6.13$, and $D_3^T = 3.37$. If we reduce the support of the distribution of X_3 by 10 percent, $X_3^* \sim U[-\pi + \pi/10, \pi - \pi/10]$, we get $D_1^{T*} = 4.05$, $D_2^{T*} = 6.13$, and $D_3^{T*} = 1.45$. The first and the last variable have clearly lost influence in comparison to the second one. This can be explained by the interaction between X_1 and X_3 , which is much stronger at the borders than in the rest of the function. The index of X_2 stays unchanged as it does not interact with X_3 . The change in the output when changing X_2 stays the same.

As very cheap-to-evaluate metamodels are often used in practice (Section 2.2), it is possible to compute more than the common one or two simple numbers, the Sobol and the total sensitivity index, for each input variable. In this chapter, an index that extends the traditional global sensitivity analysis is defined and examined. It explores the influence of different parts of the support of the input variables by applying the ideas of functional sensitivity analysis from Chapter 4 to a regular scalar situation. In a way, this method is a turn back to local sensitivity analysis (see Section 2.1), where, out of necessity, only the local influence of parameters is studied. The main difference is

that in local sensitivity analysis the influence of a variable is studied in a local space of the other inputs whereas the following method studies the local influence of a variable, but in a global space of the other variables. In addition, the local influences are studied over the whole input space, so no space is ignored. This idea, more specifically the idea to use sensitivity analysis to determine regions in the input space for which the model variation is maximum, is mentioned in Saltelli et al. (2000, p. 6), but to our knowledge not pursued further.

The following method presents a way to show the behavior of the sensitivity for different parts of the input support, so that the effects of the choice of an input distribution can be visible. In addition, it gives further insight into the function by presenting a *function of sensitivities* as extension to Sobol and total sensitivity indices. Following the definition of the support index functions, we show that they are connected to classical indices when the variable in question is restricted to an interval whose size goes to zero. Then, the method is applied to the Ishigami function, followed by a study of their expected value over the support.

5.1 Support index functions

Consider again the standard situation of a black-box function $Y = f(\mathbf{X})$, where $\mathbf{X} = (X_1, \dots, X_d)$ is a vector of independent random variables with distribution $\mu = \mu_1 \otimes \dots \otimes \mu_d$ and f a d -dimensional function $f : \Delta \rightarrow \mathbb{R}$. In addition we assume that $\Delta = [0, 1]^d$ compact, f of class C^1 , and that $X_1, \dots, X_d$ have continuous density functions with support $[0, 1]$. Notice that since f is of class C^1 on the compact set $[0, 1]^d$, f and all its first-order derivatives are bounded, ensuring that $f(\mathbf{X})$ and $\frac{\partial f}{\partial x_\bullet}(\mathbf{X})$ are in L^2 .

The support index functions are then defined as follows.

Definition 5.1. *Support index functions*

The first-order support index $D_i(t)$ of an input variable X_i at a point t , $t \in [0, 1]$, is

defined as the square of the expected value of the first derivative of f ,

$$D_i(t) = \left(E \left(\frac{\partial f}{\partial \mathbf{x}_i}(t, \mathbf{X}_{-i}) \right) \right)^2.$$

The total support index $D_i^T(t)$ of an input variable X_i at a point t , $t \in [0, 1]$, is defined as the expected value of the squared derivative of f ,

$$D_i^T(t) = E \left(\left(\frac{\partial f}{\partial \mathbf{x}_i}(t, \mathbf{X}_{-i}) \right)^2 \right).$$

The support variance $D^{X_i}(t)$ corresponding to a support index of an input variable X_i at a point t , $t \in [0, 1]$, is defined as the overall variance for $X_i = t$,

$$D^{X_i}(t) = \text{Var}(f(t, \mathbf{X}_{-i})).$$

The functions can be evaluated and plotted for a sufficiently large number of discrete points $t = 0, \dots, 1$ over the support. If the gradient of f is known the functions can be calculated directly. If not, the gradient can be approximated by finite differences. For instance the *first-order support index* of X_i at point t using Monte Carlo samples of $\mathbf{X}_{-i}$, $\mathbf{x}_{-i}^{(1)}, \dots, \mathbf{x}_{-i}^{(n)}$, and a small value of δ^* can be estimated by

$$\hat{D}_i(t) = \left(\frac{1}{n} \sum_{k=1}^n \frac{f(t + \delta^*, \mathbf{x}_{-i}^{(k)}) - f(t - \delta^*, \mathbf{x}_{-i}^{(k)})}{2\delta^*} \right)^2.$$

An interpretation of the support functions can be obtained, when we look at the first-order and total influence of the specific point in the support of the variable. Indeed, the functions are connected to the classical index estimation when we restrict X_i to vary only over a small interval around t and let the size of the interval go to zero.

Proposition 5.1. *Let $X_1, \dots, X_d$ be independent random variables with continuous density functions with support $[0, 1]$. Let f be a function defined on $\Delta = [0, 1]^d$ compact and assume that f is of class C^2 , ensuring that $f(\mathbf{X})$, $\frac{\partial f}{\partial \mathbf{x}_\bullet}(\mathbf{X})$, $\frac{\partial^2 f}{\partial \mathbf{x}_\bullet \partial \mathbf{x}_\bullet}(\mathbf{X})$ are in L^2 . Now let one of the random variables, say X_i , be restricted to $X_i^h \sim X_i | X_i \in [t - h/2, t + h/2]$ for $t \in]0, 1[$ and $h > 0$.*

Then we have for the overall variance, the Sobol index of X_i^h , and the total sensitivity index of X_i^h that

1. $\text{Var}(f(X_i^h, \mathbf{X}_{-i})) \xrightarrow{h \rightarrow 0} \text{Var}(f(t, \mathbf{X}_{-i})) = D^{X_i}(t),$
2. $\text{Var}(E[f(X_i^h, \mathbf{X}_{-i})|X_i^h]) / h^2 \xrightarrow{h \rightarrow 0} \frac{1}{12} \left(E \left(\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}) \right) \right)^2 = D_i(t)/12,$
3. $E(\text{Var}[f(X_i^h, \mathbf{X}_{-i})|\mathbf{X}_{-i}]) / h^2 \xrightarrow{h \rightarrow 0} \frac{1}{12} E \left(\left(\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}) \right)^2 \right) = D_i^T(t)/12.$

Proof. Let us first show that

$$\frac{\text{Var}(X_i^h - t)}{h^2} \xrightarrow{h \rightarrow 0} \frac{1}{12}.$$

Let g be the density function of $X_i - t$, $g(x) = f_{X_i}(x + t)$. By assumption, g is continuous in 0, $g(0) > 0$. Consider the density of the truncated variable $X_i^h - t$, $g_h(x) = \frac{g(x)}{\int_{-\frac{h}{2}}^{\frac{h}{2}} g(t) dt} 1_{[-\frac{h}{2}, \frac{h}{2}]}(x)$.

Let $\ell_h = \inf_{x \in [-\frac{h}{2}, \frac{h}{2}]} g(x)$ and $u_h = \sup_{x \in [-\frac{h}{2}, \frac{h}{2}]} g(x)$. With continuity of f in 0, we have $u_h \xrightarrow{h \rightarrow 0} g(0)$ and $\ell_h \xrightarrow{h \rightarrow 0} g(0)$. By standard integral computations we get

$$h\ell_h \leq \int_{-\frac{h}{2}}^{\frac{h}{2}} g(x) dx \leq hu_h \quad \text{and} \quad \frac{h^3}{12}\ell_h \leq \int_{-\frac{h}{2}}^{\frac{h}{2}} g(x)x^2 dx \leq \frac{h^3}{12}u_h$$

so that for $E((X_i^h - t)^2) = \frac{\int_{-\frac{h}{2}}^{\frac{h}{2}} g(x)x^2 dx}{\int_{-\frac{h}{2}}^{\frac{h}{2}} g(x) dx}$ it holds that

$$\frac{\ell_h}{u_h} \frac{h^3/12}{h} \leq E((X_i^h - t)^2) \leq \frac{u_h}{\ell_h} \frac{h^3/12}{h}.$$

By dividing by h^2 this leads to

$$\frac{E((X_i^h - t)^2)}{h^2} \xrightarrow{h \rightarrow 0} \frac{1}{12}.$$

It remains to show that $\frac{E(X_i^h - t)}{h} \xrightarrow{h \rightarrow 0} 0$. As

$$\int_{-\frac{h}{2}}^{\frac{h}{2}} g(x)x dx = \int_0^{\frac{h}{2}} (g(x) - g(-x))x dx$$

it holds that

$$\begin{aligned} |E(X_i^h - t)| &\leq \frac{1}{\ell_h h} \int_0^{\frac{h}{2}} |g(x) - g(-x)| x dx \leq \frac{1}{\ell_h h} \int_0^{\frac{h}{2}} (u_h - \ell_h) x dx \\ &= \frac{u_h - \ell_h}{\ell_h h} \int_0^{\frac{h}{2}} x dx = \frac{u_h - \ell_h}{\ell_h h} \frac{h^2}{8}. \end{aligned}$$

Thus

$$\frac{|E(X_i^h - t)|}{h} \leq \frac{u_h - \ell_h}{8\ell_h} \xrightarrow{h \rightarrow 0} 0.$$

In the following, denote $B_i = \frac{X_i^h - t}{h}$. Remark that B_i is bounded by $|B_i| \leq \frac{1}{2}$ and that $\text{Var}(B_i) \xrightarrow{h \rightarrow 0} \frac{1}{12}$.

Proof of 1. Write the mean value theorem between a real number $x_i \in [t - h/2, t + h/2] \subseteq [0, 1]$ and t

$$f(\mathbf{x}) = f(x_i, \mathbf{x}_{-i}) = f(t, \mathbf{x}_{-i}) + (x_i - t) \frac{\partial f}{\partial x_i}(c, \mathbf{x}_{-i})$$

for a c between x_i and t , thus $c \in [t - h/2, t + h/2]$. Replacing by random variables, we get

$$f(X_i^h, \mathbf{X}_{-i}) = f(t, \mathbf{X}_{-i}) + (X_i^h - t) \frac{\partial f}{\partial x_i}(C, \mathbf{X}_{-i}),$$

where C is a random variable such that $C \in [t - h/2, t + h/2]$. Denoting $Q = B_i \frac{\partial f}{\partial x_i}(C, \mathbf{X}_{-i})$, we then have

$$f(X_i^h, \mathbf{X}_{-i}) = f(t, \mathbf{X}_{-i}) + hQ.$$

Now remark that Q is bounded since both B_i and $\frac{\partial f}{\partial x_i}$ are bounded (by continuity of $\frac{\partial f}{\partial x_i}$ on the compact set $[0, 1]^d$). This implies that $E(f(X_i^h, \mathbf{X}_{-i})) \xrightarrow{h \rightarrow 0} E(f(t, \mathbf{X}_{-i}))$ by the Lebesgue dominated convergence theorem as well as $E(f(X_i^h, \mathbf{X}_{-i})^2) \xrightarrow{h \rightarrow 0} E(f(t, \mathbf{X}_{-i})^2)$. The result follows.

Proof of 2. The proof is similar to the one of 1. Write the Taylor-Lagrange expansion of f between a real number $x_i \in [t - h/2, t + h/2] \subseteq [0, 1]$ and t

$$f(\mathbf{x}) = f(x_i, \mathbf{x}_{-i}) = f(t, \mathbf{x}_{-i}) + (x_i - t) \frac{\partial f}{\partial x_i}(t, \mathbf{x}_{-i}) + \frac{1}{2}(x_i - t)^2 \frac{\partial^2 f}{\partial x_i^2}(c, \mathbf{x}_{-i})$$

for a c between x_i and t , thus $c \in [t - h/2, t + h/2]$. Replacing by random variables, we have

$$f(X_i^h, \mathbf{X}_{-i}) = f(t, \mathbf{X}_{-i}) + hB_i \frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}) + h^2 \tilde{R},$$

with $\tilde{R} = \frac{1}{2} B_i^2 \frac{\partial^2 f}{\partial x_i^2}(C, \mathbf{X}_{-i})$, where C is a random variable, $C \in [t - h/2, t + h/2]$. Then, using the independence between X_i^h and $\mathbf{X}_{-i}$, we have

$$\mathbb{E} [f(X_i^h, \mathbf{X}_{-i}) | X_i^h] = \beta_0 + h\beta_1 B_i + h^2 R,$$

with $\beta_0 = \mathbb{E}(f(t, \mathbf{X}_{-i}))$, $\beta_1 = \mathbb{E}\left(\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i})\right)$, and $R = \mathbb{E}[\tilde{R} | X_i^h]$. Finally

$$\text{Var}(\mathbb{E}[f(X_i^h, \mathbf{X}_{-i}) | X_i^h]) = h^2 \beta_1^2 \text{Var}(B_i) + 2h^3 \beta_1 \text{Cov}(B_i, R) + h^4 \text{Var}(R).$$

Now remark that $\tilde{R}$ is a bounded random variable by continuity of $\frac{\partial^2 f}{\partial x_i^2}$ on the compact set $[0, 1]^d$, and thus R is bounded as well. This implies that $\text{Var}(R) = O(1)$ and $\text{Cov}(B_i, R) = O(1)$, and the result follows.

Proof of 3. With the same notations as in the proof of 2., we have

$$\mathbb{E}[f(X_i^h, \mathbf{X}_{-i}) | \mathbf{X}_{-i}] = f(t, \mathbf{X}_{-i}) + h\mathbb{E}(B_i) \frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}) + h^2 \mathbb{E}[\tilde{R} | \mathbf{X}_{-i}]$$

and thus

$$f(X_i^h, \mathbf{X}_{-i}) - \mathbb{E}[f(X_i^h, \mathbf{X}_{-i}) | \mathbf{X}_{-i}] = h(B_i - \mathbb{E}(B_i)) \frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}) + h^2 S,$$

with $S = \tilde{R} - \mathbb{E}[\tilde{R} | \mathbf{X}_{-i}]$. Since S is centered, we have by independence of B_i and $\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i})$

$$\begin{aligned} \mathbb{E}\left(\left(f(X_i^h, \mathbf{X}_{-i}) - \mathbb{E}[f(X_i^h, \mathbf{X}_{-i}) | \mathbf{X}_{-i}]\right)^2\right) = \\ \beta_2 \text{Var}(B_i) h^2 + 2h^3 \text{Cov}\left((B_i - \mathbb{E}(B_i)) \frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i}), S\right) + h^4 \text{Var}(S), \end{aligned}$$

with $\beta_2 = \mathbb{E}\left(\left(\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i})\right)^2\right)$. The result follows in the same way as in the proof of 2. by using that $\frac{\partial f}{\partial x_i}(t, \mathbf{X}_{-i})$ and S are bounded random variables. $\square$

5.2 Analytical example

Applying the described support analysis to the Ishigami function, $\sin(X_1) + 7 \sin^2(X_2) + 0.1 X_3^4 \sin(X_1)$, reveals the effect observed in the introduction. For the estimation, the finite element estimators $\widehat{D}_i(t)$ and $\widehat{D}_i^T(t)$ were computed at 40 regularly spaced points in $]-\pi, \pi[$ for Monte Carlo samples of size $n = 5\,000$. See Fig. 5.1 for the result. The influence of the input variable X_3 is increasing strongly at the borders, which explains the observed strong change in the total effect, when the support size is changed. Beside this the sinusoidal behavior of X_1 and X_2 can be observed as well as the facts already known from the standard analysis that X_1 and X_3 interact and that the main influence of X_2 is the strongest.

5.3 Expected values

Taking the expected value over the support index functions reveals interesting connections to other indices, which give a little more insight into the interpretation of the functions.

For the total support index $D_i^T(t)$, the expected value is equal to the DGSM ν_i (Section 2.4.2),

$$\mathbb{E} [D_i^T(X_i)] = \mathbb{E} \left[\mathbb{E} \left(\left(\frac{\partial f}{\partial \mathbf{x}_i}(X_i, \mathbf{X}_{-i}) \right)^2 \middle| X_i \right) \right] = \nu_i.$$

It follows that the expected value of the total support index serves as an upper bound for the total sensitivity index,

$$C(\mu_i) \mathbb{E} [D_i^T(X_i)] \geq D_i^T,$$

with $C(\mu_i)$ the Poincaré constant of μ_i as before.

On the other hand, the expected value over the first-order support index function can be connected to variance-based indices. Denote the global mean of the derivative

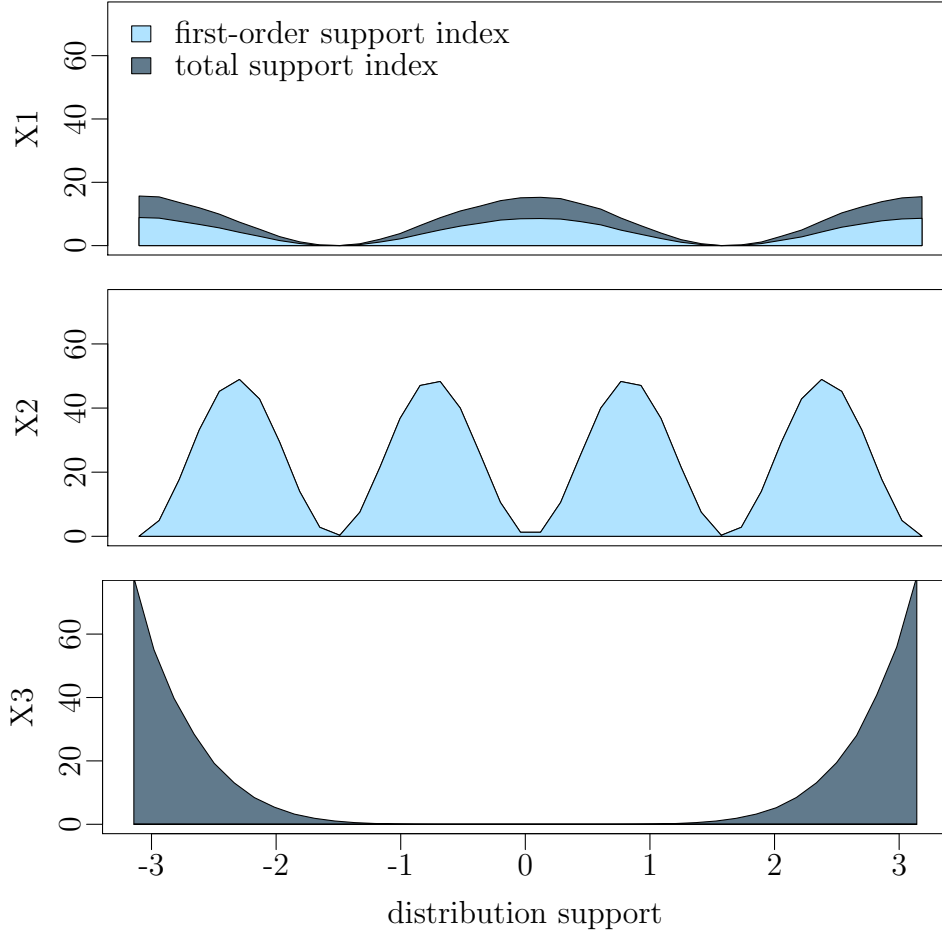


Figure 5.1: Support analysis of the Ishigami function.

function $\frac{\partial f}{\partial x_i}(\cdot)$ by

$$\left(\frac{\partial f}{\partial x_i}\right)_0 = \mathbb{E} \left(\frac{\partial f}{\partial x_i}(\mathbf{X}) \right).$$

Now compute the first-order Sobol index of $\frac{\partial f}{\partial x_i}(\cdot)$ corresponding to the input variable x_i by using the pick-freeze formula,

$$\begin{aligned} D_i \left(\frac{\partial f}{\partial x_i}(\mathbf{X}) \right) &= \mathbb{E} \left[\frac{\partial f}{\partial x_i}(X_i, \mathbf{X}_{-i}) \frac{\partial f}{\partial x_i}(X_i, \mathbf{Z}_{-i}) \right] - \left(\frac{\partial f}{\partial x_i} \right)_0^2 \\ &= \mathbb{E} \left[\left(\mathbb{E} \left[\frac{\partial f}{\partial x_i}(X_i, \mathbf{X}_{-i}) \mid X_i \right] \right)^2 \right] - \left(\frac{\partial f}{\partial x_i} \right)_0^2. \end{aligned}$$

Thus we obtain

$$E[D_i(X_i)] = D_i\left(\frac{\partial f}{\partial x_i}(\mathbf{X})\right) + \left(\frac{\partial f}{\partial x_i}\right)_0^2,$$

the expected value over the first-order support index function can be obtained as first-order index of the derivative function plus its global mean.

Chapter discussion

Two new indices, the first-order support index and the total support index, have been presented that extend the Sobol index and the total sensitivity index, respectively. They return a function of sensitivities that give insight into the local behavior of input variables, and can be especially helpful in the specification of the input distribution. It was shown that the support index functions are connected to the limits of their corresponding scalar indices when the support approaches zero. A link to DGSM and variance-based indices was made.

6. APPLICATION TO A SHEET METAL FORMING PROCESS

The industrial process of deep drawing is a fundamental procedure for forming sheet metal into desired shapes. It is for example extensively applied in the automotive industry in the production of car bodies. In the deep drawing process, a flat sheet metal is pressed with a punch into a die while so-called blankholders keep the metal fixed at the metal borders. Three time points of the forming process are illustrated in Fig. 6.1.

The applications are performed within the Collaborative Research Center SFB 708 at the “Institute of Forming Technology and Lightweight Construction” (IUL) at TU Dortmund University. Several presses are employed in the research, one shown in

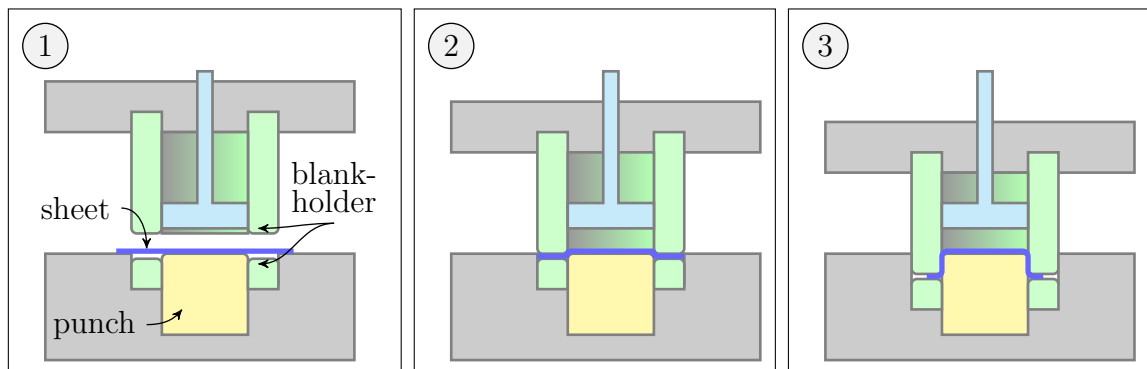


Figure 6.1: Illustration of the deep drawing forming process at three time points.

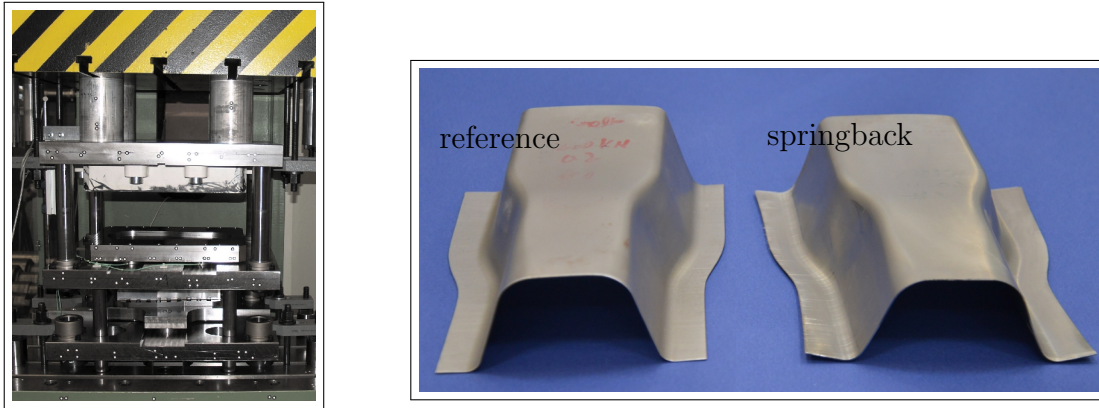


Figure 6.2: Sheet metal forming press at IUL (left), formed parts without and with springback (right).

Fig. 6.2 on the left. Examples of the target workpiece, inspired by the shape of B-pillar, can be seen in Fig. 6.2 on the right. One of the main problems to investigate in the deep drawing analysis is the shape accuracy of the formed part. Tearing or wrinkling can occur during the forming process. Another serious problem is springback, the reforming of the flange after the removal of the punch, which can lead to strong difficulties in the assembly. Such inaccuracy is exemplarily depicted in Fig. 6.2 on the right, springback can be observed in the right piece. The task of analyzing and optimizing the shape accuracy is growing in importance and complexity with the increasing requirements on the formed parts, especially in the automotive industry. In addition, advanced high strength steels are used more and more due to their advantages of light weight and higher strength, which are even more prone to springback.

The forming process can be adjusted by several parameters, which concern for instance the specific properties of the formed material, like thickness or strength, parameters that can be adjusted during the process like the blankholder force or the speed of the punch, or the geometry of the die. It is common practice to adjust and study the process on Finite-Element models, for instance by the popular Software tool AutoForm (AutoForm, 2004). In this work, the more specialized Software LS-DYNA (Livermore

Software Technology Corporation, 2005) is used for the deep drawing simulations. All other computations in the study are performed by the software R (R Core Team, 2014). The aim of the study is to show how the forming process can be analyzed by the methods presented in this work, how they help to gain a deeper insight into the process and improve subsequent analysis techniques.

6.1 Thickness reduction

The first application aims at analyzing the influence of several input parameters on the thickness reduction of the formed part. A high thickness reduction at a region of the sheet metal means a strong local thinning, which can lead to structural deformations or actual tearing. Is it thus of interest to know, which parameters influence the thinning and in which way as well as modeling and optimization. The part under study is the demonstrator part already shown in Fig. 6.2 on the right. Eight input variables are varied in this study, listed in Tab. 6.1 together with their ranges. The variables flow stress, initial sheet thickness, hardening exponent and sheet layout concern the material, blankholder force and friction can be varied during the process. In practice, friction can be changed during the process by adding or removing lubricant with high pressure air and oil removing agents. To roughly integrate this change of friction over the process time, friction is modeled as three independent variables corresponding to three stages of the process time over which friction is kept constant. The more elaborate approach to the analysis of functional inputs of Chapter 4 is applied later in Section 6.2. The output value is the maximal thickness reduction, that is the thickness reduction of the point of the formed part with the strongest thinning. It is given as the ratio between thickness reduction and initial thickness which leads to a scalar value between 0 and 1.

	input variable	feasible region
X_1	flow stress	100-200 MPa
X_2	initial sheet thickness	0.5-1.7 mm
X_3	blankholder force	50-200 kN
X_4	friction; 1 st third of process time	0-0.14
X_5	friction: 2 nd third of process time	0-0.14
X_6	friction: 3 rd third of process time	0-0.14
X_7	hardening exponent	0.1-0.3
X_8	sheet layout	100-150%

Table 6.1: Input variables of the thickness reduction application.

As the duration of the simulation lies between 2 and 6 hours per run, it is out of the scope of a direct analysis. Instead, a Latin hypercube design with 50 runs is applied as initial design for metamodeling. The design and the corresponding thickness reduction results can be found in the Appendix in Tab. B.7. It shows that the output values fall indeed into $[0, 1]$ and are quite well spread over this interval. For the further analysis, the input variables are scaled to lie between 0 and 1. While this does not affect the sensitivity analysis, it simplifies computations and is necessary for the comparisons in the support analysis. On the results, a Kriging metamodel (2.1) with a standard product kernel, a single constant trend, and Matérn 5/2 covariance functions is fit using the R package `DiceKriging` (Roustant et al., 2012). Its prediction function is used as a metamodel.

The first-order Sobol and total sensitivity indices are calculated by the frequency-based methods FAST and EFAST (Section 2.3.3) using a number of $n_{\text{FAST}} = 2000$ sample points. For the computation and visualization the R package `sensitivity` (Pujol et al., 2014) is used. The results are shown in Fig. 6.3. According to this, variable X_8 , the sheet layout, has the strongest influence followed by the two last parts of the friction

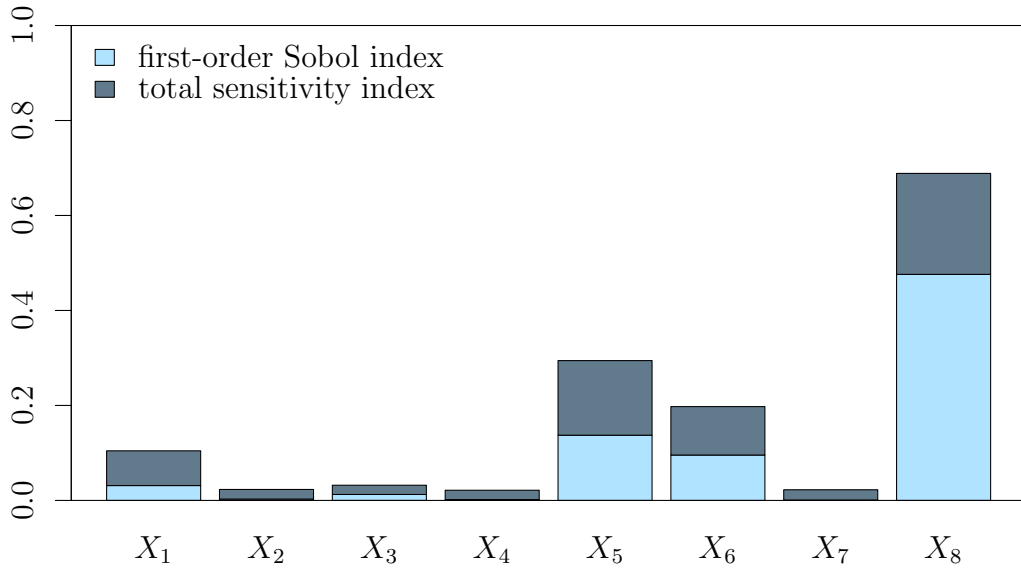


Figure 6.3: Thickness reduction application, first-order Sobol and total sensitivity indices.

and the flow stress. The variables X_2 , X_4 , and X_7 — the initial sheet thickness, the first third of the friction, and the hardening exponent — show only negligible influence and can thus be removed from the analysis. Nevertheless, X_4 , the first third of friction, is kept in order to be consistent in keeping friction as a parameter. The total sensitivity indices show strong interactions for all active parameters. A new Kriging model on the same data but without X_2 and X_7 is set up as updated metamodel.

As we are working with the very fast-to-evaluate Kriging metamodel, we can explore this main effect analysis further by looking at the support indices presented in Chapter 5. This analysis can in particular check if the chosen input ranges influence the results critically. Support indices are estimated for the six remaining variables of interest at 40 equally distributed points $t = 0, \dots, 1$ over the normalized domain via

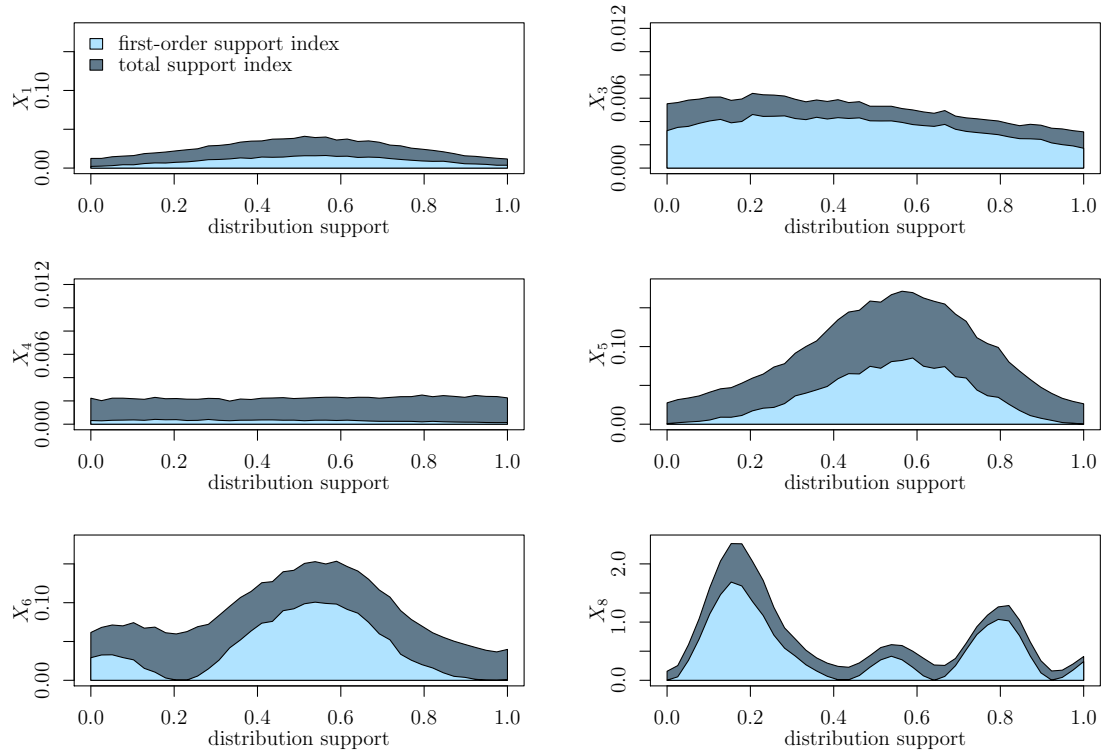


Figure 6.4: Thickness reduction application, support analysis. Note that the vertical axes are not equally scaled in order to visualize the behavior of less influential input variables.

finite differences with a Monte Carlo sample of size 2000 for each point. The resulting support functions are drawn in Fig. 6.4. The plots show that for the most influential variables X_8 , X_6 , X_5 , and X_1 the specific range indeed does not affect the result strongly, as the output variation is not strong at the borders of the input support. The wavy shape of the support index function of variable X_8 reveals strong nonlinearities in the relation between this variable and the thickness reduction. The two friction parameters X_5 and X_6 show similar support functions, underlining that both describe the same parameter. For X_3 and X_2 , the function of the input influence over the domain does not change much, possibly due to their small overall influence. For all variables, the interactions seem to be rather equally spread over the support.

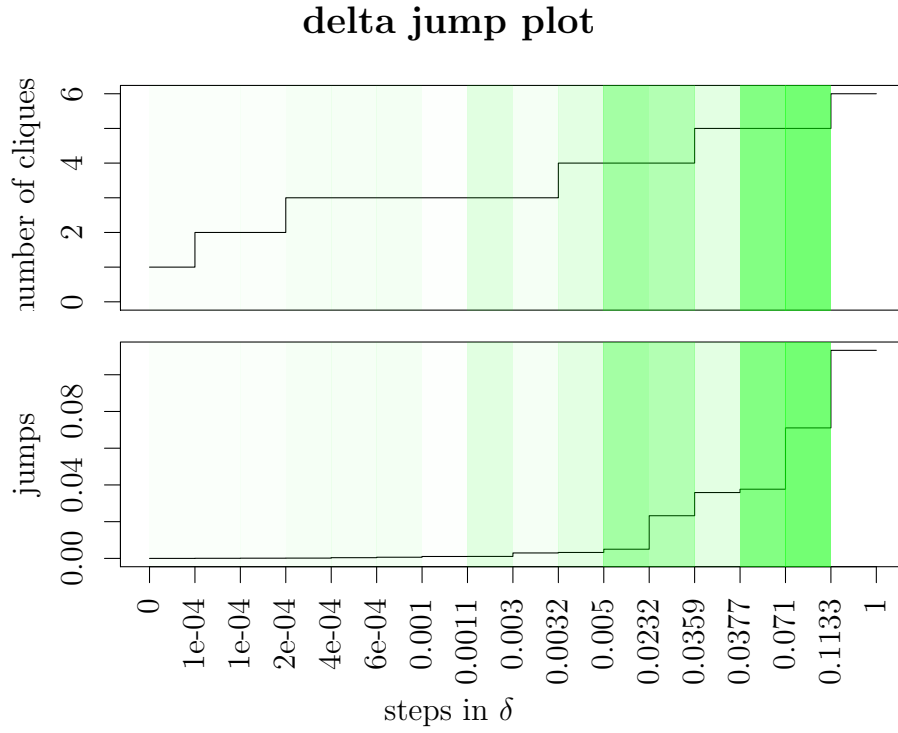
threshold cut	0.000	0.010	0.030	0.050	0.100	1.000
RMSE	0.793	0.585	0.780	0.735	0.786	0.860

Table 6.2: Thickness reduction application, threshold identification with standard product kernel as initial kernel, leave-one-out RMSE values for different thresholds.

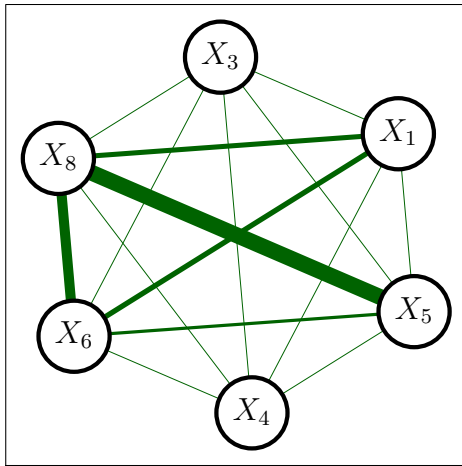
In terms of interaction, the analysis so far has shown that the variables do interact, but not in which way. Thus, the specific structure of the interactions shall be examined closer using the methods presented in Chapter 3 via R package `fanovaGraph`. A FANOVA graph is set up with the TII estimated by the *Liu and Owen method*. A total number of 20 000 evaluations is used. The resulting FANOVA graph can be seen in Fig. 6.5 (b). All interactions are active, although the majority to a very small, possibly negligible, amount. The strongest interaction can be seen between X_5 and X_8 , and generally, X_5 , X_6 and X_8 seem to interact the most.

To apply the graph in subsequent analysis, a threshold to pick out the relevant interactions has to be determined. The *delta jump plot* (see Section 3.6) is shown in Fig. 6.5 (a). It reveals four main jumps around the values 0.01, 0.03, 0.05, and 0.1. The threshold decision procedure by Kriging kernel comparison is performed for those four threshold candidates and 0 and 1, first by using the standard product kernel as initial kernel. The resulting leave-one-out RMSE values can be seen in Tab. 6.2. The table shows that the smallest RMSE can be achieved when removing all interactions with a TII lower than 0.01. Figure 6.5 (c) shows the corresponding FANOVA graph. The variables X_3 , the blankholder force, and X_4 , the first third of friction, are separated, while the other variables fully interact, except for X_1 and X_5 .

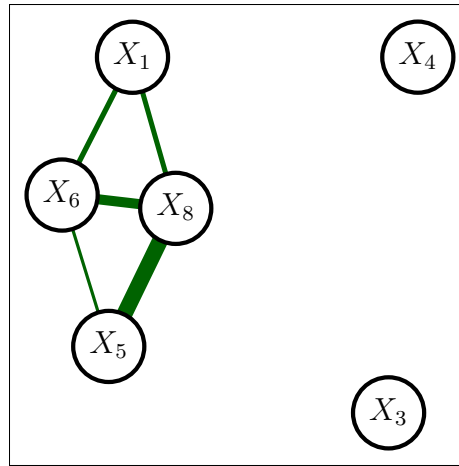
The threshold determination procedure is repeated with the ANOVA kernel by Durand et al. (2013) as initial kernel, which showed superiority in the simulation study in Section 3.6. The results can be seen in Fig. 6.6. With this kernel, the initial unthresholded FANOVA graph only shows three active interactions, with the interac-



(a)



(b)



(c)

Figure 6.5: Thickness reduction application with standard product kernel as initial kernel. Delta jump plot (a). FANOVA graphs of the complete interaction structure (b) and the thresholded values (c).

threshold cut	0.000000	0.000003	0.000200	1.000000
RMSE	0.792732	0.725182	0.723123	0.859808

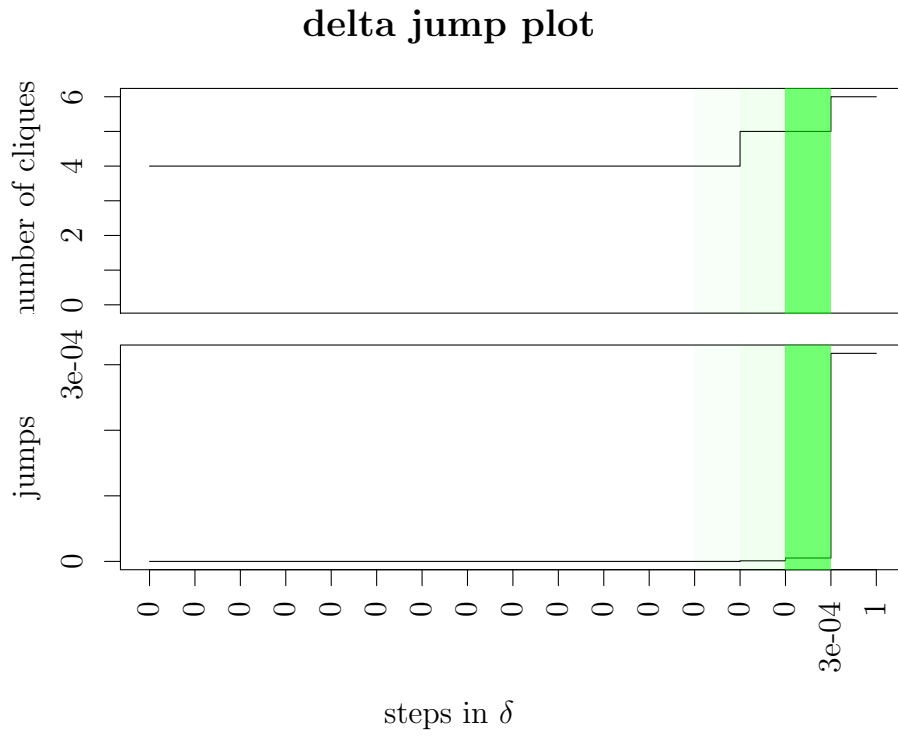
Table 6.3: Thickness reduction application, threshold identification with ANOVA kernel as initial kernel, leave-one-out RMSE values for different thresholds.

tion between X_5 and X_8 being by far the strongest, indicated by a huge jump in the *delta jump plot*. The threshold identification procedure on Kriging kernel comparison (Tab. 6.3) even leads to an emptier graph which contains the interaction between X_5 and X_8 only.

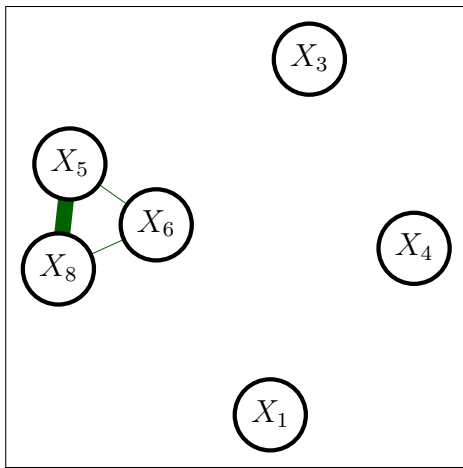
The results of the two kernels differ clearly but at least are consistent in identifying the interaction between X_5 and X_8 as strongest. The ANOVA kernel showed superiority in the simulation study, but in this application, the graph based on the standard product kernel lead to a smaller leave-one-out RMSE value (Tab. 6.2). To obtain more reliable results on the block-additive structure, a larger study on a new test data set seems to be necessary.

Conclusion

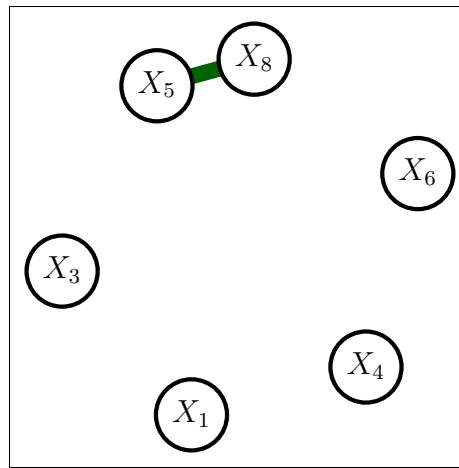
A complete sensitivity analysis of the thickness reduction was performed where the newly developed methods of interaction screening, support analysis, and threshold determination were applied. A strong total interaction was found between the sheet layout and the friction of the second third of the process time. Smaller interactions were identified between the sheet layout and the friction of the last third of the process time and also between both friction variables. For the identification of a thresholded FANOVA graph for subsequent studies, a deeper investigations is necessary.



(a)



(b)



(c)

Figure 6.6: Thickness reduction application with ANOVA kernel as initial kernel. Delta jump plot (a). FANOVA graphs of the complete interaction structure (b) and the thresholded values (c).

6.2 Functional input application

Adjustments to the deep drawing press and to the simulation tools at the IUL allow the variation of the process parameters blankholder force and friction during the forming process, a possibility to further improve the shape accuracy of the formed part by more specialized parameter settings. The effect of functional friction has already been roughly examined in the thickness reduction application where it was separated in three parts considered as independent variables. This analysis shall be extended to a detailed functional sensitivity analysis of friction and blankholder force that gives insight into the functional behavior of both parameters and screens out negligible time intervals for further analysis and optimization. A small previous study showing the usefulness of the functional examination has already been shown in the motivational example in the beginning of Chapter 4. A part of the results is published in ul Hassan et al. (2013).

We apply the methods presented in Chapter 4 for functional sensitivity analysis to a situation with two functional inputs (friction and blankholder force) and no scalar input. This time, the scalar output is represented by the mean springback as a measure of the geometrical accuracy. Three steps of functional sensitivity analysis are performed following the design based on sequential bifurcation (Section 4.3). Mirror runs are added to ensure unbiased estimates, which results in 24 runs in total. See the Appendix for the complete design (Tab. B.8). Figure 6.7 shows the normalized regression indices of the three steps. At first, two intervals are considered for both input variables, which reveals a stronger influence of friction than blankholder force as well as a general tendency towards a positive influence in the first half and negative influence in the second half. At this point, all intervals are interesting to the engineers and thus are split for the second step. There, the last intervals of each half show to have the greatest impact whereas the first and third intervals have relatively small influence. Therefore, in the third step of the sequential analysis, these intervals are not explored further, but

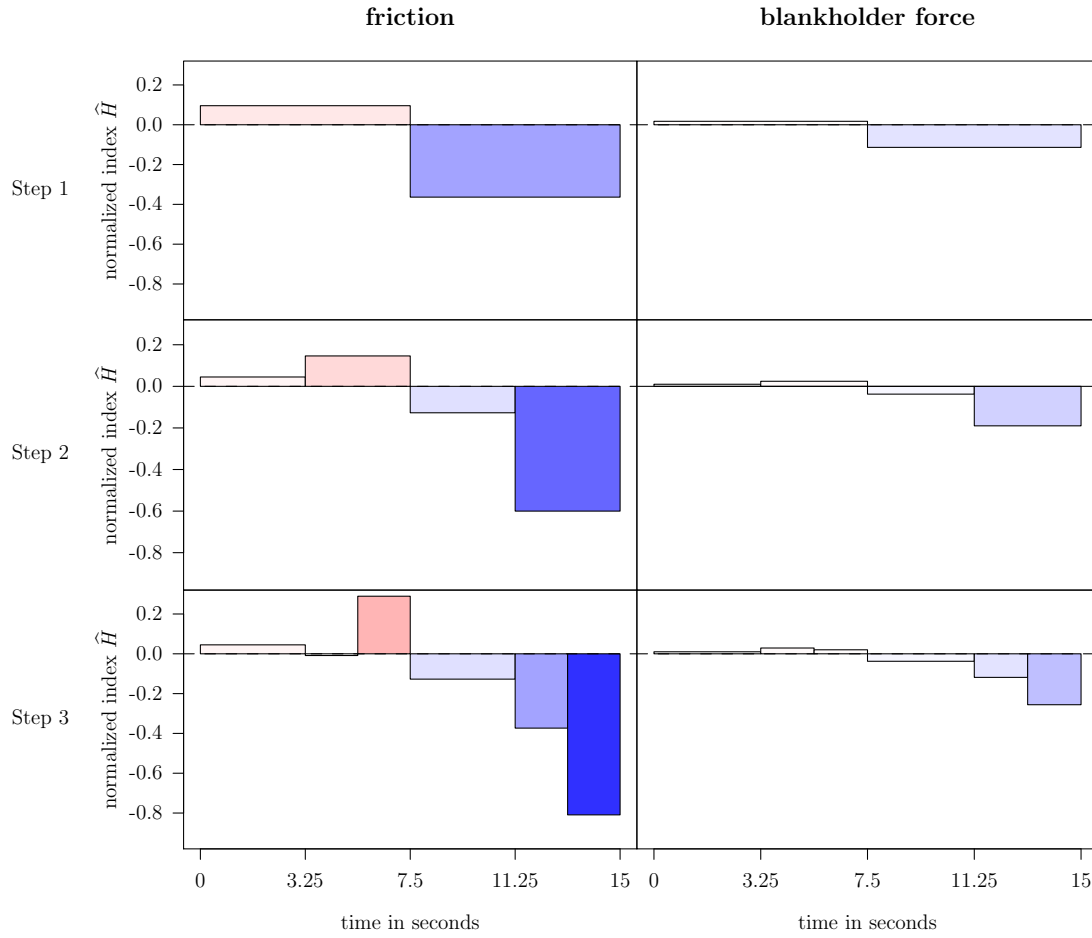


Figure 6.7: Functional input application, normalized regression indices of three steps (top, middle, bottom). The bar color references sign and amplitude of the bar.

each second and fourth interval is split again. After these three steps, the functional sensitivities are explored sufficiently for the engineers.

To confirm the result, a further analysis using factorial designs is conducted, which requires a higher number of runs than sequential bifurcation but allows for the estimation of interactions. For the first step, a full factorial 2^4 -design with 16 runs is chosen. The second and the third step are both performed on a 2^{8-3} -design with 32 runs, in which the confounding structure is chosen such that the confounding of influential time intervals is avoided. The plot of normalized regression indices in Fig. 6.8

shows only a few dissimilarities compared to the sequential bifurcation analysis which overall confirms the chosen model. Table B.9 in the Appendix shows the estimated interaction indices together with the confounding structure. Generally, interactions are stronger for friction than for blankholder force and generally close intervals interact stronger than remote ones. The largest interaction can be observed between the intervals [11.25, 13.125] and [13.125, 15] of friction, which are directly consecutive and which are the two intervals with the strongest main effects. The second largest interaction influence is shared by the interactions between the intervals [3.75, 5.625] and [5.625, 7.5] and the intervals [11.25, 13.125] and [13.125, 15], both corresponding to blankholder force. These intervals are again consecutive and comparatively strong.

When interpreting the general results, at first an overall similar behavior of both variables can be clearly seen, which confirms a presumption of the engineers as both parameters affect the springback in the similar way. A clear break in about half of the time can be observed. It can also be noted that the magnitude of the influence increases towards the end of both halves. These intervals are most important for the springback behavior.

Conclusion

For the essential task of springback compensation in sheet metal forming a functional sensitivity analysis was performed. To the best of our knowledge, this is the first time that the influence of time intervals is systematically investigated in sheet metal forming. From the engineering point of view, the main conclusion is that a high value of blankholder force and friction coefficient for the last 15 percent of the punch travel can strongly reduce springback in the final part. This effect can be explained by the increase in the plastic strains in the part at the end of the process, which reduces the springback. A further insight is that the parameters should have a small value during 50-60 percent of the punch travel as the sheet is being formed over the larger radius

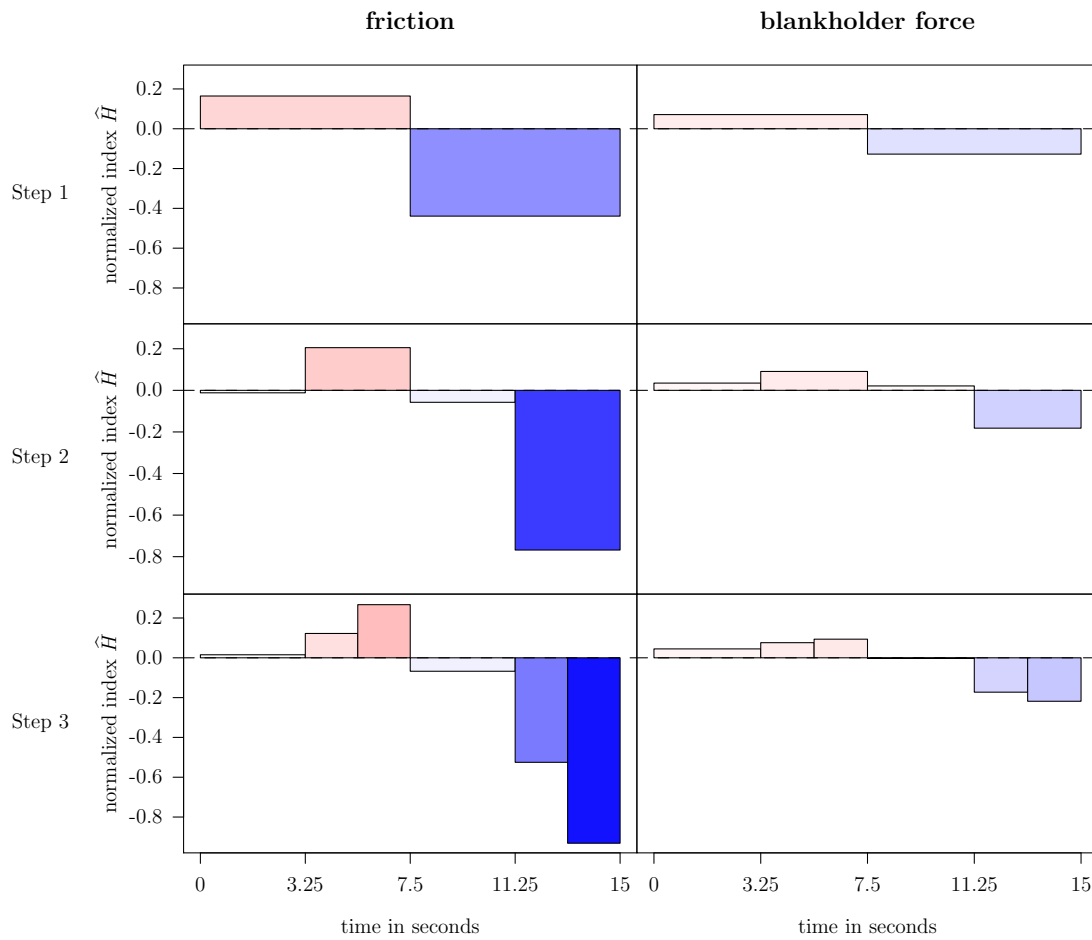


Figure 6.8: Functional input application in three steps via factorial designs. First step (top). Second step (middle). Third step (bottom). The bar color references sign and amplitude of the bar.

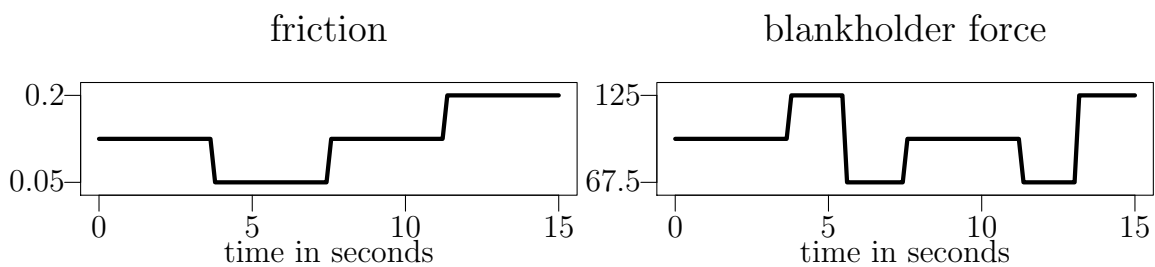


Figure 6.9: Functional input setting that leads to the best springback value of 0.47 mm.

of the punch at that time. The input setting that resulted in the lowest springback of 0.47 mm, run 21 of the factorial design, is pictured in Fig. 6.9. The best run with constant settings, run 2 of the sequential bifurcation design in Tab. B.8, resulted in a springback of 0.74 mm. This reduction of the springback by around 35 percent underlines the benefits of the functional variation approach and gives an impression of its possibilities for subsequent springback optimization.

7. CONCLUSION AND OUTLOOK

The general field of this thesis was sensitivity analysis for black-box functions. Three different concepts that extend the basic sensitivity analysis have been presented. They all share the ability to give deeper insight into the function's behavior on different levels, be it interactions, the domain of functional inputs, or the distribution support of inputs.

The sensitivity analysis for interaction screening enables a user to explore the interaction structure of the function and to find inactive interactions. This provides a block-additive decomposition of the black-box function, which in turn can be of use, for instance in prediction modeling and optimization.

The sensitivity analysis for functional inputs gives insight into the influence of the inputs on the functional domain level. The method explores the influence of different parts of the functional domain, possibly leading to enhanced insights about the influence of the input variables compared to scalar analysis. The screening of inactive parts of the functional domain can again serve as a starting point for modeling and optimization.

The support analysis, finally, provides an extension from single sensitivity indices to sensitivity functions by returning a sensitivity for each point of the distribution support of the input. An insight into the influence of the inputs for different supports can be gained, showing the potential consequences of input distribution choices.

The presented methods were applied in two forming process studies, concerning the thickness of the formed metal and the springback, respectively. They helped to improve the understanding of the forming process and gave more insight into the way input sources influence the deformation yielding an important prerequisite to subsequent research in the field of sheet metal forming.

All three methods provide complete procedures for the exploration of computer experiments, which can lead to new insights and can substantially support further analyses. There are, nevertheless, several possibilities for improvement. For the sensitivity analysis of interaction screening, several extensions are imaginable, e.g. to allow for the employment of dependent variables or to explore metamodels that enable the direct computation of the TIIs, possibly via ANOVA zero-mean kernels or Polynomial Chaos Expansion. In the sensitivity analysis for functional input, aims for future research could be to develop methods that allow for smooth design functions, non-linear influences, and models that avoid the problem of canceling out variables of opposite sign. However, those will be challenging tasks due to the necessary strong dimension reduction and the focus on sensitivity analysis. One further point of improvement is to consider not only the value of the input at the time points but also how the values between the time points evolve, and thus to set up models that incorporate this evolution in time. A straightforward outlook is the extension to spatial input variables. One idea could be to treat each spatial dimension as a functional input and thus lead the problem back to functional analysis. Another idea is to split the input space into simplices similar to the splitting of intervals in the functional analysis.

APPENDIX A

A. NOTATIONS

symbol	description
mathematical notations	
$L^p(\mu)$	space of functions that are p times integrable with respect to measure μ
C^k	class of functions with continuous derivatives up to order k
$\ \mathbf{x}\ $	Euclidean norm over vector $\mathbf{x}$
$\nabla f(\mathbf{X})$	gradient of f
$ Q $	for a set Q , number of elements in Q
$A - B$	$\{x x \in A \text{ and } x \notin B\}$, difference of sets A and B
$\mathcal{F}_{[0,1]}$	space of functions with domain $[0, 1]$
$C(\mu)$	Poincaré constant of a measure μ
defined variable names (general)	
n	size of Monte Carlo sample
d	number of input variables
$\mathbf{x} = (x_1, \dots, x_d)'$	vector of input setting (for one run)
$\mathbf{X} = (X_1, \dots, X_d)'$	random vector input variables
$k(\mathbf{h})$	Kriging covariance kernel for input distance $\mathbf{h}$ (2.3)
δ^*	small real-valued spacing, e.g. in finite differences
I	set of variable indices, subset of $\{1, \dots, d\}$
$\mathbf{x}_I, \mathbf{X}_I$	input setting and variable for all input variables in I
$\mathbf{x}_{-I}, \mathbf{X}_{-I}$	input setting and variable for all input variables but the ones in I
y, Y	specific output value and corresponding random variable
f	underlying function, either computer experiment or metamodel
$\mu = \mu_1 \otimes \dots \otimes \mu_d$	probability measure of $\mathbf{X}$
Δ	domain of f
f_I	additive term of f in FANOVA decomposition (2.4)

D	overall variance of Y (2.7)
D_I, S_I	unscaled (2.8) and scaled (2.9) standard Sobol index
D_I^T, D_I^C	total (2.10) and closed (2.11) sensitivity index
${}^{\text{pf}}\widehat{D}_I^C, {}^{\text{Cor1}}\widehat{D}_I^C,$ ${}^{\text{Cor2}}\widehat{D}_I^C$	closed sensitivity index estimators (Section 2.3.2)
${}^{\text{Jan}}\widehat{D}_I^T$	total sensitivity index estimator by Jansen formula (3.7)
$\widehat{S}_n^1, \widehat{S}_n^2$	normalized closed sensitivity index estimators considered by Janon et al. (2013) (2.24,2.25)
${}^{\text{FAST}}\widehat{D}_i, {}^{\text{FAST}}\widehat{D},$ ${}^{\text{EFAST}}\widehat{D}^T_i, {}^{\text{RBD}}\widehat{D}^T_I$	frequency-based estimators (Section 2.3.3)
ν_i	DGSM (2.30)
defined variable names in Chapter 3	
$\mathfrak{D}_{i,j}$	total interaction index (3.1)
${}^{\text{Jan}}\widehat{\mathfrak{D}}_{i,j}, {}^{\text{RBD}}\widehat{\mathfrak{D}}_{i,j},$ ${}^{\text{pf}}\widehat{\mathfrak{D}}_{i,j}, {}^{\text{LO}}\widehat{\mathfrak{D}}_{i,j},$ ${}^{\text{fix}}\widehat{\mathfrak{D}}_{i,j}$	proposed TII estimators (Section 3.2)
δ	threshold cut (3.12)
$\nu_{i,j}$	crossed DGSM (3.13)
defined variable names in Chapter 4	
$d_{\text{scal}}, d_{\text{fun}}$	number of scalar and functional input variables
$g_1, \dots, g_{d_{\text{fun}}}$	functional input variables
p_j	number of intervals of the domain decomposition of input g_j
$\mathbf{a}_j = (a_j^0, \dots, a_j^{p_j})$	splitting points of the domain decomposition of input g_j
$V_{\mathbf{a}_j}$	space of piecewise constant functions (4.1)
$Z_j^{(k)}$	level of functional input g_j over interval $[a_{k-1}, a_k]$ (4.1)
$\beta_j^{(k)}$	regression coefficients corresponding to Z_j^k (4.2)

$\beta_j^{(k,k')}$	regression coefficient corresponding to the interaction between $Z_j^{(k)}$ and $Z_j^{(k')}$ (4.3)
$\hat{H}_j^k, \hat{H}_j^{k,k'}$	normalized regression index of Z_j^k and normalized interaction regression index of $Z_j^{(k)}$ and $Z_j^{(k')}$ (Def. 4.1)
	defined variable names in Chapter 5
$D_i(\cdot), D_i^T(\cdot)$	first-order and total support index function (Def. 5.1)
$D^{X_i}(\cdot)$	support variance (Def. 5.1)

APPENDIX B

B. DATA

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	0.01050	0.01048	0.00775	0.00775
Jansen	X1*X3	0.00000	-0.00026	0.00191	0.00193
Jansen	X1*X4	0.00000	0.00005	0.00190	0.00190
Jansen	X2*X3	0.00000	-0.00008	0.00201	0.00202
Jansen	X2*X4	0.00000	0.00015	0.00192	0.00193
Jansen	X3*X4	0.01196	0.01203	0.00071	0.00071
RBD-FAST	X1*X2	0.01050	0.01080	0.00556	0.00557
RBD-FAST	X1*X3	0.00000	0.00078	0.00462	0.00469
RBD-FAST	X1*X4	0.00000	0.00041	0.00465	0.00467
RBD-FAST	X2*X3	0.00000	0.00010	0.00537	0.00537
RBD-FAST	X2*X4	0.00000	0.00007	0.00520	0.00520
RBD-FAST	X3*X4	0.01196	0.01370	0.00063	0.00185
Liu-Owen	X1*X2	0.01050	0.01045	0.00067	0.00067
Liu-Owen	X1*X3	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X1*X4	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X2*X3	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X2*X4	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X3*X4	0.01196	0.01200	0.00087	0.00087
pick-freeze	X1*X2	0.01050	0.01099	0.00571	0.00573
pick-freeze	X1*X3	0.00000	0.00039	0.00347	0.00349
pick-freeze	X1*X4	0.00000	0.00020	0.00416	0.00416
pick-freeze	X2*X3	0.00000	0.00041	0.00380	0.00383
pick-freeze	X2*X4	0.00000	0.00035	0.00375	0.00377
pick-freeze	X3*X4	0.01196	0.01230	0.00382	0.00384
fixing method	X1*X2	0.01050	0.01065	0.00004	0.00016
fixing method	X1*X3	0.00000	0.00072	0.00035	0.00080
fixing method	X1*X4	0.00000	0.00072	0.00034	0.00080
fixing method	X2*X3	0.00000	0.00068	0.00030	0.00075
fixing method	X2*X4	0.00000	0.00068	0.00031	0.00075
fixing method	X3*X4	0.01196	0.01280	0.00044	0.00096

Table B.1: Simulation study on TII estimators (Section 3.5), simulation results of test function *function 1*.

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	0.00000	-0.00982	0.34844	0.34858
Jansen	X1*X3	3.37400	3.37224	0.40906	0.40907
Jansen	X2*X3	0.00000	-0.03654	0.24399	0.24671
RBD-FAST	X1*X2	0.00000	0.09898	0.57583	0.58428
RBD-FAST	X1*X3	3.37400	4.45408	0.32655	1.12836
RBD-FAST	X2*X3	0.00000	0.11056	0.36745	0.38373
Liu-Owen	X1*X2	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X1*X3	3.37400	3.36829	0.26011	0.26018
Liu-Owen	X2*X3	0.00000	0.00000	0.00000	0.00000
pick-freeze	X1*X2	0.00000	0.01049	0.28169	0.28189
pick-freeze	X1*X3	3.37400	3.31536	0.38642	0.39084
pick-freeze	X2*X3	0.00000	0.00154	0.31470	0.31471
fixing method	X1*X2	0.00000	0.04436	0.03331	0.05547
fixing method	X1*X3	3.37400	3.29961	0.01244	0.07543
fixing method	X2*X3	0.00000	0.05403	0.01874	0.05719

Table B.2: Simulation study on TII estimators (Section 3.5), simulation results of test function *Ishigami*.

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	1.00000	1.00967	0.09440	0.09489
Jansen	X1*X3	1.00000	0.99815	0.09823	0.09824
Jansen	X2*X3	1.00000	1.00765	0.08971	0.09004
RBD-FAST	X1*X2	1.00000	0.94694	0.07075	0.08844
RBD-FAST	X1*X3	1.00000	0.72957	0.06419	0.27795
RBD-FAST	X2*X3	1.00000	1.25540	0.07165	0.26526
Liu-Owen	X1*X2	1.00000	1.01280	0.11338	0.11410
Liu-Owen	X1*X3	1.00000	1.00429	0.09689	0.09699
Liu-Owen	X2*X3	1.00000	1.01740	0.11797	0.11925
pick-freeze	X1*X2	1.00000	0.99508	0.09425	0.09438
pick-freeze	X1*X3	1.00000	0.99953	0.08953	0.08953
pick-freeze	X2*X3	1.00000	1.00055	0.08620	0.08620
fixing method	X1*X2	1.00000	0.93073	0.45798	0.46319
fixing method	X1*X3	1.00000	0.97158	0.42273	0.42368
fixing method	X2*X3	1.00000	1.03120	0.44287	0.44396

Table B.3: Simulation study on TII estimators (Section 3.5), simulation results of test function *pure 3rd order*.

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	9971.73000	10263.77293	1539.20775	1566.66831
Jansen	X1*X3	4294.61400	4545.41228	1185.94330	1212.17214
Jansen	X1*X4	4959.11800	5073.41172	1231.87397	1237.16472
Jansen	X2*X3	4959.11800	5138.66469	1242.35094	1255.25809
Jansen	X2*X4	5726.43900	5795.54286	1315.01195	1316.82640
Jansen	X3*X4	9971.73000	10071.57890	1378.60814	1382.21931
RBD-FAST	X1*X2	9971.73000	10598.93981	2019.91567	2115.05354
RBD-FAST	X1*X3	4294.61400	3342.48113	3546.85770	3672.43197
RBD-FAST	X1*X4	4959.11800	5810.18917	2760.52452	2888.73982
RBD-FAST	X2*X3	4959.11800	5485.89458	2499.24078	2554.15309
RBD-FAST	X2*X4	5726.43900	4125.33919	3479.07869	3829.81842
RBD-FAST	X3*X4	9971.73000	9650.29305	2215.92131	2239.11343
Liu-Owen	X1*X2	9971.73000	9920.99781	1819.17274	1819.88000
Liu-Owen	X1*X3	4294.61400	4280.11606	615.13594	615.30677
Liu-Owen	X1*X4	4959.11800	4943.35961	638.60851	638.80291
Liu-Owen	X2*X3	4959.11800	4903.13648	612.90229	615.45361
Liu-Owen	X2*X4	5726.43900	5698.02040	621.84855	622.49758
Liu-Owen	X3*X4	9971.73000	9882.87512	1460.01897	1462.72027
pick-freeze	X1*X2	9971.73000	9970.56641	1450.59069	1450.59116
pick-freeze	X1*X3	4294.61400	4200.28742	1183.28793	1187.04163
pick-freeze	X1*X4	4959.11800	4949.93453	968.12614	968.16969
pick-freeze	X2*X3	4959.11800	4929.59915	1057.96011	1058.37184
pick-freeze	X2*X4	5726.43900	5803.13487	976.65760	979.66439
pick-freeze	X3*X4	9971.73000	9966.71909	1508.54626	1508.55458
fixing method	X1*X2	9971.73000	11119.07447	8750.83738	8825.73250
fixing method	X1*X3	4294.61400	4409.87390	2427.35963	2430.09457
fixing method	X1*X4	4959.11800	4996.25926	2424.93445	2425.21887
fixing method	X2*X3	4959.11800	4749.06173	2510.49126	2519.26378
fixing method	X2*X4	5726.43900	5662.63647	2633.24417	2634.01701
fixing method	X3*X4	9971.73000	9755.53576	7674.47413	7677.51869

Table B.4: Simulation study on TH estimators (Section 3.5), simulation results of test function *Branin4*.

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	0.00000	-0.00782	0.06759	0.06804
Jansen	X1*X14	1.00000	0.99866	0.04249	0.04251
Jansen	X3*X4	0.00000	-0.01102	0.05985	0.06086
Jansen	X4*X7	0.00000	0.00446	0.05697	0.05715
Jansen	X5*X11	1.00000	0.99753	0.04346	0.04353
Jansen	X7*X9	1.00000	1.00381	0.04263	0.04280
Jansen	X9*X11	0.00000	-0.00709	0.07341	0.07375
Jansen	X13*X14	0.00000	-0.00481	0.06506	0.06523
Liu-Owen	X1*X2	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X1*X14	1.00000	0.99587	0.01759	0.01807
Liu-Owen	X3*X4	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X4*X7	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X5*X11	1.00000	0.99780	0.01618	0.01632
Liu-Owen	X7*X9	1.00000	0.99960	0.01740	0.01740
Liu-Owen	X9*X11	0.00000	0.00000	0.00000	0.00000
Liu-Owen	X13*X14	0.00000	0.00000	0.00000	0.00000
pick-freeze	X1*X2	0.00000	0.01848	0.08117	0.08325
pick-freeze	X1*X14	1.00000	1.01468	0.08102	0.08234
pick-freeze	X3*X4	0.00000	0.01082	0.10426	0.10482
pick-freeze	X4*X7	0.00000	0.00862	0.09628	0.09666
pick-freeze	X5*X11	1.00000	1.01237	0.08412	0.08502
pick-freeze	X7*X9	1.00000	1.00725	0.08367	0.08399
pick-freeze	X9*X11	0.00000	0.00675	0.09844	0.09867
pick-freeze	X13*X14	0.00000	0.02495	0.10254	0.10553
fixing method	X1*X2	0.00000	0.01132	0.00590	0.01276
fixing method	X1*X14	1.00000	1.00831	0.00594	0.01022
fixing method	X3*X4	0.00000	0.01053	0.00508	0.01169
fixing method	X4*X7	0.00000	0.01056	0.00448	0.01147
fixing method	X5*X11	1.00000	1.00837	0.00558	0.01006
fixing method	X7*X9	1.00000	1.00810	0.00510	0.00958
fixing method	X9*X11	0.00000	0.01058	0.00534	0.01186
fixing method	X13*X14	0.00000	0.01035	0.00519	0.01158

Table B.5: Simulation study on TII estimators (Section 3.5), simulation results of test function *high 2nd order*.

method	interaction	true $\mathfrak{D}_{i,j}$	mean($\widehat{\mathfrak{D}}_{i,j}$)	sd($\widehat{\mathfrak{D}}_{i,j}$)	RMSE
Jansen	X1*X2	1.00000	1.00588	0.27857	0.27863
Jansen	X1*X14	1.00000	0.99121	0.26777	0.26792
Jansen	X3*X4	1.00000	0.95553	0.20567	0.21043
Jansen	X4*X7	1.00000	0.99382	0.26330	0.26337
Jansen	X5*X11	1.00000	0.98963	0.25441	0.25462
Jansen	X7*X9	1.00000	0.99583	0.25962	0.25965
Jansen	X9*X11	1.00000	1.01066	0.29777	0.29796
Jansen	X13*X14	1.00000	0.94978	0.21423	0.22004
Liu-Owen	X1*X2	1.00000	0.94861	0.49732	0.49997
Liu-Owen	X1*X14	1.00000	0.87443	0.40832	0.42720
Liu-Owen	X3*X4	1.00000	0.89585	0.34816	0.36340
Liu-Owen	X4*X7	1.00000	0.94836	0.38186	0.38533
Liu-Owen	X5*X11	1.00000	1.02137	0.50280	0.50325
Liu-Owen	X7*X9	1.00000	0.98915	0.50674	0.50686
Liu-Owen	X9*X11	1.00000	0.98524	0.60291	0.60309
Liu-Owen	X13*X14	1.00000	0.90036	0.37902	0.39190
pick-freeze	X1*X2	1.00000	0.97192	0.20759	0.20948
pick-freeze	X1*X14	1.00000	0.98727	0.20220	0.20260
pick-freeze	X3*X4	1.00000	0.98420	0.19992	0.20055
pick-freeze	X4*X7	1.00000	0.96162	0.18369	0.18766
pick-freeze	X5*X11	1.00000	1.01061	0.22877	0.22902
pick-freeze	X7*X9	1.00000	0.98123	0.23393	0.23468
pick-freeze	X9*X11	1.00000	1.00763	0.23578	0.23590
pick-freeze	X13*X14	1.00000	1.00515	0.20535	0.20541
fixing method	X1*X2	1.00000	0.30138	0.97913	1.20282
fixing method	X1*X14	1.00000	0.34203	1.38163	1.53031
fixing method	X3*X4	1.00000	0.40942	1.30499	1.43240
fixing method	X4*X7	1.00000	0.58874	3.65648	3.67954
fixing method	X5*X11	1.00000	0.54546	2.43655	2.47859
fixing method	X7*X9	1.00000	0.63975	3.04218	3.06343
fixing method	X9*X11	1.00000	0.43615	1.46912	1.57361
fixing method	X13*X14	1.00000	0.53907	1.92319	1.97765

Table B.6: Simulation study on TII estimators (Section 3.5), simulation results of test function *high 14th order*.

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	y
1	175.5102	0.9898	50.0000	0.0000	0.0914	0.0314	0.2102	134.6939	0.2821
2	100.0000	1.5286	147.9592	0.1171	0.0371	0.0429	0.2225	104.0816	0.5137
3	151.0204	0.8429	151.0204	0.0257	0.0857	0.0371	0.1000	102.0408	0.2536
4	153.0612	1.1367	77.5510	0.0029	0.0429	0.0714	0.1082	121.4286	0.6289
5	118.3674	1.5776	92.8571	0.0514	0.0686	0.0914	0.2878	143.8776	0.8150
6	157.1429	1.6510	187.7551	0.0743	0.0743	0.1200	0.2796	138.7755	0.8278
7	112.2449	0.7939	126.5306	0.0371	0.0571	0.0286	0.2918	110.2041	0.4281
8	148.9796	1.0878	175.5102	0.1000	0.0286	0.1343	0.1694	112.2449	0.7314
9	200.0000	1.7000	181.6327	0.0771	0.1000	0.1086	0.1449	129.5918	0.8185
10	120.4082	1.4061	184.6939	0.0714	0.1171	0.1000	0.1367	148.9796	0.8804
11	185.7143	1.3082	74.4898	0.0314	0.0971	0.1371	0.1531	119.3878	0.8598
12	155.1020	1.5041	59.1837	0.0914	0.0829	0.0771	0.2020	106.1225	0.4078
13	138.7755	1.3326	193.8776	0.0571	0.0800	0.0143	0.2388	131.6327	0.6437
14	130.6122	0.9653	53.0612	0.1057	0.0629	0.1286	0.1490	137.7551	0.7054
15	161.2245	1.6020	65.3061	0.1400	0.1229	0.1400	0.1122	118.3674	0.8476
16	122.4490	1.2347	154.0816	0.0171	0.1200	0.1057	0.2592	132.6531	0.8787
17	106.1225	1.4306	120.4082	0.1314	0.1114	0.0600	0.1163	130.6122	0.8080
18	177.5510	1.0143	172.4490	0.0057	0.1371	0.0343	0.2265	123.4694	0.8563
19	197.9592	0.9163	157.1429	0.0400	0.0171	0.1171	0.1326	101.0204	0.4143
20	183.6735	1.1857	141.8367	0.0600	0.0029	0.0629	0.2061	125.5102	0.5995
21	171.4286	1.5531	132.6531	0.0800	0.0400	0.0400	0.1204	150.0000	0.6777
22	140.8163	0.6714	80.6122	0.0829	0.0000	0.0743	0.1735	113.2653	0.4457
23	167.3469	1.0388	56.1225	0.1143	0.0229	0.0229	0.2714	122.4490	0.4120
24	114.2857	1.6755	105.1020	0.0657	0.0143	0.0971	0.1571	126.5306	0.5424
25	189.7959	1.6265	114.2857	0.0429	0.1029	0.0057	0.1776	120.4082	0.6269
26	134.6939	0.8918	102.0408	0.0629	0.0543	0.0086	0.1408	136.7347	0.7283
27	146.9388	1.2837	135.7143	0.1086	0.1314	0.0800	0.2837	116.3265	0.8093
28	128.5714	0.5735	200.0000	0.0486	0.1343	0.0943	0.2184	114.2857	0.8677
29	136.7347	1.0633	111.2245	0.0943	0.0314	0.1257	0.3000	108.1633	0.6944
30	191.8367	0.5000	178.5714	0.1114	0.0457	0.0457	0.2755	139.7959	0.4572
31	195.9184	0.5490	68.3674	0.1257	0.0600	0.0200	0.1653	124.4898	0.3742
32	159.1837	1.2102	117.3469	0.1343	0.0086	0.0029	0.1245	115.3061	0.5054
33	104.0816	1.3571	144.8980	0.0686	0.1400	0.0829	0.1816	107.1429	0.4304
34	173.4694	0.6224	83.6735	0.0143	0.0114	0.0657	0.2633	140.8163	0.4466
35	181.6327	1.1122	166.3265	0.0114	0.0714	0.1143	0.1857	142.8571	0.8553
36	193.8776	0.8184	108.1633	0.0543	0.1257	0.1114	0.2306	100.0000	0.2416
37	126.5306	0.7449	98.9796	0.0886	0.1143	0.0686	0.2959	146.9388	0.8804
38	142.8571	0.9408	71.4286	0.1286	0.1286	0.0171	0.1980	141.8367	0.8622
39	110.2041	0.5980	89.7959	0.0857	0.1057	0.0886	0.2551	111.2245	0.6971
40	116.3265	0.6959	95.9184	0.0229	0.0943	0.0857	0.1041	147.9592	0.8663
41	102.0408	1.1612	138.7755	0.0086	0.1086	0.0257	0.1612	127.5510	0.8061
42	179.5918	0.7204	169.3878	0.0457	0.0514	0.0114	0.2429	105.1020	0.3377
43	163.2653	1.2592	190.8163	0.0343	0.0057	0.0000	0.1286	117.3469	0.4839
44	132.6531	0.7694	196.9388	0.1229	0.0657	0.0514	0.2143	103.0612	0.4160
45	169.3878	1.4551	62.2449	0.1371	0.0200	0.1029	0.1898	133.6735	0.4775
46	187.7551	1.4796	123.4694	0.1200	0.0771	0.0543	0.2469	144.8980	0.7633
47	144.8980	1.3816	163.2653	0.0200	0.0343	0.0571	0.2673	109.1837	0.5016
48	108.1633	0.8673	160.2041	0.0971	0.0257	0.0486	0.2347	145.9184	0.7730
49	124.4898	0.6469	129.5918	0.0286	0.0486	0.1229	0.1939	128.5714	0.4776
50	165.3061	0.5245	86.7347	0.1029	0.0886	0.1314	0.2510	135.7143	0.8165

Table B.7: Thickness reduction application (Section 6.1), Latin hypercube design.

step	run	time intervals: friction								time intervals: blankholder force								y
		1	2	3	4	5	6	7	8	1	2	3	4	5	6	7	8	
1	1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	6.2
	2	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	0.7
	3	+	+	+	+	+	+	+	+	-	-	-	-	-	-	-	-	2.1
	4	-	-	-	-	-	-	-	-	+	+	+	+	+	+	+	+	4.7
	5	+	+	+	+	-	-	-	-	-	-	-	-	-	-	-	-	9.1
	6	-	-	-	-	+	+	+	+	+	+	+	+	+	+	+	+	0.7
	7	+	+	+	+	+	+	+	+	+	+	+	+	-	-	-	-	2.3
	8	-	-	-	-	-	-	-	-	-	-	-	-	+	+	+	+	4.3
2	9	+	+	-	-	-	-	-	-	-	-	-	-	-	-	-	-	6.8
	10	-	-	+	+	+	+	+	+	+	+	+	+	+	+	+	+	0.7
	11	+	+	+	+	+	+	-	-	-	-	-	-	-	-	-	-	6.9
	12	-	-	-	-	-	-	+	+	+	+	+	+	+	+	+	+	0.5
	13	+	+	+	+	+	+	+	+	+	+	-	-	-	-	-	-	2.2
	14	-	-	-	-	-	-	-	-	-	-	+	+	+	+	+	+	4.7
	15	+	+	+	+	+	+	+	+	+	+	+	+	+	+	-	-	1.9
	16	-	-	-	-	-	-	-	-	-	-	-	-	-	-	+	+	4.5
3	17	+	+	+	-	-	-	-	-	-	-	-	-	-	-	-	-	6.8
	18	-	-	-	+	+	+	+	+	+	+	+	+	+	+	+	+	0.7
	19	+	+	+	+	+	+	+	-	-	-	-	-	-	-	-	-	4.2
	20	-	-	-	-	-	-	-	+	+	+	+	+	+	+	+	+	0.5
	21	+	+	+	+	+	+	+	+	+	+	+	-	-	-	-	-	2.2
	22	-	-	-	-	-	-	-	-	-	-	-	+	+	+	+	+	4.4
	23	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	-	0.7
	24	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	+	4.2

Table B.8: Functional input application (Section 6.2), sequential bifurcation design in three steps. Settings: 0.05 (-) and 0.2 (+) for friction and 67.5 (-) and 125 (+) for blankholder force. Mirror runs are shaded.

step 1	$\hat{H}$	step 2	$\hat{H}$	step 3	$\hat{H}$	confounding
$Z_F^{(1,1)} \times Z_F^{(2,1)}$	-0.018	$Z_F^{(1,2)} \times Z_F^{(2,2)}$	-0.007	$Z_F^{(1,3)} \times Z_F^{(2,3)}$	-0.031	1
$Z_F^{(1,1)} \times Z_B^{(1,1)}$	0.006	$Z_F^{(1,2)} \times Z_F^{(3,2)}$	0.008	$Z_F^{(1,3)} \times Z_F^{(3,3)}$	-0.008	–
$Z_F^{(1,1)} \times Z_B^{(2,1)}$	-0.005	$Z_F^{(1,2)} \times Z_F^{(4,2)}$	0.003	$Z_F^{(1,3)} \times Z_F^{(4,3)}$	-0.018	–
$Z_F^{(2,1)} \times Z_B^{(1,1)}$	-0.009	$Z_F^{(1,2)} \times Z_B^{(1,2)}$	-0.023	$Z_F^{(1,3)} \times Z_B^{(1,3)}$	-0.023	2
$Z_F^{(2,1)} \times Z_B^{(2,1)}$	0.008	$Z_F^{(1,2)} \times Z_B^{(2,2)}$	-0.022	$Z_F^{(1,3)} \times Z_B^{(2,3)}$	-0.021	3
$Z_B^{(1,1)} \times Z_B^{(2,1)}$	-0.005	$Z_F^{(1,2)} \times Z_B^{(3,2)}$	0.023	$Z_F^{(1,3)} \times Z_B^{(3,3)}$	0.014	4
		$Z_F^{(1,2)} \times Z_B^{(4,2)}$	0.005	$Z_F^{(1,3)} \times Z_B^{(4,3)}$	-0.004	5
		$Z_F^{(2,2)} \times Z_F^{(3,2)}$	-0.034	$Z_F^{(2,3)} \times Z_F^{(3,3)}$	-0.028	–
		$Z_F^{(2,2)} \times Z_F^{(4,2)}$	-0.038	$Z_F^{(2,3)} \times Z_F^{(4,3)}$	-0.043	–
		$Z_F^{(2,2)} \times Z_B^{(1,2)}$	0.005	$Z_F^{(2,3)} \times Z_B^{(1,3)}$	-0.004	5
		$Z_F^{(2,2)} \times Z_B^{(2,2)}$	0.023	$Z_F^{(2,3)} \times Z_B^{(2,3)}$	0.014	4
		$Z_F^{(2,2)} \times Z_B^{(3,2)}$	-0.022	$Z_F^{(2,3)} \times Z_B^{(3,3)}$	-0.021	3
		$Z_F^{(2,2)} \times Z_B^{(4,2)}$	-0.023	$Z_F^{(2,3)} \times Z_B^{(4,3)}$	-0.023	2
		$Z_F^{(3,2)} \times Z_F^{(4,2)}$	0.026	$Z_F^{(3,3)} \times Z_F^{(4,3)}$	0.241	–
		$Z_F^{(3,2)} \times Z_B^{(1,2)}$	-0.022	$Z_F^{(3,3)} \times Z_B^{(1,3)}$	-0.021	–
		$Z_F^{(3,2)} \times Z_B^{(2,2)}$	-0.005	$Z_F^{(3,3)} \times Z_B^{(2,3)}$	0.005	–
		$Z_F^{(3,2)} \times Z_B^{(3,2)}$	0.013	$Z_F^{(3,3)} \times Z_B^{(3,3)}$	-0.023	–
		$Z_F^{(3,2)} \times Z_B^{(4,2)}$	0.012	$Z_F^{(3,3)} \times Z_B^{(4,3)}$	0.009	–
		$Z_F^{(4,2)} \times Z_B^{(1,2)}$	-0.008	$Z_F^{(4,3)} \times Z_B^{(1,3)}$	-0.008	–
		$Z_F^{(4,2)} \times Z_B^{(2,2)}$	-0.025	$Z_F^{(4,3)} \times Z_B^{(2,3)}$	-0.029	–
		$Z_F^{(4,2)} \times Z_B^{(3,2)}$	-0.001	$Z_F^{(4,3)} \times Z_B^{(3,3)}$	0.036	–
		$Z_F^{(4,2)} \times Z_B^{(4,2)}$	0.013	$Z_F^{(4,3)} \times Z_B^{(4,3)}$	0.021	–
		$Z_B^{(1,2)} \times Z_B^{(2,2)}$	0.028	$Z_B^{(1,3)} \times Z_B^{(2,3)}$	0.075	6
		$Z_B^{(1,2)} \times Z_B^{(3,2)}$	-0.006	$Z_B^{(1,3)} \times Z_B^{(3,3)}$	0.006	7
		$Z_B^{(1,2)} \times Z_B^{(4,2)}$	-0.007	$Z_B^{(1,3)} \times Z_B^{(4,3)}$	-0.031	1
		$Z_B^{(2,2)} \times Z_B^{(3,2)}$	-0.007	$Z_B^{(2,3)} \times Z_B^{(3,3)}$	-0.031	1
		$Z_B^{(2,2)} \times Z_B^{(4,2)}$	-0.006	$Z_B^{(2,3)} \times Z_B^{(4,3)}$	0.006	7
		$Z_B^{(3,2)} \times Z_B^{(4,2)}$	0.028	$Z_B^{(3,3)} \times Z_B^{(4,3)}$	0.075	6

Table B.9: Functional input application, normalized interaction indices of friction (F) and blankholder force (B). The superscripts stand for the following intervals: step 1 (1, 1), (2, 1): [0, 7.5], [7.5, 15], step 2 (1, 2), ..., (4, 2): [0, 3.75], [3.75, 7.5], [7.5, 11.25], [11.25, 15], and step 3 (1, 3), ..., (4, 3): [3.75, 5.625], [5.625, 7.5], [11.25, 13.125], [13.125, 15]. The confounding structure of the fractional factorial designs in steps two and three is given in the last column: confounded interactions are indicated by the same number, no confounding is indicated by –. The strongest interactions are in bold characters.

BIBLIOGRAPHY

- Archer, G. E. B., A. Saltelli, and I. M. Sobol' (1997). Sensitivity measures, ANOVA-like techniques and the use of bootstrap. *Journal of Statistical Computation and Simulation* 58(2), 99–120.
- AutoForm (2004). *AutoForm users manual*. AutoForm Inc.
- Bettonvil, B. (1995). Factor screening by sequential bifurcation. *Communications in Statistics-Simulation and Computation* 24(1), 165–185.
- Blatman, G. and B. Sudret (2010). Efficient computation of global sensitivity indices using sparse polynomial chaos expansions. *Reliability Engineering & System Safety* 95(11), 1216–1229.
- Caflisch, R. E. (1998). Monte Carlo and quasi-Monte Carlo methods. *Acta numerica* 7, 1–49.
- Confalonieri, R., G. Bellocchi, S. Bregaglio, M. Donatelli, and M. Acutis (2010). Comparison of sensitivity analysis techniques: A case study with the rice model warm. *Ecological Modelling* 221, 1897–1906.
- Cornfield, J. and J. W. Tukey (1956). Average values of mean squares in factorials. *The Annals of Mathematical Statistics*, 907–949.
- Csárdi, G. and T. Nepusz (2006). The igraph software package for complex network research. *InterJournal, Complex Systems* 1695.

- Cukier, R. I., H. B. Levine, and K. E. Shuler (1978). Nonlinear sensitivity analysis of multiparameter model systems. *Journal of Computational Physics* 26(1), 1–42.
- Daniel, C. (1973). One-at-a-time plans. *Journal of the American Statistical Association* 68(342), 353–360.
- de Boor, C. (2001). *A practical guide to splines*, Volume 27 of *Applied mathematical sciences*. New York: Springer.
- Durrande, N., D. Ginsbourger, O. Roustant, and L. Carraro (2013). ANOVA kernels and RKHS of zero mean functions for model-based sensitivity analysis. *Journal of Multivariate Analysis* 115, 57–67.
- Efron, B. and C. Stein (1981). The jackknife estimate of variance. *The Annals of Statistics* 9(3), 586–596.
- Fang, K., R. Li, and A. Sudjianto (2006). *Design and modeling for computer experiments*. Computer science and data analysis series. Boca Raton and London: Chapman & Hall/CRC.
- Fort, J.-C., T. Klein, A. Lagnoux, and B. Laurent (2013). Estimation of the Sobol indices in a linear functional multidimensional model. *Journal of Statistical Planning and Inference* 143(9), 1590–1605.
- Fort, J.-C., T. Klein, and N. Rachdi (2014). New sensitivity analysis subordinated to a contrast. *Preprint*. [arXiv:1305.2329](https://arxiv.org/abs/1305.2329) [stat.ME].
- Franco, J., D. Dupuy, O. Roustant, G. Damblin, and B. Iooss. (2014). *DiceDesign: Designs of Computer Experiments*. R package version 1.4.
- Friedman, J. H. and B. E. Popescu (2008). Predictive learning via rule ensembles. *The Annals of Applied Statistics* 2(3), 916–954.

- Fruth, J. and M. Jastrow (2014). *seqSAFI: sequential sensitivity analysis for functional inputs*. R package version 1.0, available on <http://www.statistik.tu-dortmund.de/1552.html>.
- Fruth, J., O. Roustant, and S. Kuhnt (2014a). Sequential designs for sensitivity analysis of functional inputs in computer experiments. *Reliability Engineering & System Safety*, doi: 10.1016/j.ress.2014.07.018.
- Fruth, J., O. Roustant, and S. Kuhnt (2014b). Total interaction index: A variance-based sensitivity index for second-order interaction screening. *Journal of Statistical Planning and Inference* 147, 212–223.
- Fruth, J., O. Roustant, and T. Mühlenstädt (2013). The fanovagraph package: Visualization of interaction structures and construction of block-additive Kriging models. <http://hal.archives-ouvertes.fr/hal-00795229>.
- Gamboa, F., A. Janon, T. Klein, and A. Lagnoux (2013). Sensitivity indices for multivariate outputs. *Comptes Rendus Mathématique* 351(7–8), 307–310.
- Garcia-Cabrejo, O. and A. Valocchi (2014). Global sensitivity analysis for multivariate output using polynomial chaos expansion. *Reliability Engineering & System Safety* 126, 25–36.
- Hoeffding, W. (1948). A class of statistics with asymptotically normal distribution. *The Annals of Mathematical Statistics* 19(3), 293–325.
- Homma, T. and A. Saltelli (1996). Importance measures in global sensitivity analysis of nonlinear models. *Reliability Engineering & System Safety* 52(1), 1–17.
- Hooker, G. (2004). Discovering additive structure in black box functions. In *Proceedings of KDD 2004*, 575–580.
- Iooss, B. and M. Ribatet (2009). Global sensitivity analysis of computer models with functional inputs. *Reliability Engineering & System Safety* 94(7), 1194–1204.

- Ishigami, T. and T. Homma (1990). An importance quantification technique in uncertainty analysis for computer models. In B. M. Ayyub (Ed.), *Proceedings of the ISUMA '90*, 398–403.
- Ivanov, M. and S. Kuhnt (2014). A parallel optimization algorithm based on FANOVA decomposition. *Quality and Reliability Engineering International*, doi: 10.1002/qre.1710.
- James, G. M., J. Wang, and J. Zhu (2009). Functional linear regression that’s interpretable. *The Annals of Statistics* 37(5A), 2083–2108.
- Janon, A., T. Klein, A. Lagnoux, M. Nodet, and C. Prieur (2013). Asymptotic normality and efficiency of two sobol index estimators. *ESAIM: Probability and Statistics*, doi: 10.1051/ps/2013040.
- Jansen, M. J. W. (1999). Analysis of variance designs for model output. *Computer Physics Communications* 117(1–2), 35–43.
- Katz, R. W. (2002). Techniques for estimating uncertainty in climate change scenarios and impact studies. *Climate Research* 20(2), 167–185.
- Kleijnen, J. P. C. (1997). Sensitivity analysis and related analyses: a review of some statistical techniques. *Journal of Statistical Computation and Simulation* 57(1), 111–142.
- Krige, D. G. (1951). A statistical approach to some basic mine valuation problems on the witwatersrand. *Journal of the Chemical, Metallurgical and Mining Society of South Africa* 52(6), 119–139.
- Kucherenko, S., B. Delpuech, B. Iooss, and S. Tarantola (2014). Application of the control variate technique to estimation of total sensitivity indices. *Reliability Engineering & System Safety*, doi: 10.1016/j.res.2014.07.008.

- Kucherenko, S., M. Rodriguez-Fernandez, C. Pantelides, and N. Shah (2009). Monte Carlo evaluation of derivative-based global sensitivity measures. *Reliability Engineering & System Safety* 94(7), 1135–1148.
- Lamboni, M., B. Iooss, A.-L. Popelin, and F. Gamboa (2013). Derivative-based global sensitivity measures: general links with Sobol’ indices and numerical tests. *Mathematics and Computers in Simulation* 87, 45–54.
- Lamboni, M., H. Monod, and D. Makowski (2011). Multivariate sensitivity analysis to measure global contribution of input factors in dynamic models. *Reliability Engineering & System Safety* 96(4), 450–459.
- Lehman, J. S., T. J. Santner, and I. N. William (2004). Designing computer experiments to determinate robust control variables. *Statistica Sinica* 14, 571–590.
- Lilburne, L. and S. Tarantola (2009). Sensitivity analysis of spatial models. *International Journal of Geographical Information Science* 23(2), 151–168.
- Liu, R. and A. B. Owen (2006). Estimating mean dimensionality of analysis of variance decompositions. *Journal of the American Statistical Association* 101(474), 712–721.
- Livermore Software Technology Corporation (2005). *LS-DYNA 970*. California, USA.
- Mara, T. A. (2009). Extension of the RBD-FAST method to the computation of global sensitivity indices. *Reliability Engineering & System Safety* 94(8), 1274–1281.
- Marrel, A., B. Iooss, B. Laurent, and O. Roustant (2009). Calculations of Sobol indices for the Gaussian process metamodel. *Reliability Engineering & System Safety* 94(3), 742–751.
- Mauntz, W. (2002). Global sensitivity analysis of general nonlinear systems. Master’s thesis, Imperial College London, London.

- McKay, M. D. (1995). Evaluating prediction uncertainty. Technical report, Nuclear Regulatory Commission, Washington, DC (United States). Div. of Systems Technology.
- Monod, H., C. Naud, and D. Makowski (2006). Uncertainty and sensitivity analysis for crop models. In D. Wallach, D. Makowski, and J. W. Jones (Eds.), *Working with dynamic crop models*, 55–99. Amsterdam and Boston: Elsevier.
- Morris, M. D. (1991). Factorial sampling plans for preliminary computational experiments. *Communications in Statistics - Theory and Methods* 33(2), 161–174.
- Morris, M. D. (2006). An overview of group factor screening. In A. Dean and S. Lewis (Eds.), *Screening methods for experimentation in industry, drug discovery, and genetics*, 191–206. New York: Springer.
- Morris, M. D. (2012). Gaussian surrogates for computer models with time-varying inputs and outputs. *Technometrics* 54(1), 42–50.
- Morris, M. D. and T. J. Mitchell (1995). Exploratory designs for computational experiments. *Journal of Statistical Planning and Inference* 43, 381–402.
- Mühlenstädt, T., J. Fruth, and O. Roustant (2014). Computer experiments with functional inputs and scalar outputs by a norm-based approach. *Preprint*. [arXiv:1410.0403](https://arxiv.org/abs/1410.0403) [stat.ME].
- Mühlenstädt, T., O. Roustant, L. Carraro, and S. Kuhnt (2012). Data-driven Kriging models based on FANOVA-decomposition. *Statistics and Computing* 22(3), 723–738.
- Owen, A. B. (2013a). Better estimation of small Sobol’ sensitivity indices. *ACM Transactions on Modeling and Computer Simulation* 23(2), 1–17.
- Owen, A. B. (2013b). Variance components and generalized Sobol’ indices. *SIAM/ASA Journal on Uncertainty Quantification* 1(1), 19–41.

- Plischke, E. (2010). An effective algorithm for computing global sensitivity indices (EASI). *Reliability Engineering & System Safety* 95(4), 354–360.
- Pujol, G., B. Iooss, A. J. with contributions from Laurent Gilquin, L. L. Gratiot, and P. Lemaitre (2014). *sensitivity: Sensitivity Analysis*. R package version 1.8-2.
- Punzo, V. and B. Ciuffo (2011). Sensitivity analysis of car-following models: methodology and application. In *Proc. 90th Annual TRB Meeting, Washington, USA*.
- R Core Team (2014). *R: A Language and Environment for Statistical Computing*. Vienna, Austria: R Foundation for Statistical Computing.
- Rabitz, H., O. F. Alis, J. Shorter, and K. Shim (1999). Efficient input–output model representations. *Computer Physics Communications* 117(1-2), 11–20.
- Ramsay, J. O. and B. W. Silverman (1997). *Functional data analysis*. Springer series in statistics. New York: Springer.
- Rasmussen, C. E. and C. K. I. Williams (2006). *Gaussian processes for machine learning*. Cambridge and Mass: MIT Press.
- Roustant, O., J. Fruth, B. Iooss, and S. Kuhnt (2014). Crossed-derivative based sensitivity measures for interaction screening. *Mathematics and Computers in Simulation* 105, 105–118.
- Roustant, O., D. Ginsbourger, and Y. Deville (2012). DiceKriging, DiceOptim: Two R packages for the analysis of computer experiments by Kriging-based metamodeling and optimization. *Journal of Statistical Software* 51(1), 1–55.
- Saltelli, A. (2002). Making best use of model evaluations to compute sensitivity indices. *Computer Physics Communications* 145(2), 280–297.
- Saltelli, A. and R. Bolado (1998). An alternative way to compute Fourier amplitude sensitivity test (FAST). *Computational Statistics & Data Analysis* 26(4), 445–460.

Saltelli, A., K. Chan, and E. M. Scott (2000). *Sensitivity analysis*. Wiley series in probability and statistics. Chichester: Wiley.

Saltelli, A., M. Ratto, S. Tarantola, and F. Campolongo (2006). Sensitivity analysis practices: Strategies for model-based inference. *Reliability Engineering & System Safety* 91(10), 1109–1125.

Saltelli, A., S. Tarantola, and K. P. S. Chan (1999). A quantitative model-independent method for global sensitivity analysis of model output. *Technometrics* 41(1), 39–56.

Santner, T. J., B. J. Williams, and W. I. Notz (2003). *The design and analysis of computer experiments*. Springer series in statistics. New York: Springer.

Schaibly, J. H. and K. E. Shuler (1973). Study of the sensitivity of coupled reaction systems to uncertainties in rate coefficients. II Applications. *The Journal of Chemical Physics* 59(8), 3879–3888.

Sobol', I. M. (1993). Sensitivity estimates for non linear mathematical models. *Mathematical Modelling and Computational Experiments* 1, 407–414.

Sobol', I. M. (2001). Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Mathematics and Computers in Simulation* 55(1-3), 271–280.

Sobol', I. M. and A. Gershman (1995). On an alternative global sensitivity estimator. In A. Saltelli and H. v. Maravic (Eds.), *Samo 95: Theory and Applications of Sensitivity Analysis of Model Output in Computer Simulation*, 40–42. European Commission.

Sobol', I. M. and S. Kucherenko (2009). Derivative based global sensitivity measures and their link with global sensitivity indices. *Mathematics and Computers in Simulation* 79(10), 3009–3017.

- Tarantola, S., D. Gatelli, and T. A. Mara (2006). Random balance designs for the estimation of first order global sensitivity indices. *Reliability Engineering & System Safety* 91(6), 717–727.
- Tissot, J.-Y. and C. Prieur (2012). Bias correction for the estimation of sensitivity indices based on random balance designs. *Reliability Engineering & System Safety* 107, 205–213.
- ul Hassan, H., J. Fruth, A. Güner, and A. E. Tekkaya (2013). Finite element simulations for sheet metal forming process with functional input for the minimization of springback. In *IDDRG conference 2013*, Zürich, 393–398.
- van der Vaart, A. W. (1998). *Asymptotic statistics*. Cambridge: Cambridge University Press.
- Walker, J. S. (1999). *A primer on wavelets and their scientific applications*. London: Chapman & Hall/CRC.
- Watson, G. S. (1961). A study of the group screening method. *Technometrics* 3(3), 371–388.