



**Champs aléatoires et problèmes statistiques inverses
associés pour la quantification des incertitudes :
application à la modélisation de la géométrie des voies
ferrées pour l'évaluation de la réponse dynamique des
trains à grande vitesse et l'analyse**

Guillaume Perrin

► **To cite this version:**

Guillaume Perrin. Champs aléatoires et problèmes statistiques inverses associés pour la quantification des incertitudes : application à la modélisation de la géométrie des voies ferrées pour l'évaluation de la réponse dynamique des trains à grande vitesse et l'analyse. Other. Université Paris-Est, 2013. English. NNT : 2013PEST1137 . pastel-01001045

HAL Id: pastel-01001045

<https://pastel.hal.science/pastel-01001045>

Submitted on 4 Jun 2014

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

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



École doctorale Sciences, Ingénierie et Environnement

Doctoral Thesis
Speciality: Mechanics

presented by
Guillaume PERRIN

Random fields and associated statistical inverse problems for
uncertainty quantification - Application to railway track
geometries for high-speed trains dynamical responses and risk
assessment

Didier CLOUTEAU	ECP - MSSMAT	Examiner
Denis DUHAMEL	ENPC - Navier	Supervisor
Christine FUNFSCHILLING	SNCF - I&R	Examiner
Jean GIORLA	CEA - DCSA	Examiner
Olivier LE MAÎTRE	CNRS - LIMSI	Reporter
Anthony NOUY	ECN - GeM	Reporter
Christian SOIZE	UPEM - MSME	Supervisor

Acknowledgments

This doctoral thesis arised within the framework of a contract between the laboratories Navier and Modélisation Simulation Multi-Échelle at Université Paris Est and the research department of SNCF, and was funded by SNCF and the French ministry of ecology, sustainment development and energy.

I feel very grateful to everybody who contributed to this work. In particular, I would like to thank my supervisors at SNCF and at Université Paris Est. With their strong involvement in this work, their high interest for the subject and their wise advice, they created a very constructive and stimulating environment for this thesis to keep bringing forward. I would like to thank them for all their ideas and proposals, for their required level and precision, which have been endless sources of motivation and innovation during this work. I also would like to thank them for the cordial relations we had and for the taste for research they gave to me.

I would like to express my gratitude to all the members of the jury. It was a pleasure to exchange ideas with such distinguished researchers.

Finally, I would like to address my thanks to all my colleagues at SNCF and at the Université Paris Est for their support, but also for all the good time we had.

Paris, September 2013.

Contents

Acknowledgments	2
Introduction and objectives	8
Industrial objectives	8
Scientific objectives	8
State of the art	9
Main scientific and industrial contributions	11
Outline of the thesis	12
General theoretical frame and corresponding notations	12
1 Short review of the methods for modeling random fields	15
1.1 Introduction	15
1.2 Classical methods to generate random fields	15
1.3 The optimality of the Karhunen-Loève expansion	16
1.3.1 Definition of the Karhunen-Loève expansion	16
1.3.2 Optimality of the KL expansion	17
1.3.3 Practical solving of the Fredholm equation	17
1.3.4 Approximated KL expansion	18
1.4 Identification and generation of random vectors	18
1.5 PCE identification of random vectors	21
1.5.1 Theoretical frame	21
1.5.2 Identification of the polynomial chaos expansion coefficients	22
1.5.3 Practical solving of the log-likelihood maximization	23
1.5.4 Identification of the PCE truncation parameters	25
1.6 Conclusions	27
2 Optimal reduced basis for random fields defined by a set of realizations	29
2.1 Introduction	29
2.2 Theoretical frame	29
2.2.1 Quantification of the relevance of a projection basis	29
2.2.2 Optimality of the Karhunen-Loève expansion	31
2.2.3 Identification difficulties when the information is partial	31
2.3 Identification of optimal basis from a finite set of independent realizations	34
2.3.1 Reformulation of the projection error minimization	34
2.3.2 Restriction of the search space	35
2.3.3 A posteriori evaluation of the representativeness error	37
2.4 Applications	39
2.4.1 Generation of independent realizations of the random field	39
2.4.2 Improvement of the projection basis	42

2.4.3	Optimized basis when few realizations are available	42
2.4.4	Relevance of the LOO error	44
2.5	Conclusions	45
3	PCE identification in high dimension from a set of realizations	46
3.1	Introduction	46
3.2	Optimized trials of independent realizations of random matrices	47
3.2.1	Reformulation of the correlation constraints	47
3.2.2	Notations and definitions	48
3.2.3	Theoretical frame	50
3.2.4	Iterative algorithms	51
3.2.5	Adaptations to the case $\nu < M$	53
3.3	Numerical stabilization of the polynomial basis in high dimension	55
3.3.1	Decomposition of the matrix of independent realizations	56
3.3.2	Influence of the truncation parameters	58
3.3.3	Adaptation of the optimization problem	60
3.3.4	Remarks on the new optimization problem	62
3.4	Application	63
3.4.1	Application in low dimension	63
3.4.2	Application in high dimension	74
3.4.3	Application in very high dimension with limited information	77
3.5	Conclusions	80
4	Karhunen-Loève expansion revisited for vector-valued random fields	82
4.1	Introduction	82
4.2	Scaled expansion	83
4.2.1	Local-global errors and optimal basis	83
4.2.2	Scaled expansion	85
4.2.3	Properties of the scaled expansion	86
4.2.4	Minimization of a weighted sum of local errors	89
4.2.5	Minimization of the maximal value of the local errors	90
4.3	Application	92
4.3.1	Generation of a vector-valued random field	92
4.3.2	Influence of the scaling vector on the local errors	93
4.3.3	Identification of the optimal basis	93
4.4	Conclusions	97
5	Experimental identification of the railway track stochastic modeling	100
5.1	Introduction	100
5.2	Experimental measurements and signal processing	102
5.2.1	Collection of the experimental inputs for the modeling	102
5.2.2	Local-global approach	102
5.3	Optimal reduced basis	106
5.3.1	Direct KL expansion and projection biases	106
5.3.2	Optimization of the projection basis	107
5.3.3	Choice of the dimension of the spatial projection parameter	109
5.4	PCE identification in very high dimension	109
5.4.1	Sorting with respect to the horizontal curvature	109
5.4.2	PCE identification	111
5.5	Generation of a whole track geometry	113

5.6	Statistical and frequency validations	116
5.7	Conclusions	121
6	Stochastic dynamics of high-speed trains and risk assessment	122
6.1	Introduction	122
6.2	Description of the railway dynamic problem	122
6.2.1	Deterministic railway problem	122
6.2.2	Domain of validity for the deterministic problem	125
6.3	Definition of the stochastic problem and validation of the modeling	128
6.3.1	Stochastic problem	128
6.3.2	Validation of the stochastic problem	128
6.4	Propagation of the variability	132
6.4.1	Influence of the track design	132
6.4.2	Influence of an increase of the speed on the quantities of interest	134
6.4.3	Comparison of three high speed trains	136
6.5	Sensitivity analysis	136
6.5.1	Nonlinearities and importance of the conjunction of track irregularities	137
6.5.2	KL-based sensitivity analysis	138
6.6	Conclusions	141
	Conclusions and prospects	144
	Summary of the industrial context	144
	Scientific and industrial contributions	145
	Prospects	146
	Publications and external communications	150
	Appendix	153
A	Proof of Lemma 2	153
B	Generation of the matrix-valued autocorrelation matrix	155
C	Definition of the local-global error functions	156

Introduction and objectives

Industrial objectives

High speed trains are currently meant to run faster and to carry heavier loads, while being less energy consuming and still ensuring the safety and comfort certification criteria. In order to optimize the conception of such high technology trains, we need a precise knowledge of the realm of possibilities of track conditions that the train is likely to be confronted to during its life cycle.

In parallel, since 2012, European high speed railway networks are meant to have gone to market. Several high speed trains, such as ICE, TGV, ETR 500..., for which mechanical properties and structures are different, are likely to run on the same tracks, whereas they may have been originally designed for specific and different railway networks. European high speed railway networks are therefore bound to be subjected to an increasing variability of mechanical loads. To optimize the track maintenance and to adjust the tolls according to the aggressiveness of a particular train toward the track, a better understanding of the interaction between the train dynamic behavior and the track geometry is necessary.

Simulation is a very useful tool to face these challenges. However, it has to be very representative of the physical behavior of the system. The models of the train, of the railway track, and of the wheel/rail contacts have thus to be fully validated and the simulations have to be raised on realistic and representative sets of excitations.

Hence, based on experimental measurements, a complete parametrization of the track geometry and of its variability would be of great concern to analyze the complex link between the train dynamics and the physical and statistical properties of the track geometry.

Scientific objectives

From a scientific point of view, a railway simulation can be seen as the dynamic response of a complex mechanical system excited by a multivariate random field, for which statistical properties are only known through a set of independent realizations. Due to the specific interactions between the train and the track, this random field is neither stationary nor Gaussian.

In order to propagate the track geometry variability to the train response, methods to identify in inverse, from a finite set of experimental data, the statistical properties of non-stationary and non-Gaussian random fields will be analyzed in this manuscript.

The train behavior being very nonlinear and very sensitive to the track geometry, the random field has to be described very precisely from frequency and statistical points of view. As a result, the statistical dimension of this random field is very high. Hence, a particular attention will be paid in this thesis to statistical reduction methods and to statistical identification methods that can be numerically applied to the high dimensional case.

State of the art

The general scheme for probabilistic analysis is usually divided in three steps (see [1, 2, 3] for further details). First, the mechanical model and the associated input parameters and output criteria (safety criteria for instance) have to be defined precisely. Then, the different sources of uncertainty have to be identified and modeled carefully. At last, the input uncertainty has to be propagated through the deterministic model, in order to characterize the statistical properties of the output quantities of interest.

These three steps are rapidly described hereunder for the studied railway system.

Mechanical model. In this work, the reactions of trains excited by the track geometry through the specific wheel/rail contacts are studied. Three kinds of inputs are therefore needed in such simulations:

- the vehicle model. Multibody simulations are usually employed to model the train dynamics (see [4]). Carbodies, bogies and wheelsets are modeled by rigid bodies linked with connections represented by rheologic models (damper, springs, ...). This leads us to several hundreds of degrees of freedom.
- the track model. A double scale parametrization is usually introduced to describe the track geometry (see [5]): each rail position is characterized by a mean-line position, which only depends on the vertical and horizontal curvatures, on the track super-elevation and the track gauge of the track, and by a deviation towards this mean-line position, which can be described by four curvilinear irregularity fields. While the mean position is decided once for all at the building of a new line, the track irregularities can evolve with respect to the track substructure, to the weather conditions and to the train dynamics. Seven curvilinear fields are needed to completely characterize the positions of the two rigid rails.
- the wheel/rail contact model. The wheel/rail contact forces are computed for any position of the train from the wheel and the rail profiles thanks to the Hertz and Kalker theories ([6, 7]). The contact properties are moreover generally recorded in a contact table.

Given these three inputs, the train response can be computed as the solution of a system of coupled equations that are strongly **nonlinear**. This system is usually solved with an explicit scheme. Once these equations have been solved, the spatial accelerations of each mass body, as well as the internal and external loads are available. These railway outputs can then be post-processed to define safety, comfort and maintenance criteria.

In this work, the commercial code Vampire is used to solve these equations. The movement equations of the railway dynamics are thus not available. Moreover, the duration of a whole railway simulation over a length of 5km is approximately 120 seconds on a standard computer.

Uncertainty quantification. Several sources of uncertainty can be categorized:

- Model uncertainty. In each model, simplifying hypotheses are introduced. In the studied system, the rigid body modeling of the train and the Hertz formulation for the wheel/rail contact are two examples of such model simplifications.

- Parameter uncertainty. The chosen model to describe the considered system is generally based on parameters, for which exact values are unknown and cannot exactly be experimentally measured. For instance, the total mass of a train is, in practice, impossible to precisely evaluate.
- Parameter variability, which comes from the physical variability of the input parameters of the model. Example is the train suspensions, for which the process of manufacturing leads us to mechanical characteristics that are not exactly as designed such that the performance can vary from one suspension to another one.
- Algorithmic uncertainty, which comes from numerical approximations. In the railway field, this uncertainty comes mostly from the time discretization in the explicit solver of the movement equations. A convergence analysis has thus to be performed to choose a relevant time step.
- Measurements uncertainty. Experimental data are images of the reality, for which biases have to be minimized as much as possible.

Due to these uncertainties, discrepancies will always be observed when comparing the modeled and measured deterministic responses of a train. On the contrary, stochastic models, which would be able to take into account these uncertainties, should lead to a better representation of the behavior of the system. This explains the very high interest for these methods that have spread for the last decades to most of the scientific fields.

Some specific fields of the probability theory have therefore focused on particular sources of uncertainty. First, the methods based on Information Theory and on the Maximal Entropy principle (see [8] and [9]) have been continuously improved to always better characterize the parameter uncertainty and variability from the only available and usable information. In the same manner, the use of methods based on the Bayesian method (see [10, 11, 12]) has kept increasing to update the input stochastic modeling in the light of new and relevant data. Then, when the movement equations of the system are available, non-parametric probability models (see [13, 14]) have been introduced to take into account not only the input uncertainty but also the model and algorithmic uncertainties.

In this work, it is supposed that a nominal model of a train is available, for which mechanical parameters are fixed and have been accurately identified. In the same manner, the contact properties are computed once for all from a new rail profile and a new wheel profile. Given a particular description of the track geometry and these contact and vehicle models, it is assumed that both the railway model and the numerical solving are sufficiently relevant to accurately compute the response of the system: the approximations introduced in the computational scheme are supposed to be controlled and the movement equations are assumed to be precise enough to represent the physical phenomena.

Hence, only the uncertainty in the track geometry is addressed in this thesis. In this prospect, this track geometry will be seen as a multivariate random field. To identify this field, a set of experimental measurements of the track geometry is used. It is assumed that the experimental uncertainties for these measurements are negligible, such that no distinction will be made between the track measurements and the real track geometry in the following. These measurements define the maximal available information about the track-geometry random field.

Uncertainty propagation and risk assessment. Once the parameter uncertainty and variability have been characterized, the variability has to be propagated through the mechanical

model. The choice of the propagation method depends on the chosen variables of interest and on the computational cost of the simulation. In this work, we will focus on the accelerations of the train mass bodies and the loads between the train and the track. We are moreover interested in probabilities for these outputs to exceed normalized thresholds.

Recall that the railway mechanical is based on a very high number of variable input parameters, that the train response is very sensitive, very non linear, and very fuzzy with respect to these input parameters, that the movement equations are not available, and that the duration of one simulation is rather cheap. The best method to compute such probabilities of exceeding thresholds for these railway outputs is therefore the Monte Carlo (MC) method ([15]). Indeed, the statistical convergence of such a method depends neither on the dimension of the input, nor on the complexity and the nonlinearity of the mechanical model, and is particularly adapted to systems that are controlled by black-box codes, that is to say, codes for which movement equations are not available, as it is the case here.

In order to get accurate results, much attention has to be paid to the modeling of the input variability, as any error on the input will be propagated to the output. In addition, the MC method asks for the generation of sets of independent realizations of the input parameters. As this work focus on the track geometry variability, methods to generate independent realistic and representative track conditions will be needed in this work.

At last, based on this MC method, each railway simulation gives access to a particular realization of the time reactions of the track. The risk assessment has therefore to be performed using statistical methods based on stochastic processes (see [16] for further details).

Main scientific and industrial contributions

The developments of this work were achieved to answer the four following questions.

- **Virtual certification.** How to develop a track generator, which would be able to generate track conditions, which are on the one hand realistic from a statistical, frequency and dynamical point of view, and from the other hand representative of a measured set of experimental data? The numerical certification indeed requires a large set of representative track conditions to capture rare events [3].
- **Optimization of the system.** How to propagate the track geometry variability to the train dynamical quantities of interest, which are mostly lateral and vertical accelerations and loads? The knowledge of the link between the track variability and the response of the train could indeed help us to propose optimized maintenance policies.
- **Railway field going to market.** How to develop a method to evaluate and compare the aggressiveness of several trains that would be likely to run on the same network?

Four scientific main scientific contributions are summarized hereunder.

1. The statistical dimension of the track-geometry random field is very high, such that advanced reduction techniques will be needed to optimally condense the statistical properties of the random field to be identified. In particular, the importance of the Karhunen-Loève (KL) expansion will be analyzed in detail in this work.

2. The available information about the track-geometry random field is very reduced compared to its statistical dimension. The statistical moments of this random field, such as the empirical estimators of the mean function or the covariance operator, on which the KL expansion is based, are not converged. A method to adapt the KL formulation to this kind of problems will thus be proposed in this thesis.
3. The track-geometry random field is multivariate, and its different components are very statistically dependent. A vectorial approach has therefore to be considered in order to accurately take into account the dependencies between these different components of the track-geometry random field. Moreover, the amplitudes of these components are different and their importances on the dynamical quantities are *a priori* unknown. An other adaptation of the classical KL expansion has thus to be introduced in order to identify a reduced basis that allows the description of each component of the random field of interest with the same precision.
4. Due to the specific interaction between the train and the track, the track-geometry random field is neither stationary nor Gaussian, such that a particular attention has to be paid to the identification of the multidimensional distribution of the coefficients of the random field on the reduced projection basis. Due to the complexity of the random field to be modeled, these coefficients define a very high dimension random vector. To this end, an adaptation to the very high dimension of the identification in inverse methods based on a polynomial chaos expansion will be presented in this work.

Outline of the thesis

From these objectives, the document is organized in six chapters that are now presented.

Chapter 1 contains a review of well-known methods for random field identification and generation. In particular, the Karhunen-Loève (KL) expansion and the polynomial chaos expansion (PCE) identification in inverse will be presented in detail.

The next chapters are devoted to the author contributions in the field of uncertainty propagation. Chapter 2 deals with the adaptation of the KL method to cases for which the maximal available information about the random field to identify is limited to a finite set of independent realizations.

Chapter 3 addresses the adaptation of the polynomial chaos expansion identification methods to the very high dimensional case.

Chapter 4 presents an original scaled KL expansion for the analysis of vector-valued random fields.

Chapter 5 considers the application of the theoretical developments of Chapters 2, 3 and 4 to identify, in inverse, from experimental data, the statistical properties of the track-geometry random field.

At last, Chapter 6 shows in what extent such a stochastic modeling of the track geometry opens new opportunities for the railway field in certification, maintenance, and safety prospects.

General theoretical frame and corresponding notations

This section aims at summarizing the main notations that will be used in this manuscript.

- $\mathbb{R}$ denotes the set of real numbers.
- $\mathbb{N}$ is the set of positive integers.

- $\Omega \subset \mathbb{R}$ refers to a subset of $\mathbb{R}$.
- $(\Theta, \mathcal{T}, \mathcal{P})$ is a probability space.
- $E[\cdot]$ is the mathematical expectation.
- $\mathcal{H} = L^2_{\mathcal{P}}(\Theta, \mathbb{R}^M)$ is the space of all the second-order random vectors defined on $(\Theta, \mathcal{T}, \mathcal{P})$ with values in $\mathbb{R}^M$, equipped with the inner product $\langle \cdot, \cdot \rangle$:

$$\langle \mathbf{A}, \mathbf{B} \rangle = \int_{\Theta} \mathbf{A}^T(\theta) \mathbf{B}(\theta) dP(\theta) = E[\mathbf{A}^T \mathbf{B}], \quad \forall \mathbf{A}, \mathbf{B} \in L^2_{\mathcal{P}}(\Theta, \mathbb{R}^M). \quad (1)$$

- $\mathcal{P}^{(Q)}([0, S])$, where $S < +\infty$, is the space of all the second-order $\mathbb{R}^Q$ -valued random fields, indexed by the compact interval $[0, S]$.
- For $Q \geq 1$, $\mathbf{X} = (X_1, \dots, X_Q) = \{(X_1(s), \dots, X_Q(s)), s \in [0, S]\}$ is in $\mathcal{P}^{(Q)}([0, S])$.
- Let $\mathbb{H} = L^2([0, S], \mathbb{R}^Q)$ be the space of square integrable functions on $[0, S]$, with values in $\mathbb{R}^Q$, equipped with the inner product $(\cdot, \cdot)$, such that, for all $\mathbf{u}$ and $\mathbf{v}$ in $\mathbb{H}$,

$$(\mathbf{u}, \mathbf{v}) = \int_{[0, S]} \mathbf{u}(s)^T \mathbf{v}(s) ds. \quad (2)$$

- $\|\cdot\|_{\mathcal{P}^{(Q)}([0, S])}$ denotes the L_2 norm in $\mathcal{P}^{(Q)}([0, S])$, such that:

$$\|\mathbf{X}\|_{\mathcal{P}^{(Q)}([0, S])}^2 = E \left[\int_{\Omega} \mathbf{X}(s)^T \mathbf{X}(s) ds \right], \quad \mathbf{X} \in \mathcal{P}^{(Q)}([0, S]). \quad (3)$$

- δ_{mp} is the kronecker symbol that is equal to 1 if $m = p$ and 0 otherwise.
- $\text{Tr}[\cdot]$ is the trace operator for square matrices.
- a, b correspond to constants in $\mathbb{R}$.
- $\mathbf{a}, \mathbf{b}$ refer to vectors with values in $\mathbb{R}^Q$, $Q \geq 1$.
- $\times$ is the vectorial product between vectors.
- $\mathbf{a}^T$ is the transpose of $\mathbf{a}$.
- $\otimes$ is the tensorial product such that $\mathbf{a} \otimes \mathbf{b} = \mathbf{a} \mathbf{b}^T$.
- A, B correspond to random variables with values in $\mathbb{R}$.
- $\mathbf{A}, \mathbf{B}$ denote random vectors with values in $\mathbb{R}^Q$, $Q \geq 1$.
- $[A], [B]$ refer to real matrices.
- $\|\cdot\|_F$ is the Frobenius norm of matrices.
- $P_{\mathbf{A}}$ and $p_{\mathbf{A}}$ denote respectively the multidimensional probability distribution and the multidimensional Probability Density Function (PDF) of random vector $\mathbf{A}$.
- If random vector $\mathbf{A}$ is of second order, we denote by $\boldsymbol{\mu}_{\mathbf{A}}$ and $[R_{\mathbf{A}\mathbf{A}}]$ the mean and the covariance matrix of $\mathbf{A}$ respectively.

- $(s, s') \mapsto [R_{\mathbf{X}\mathbf{X}}(s, s')]$ corresponds to the matrix-valued covariance function of $\mathbf{X}$, such that for all s, s' in Ω , $[R_{\mathbf{X}\mathbf{X}}(s, s')] = E[(\mathbf{X}(s) - E[\mathbf{X}(s)]) \otimes (\mathbf{X}(s') - E[\mathbf{X}(s')])]$.
- When $Q = 1$, $\mathcal{P}^{(1)}(\Omega)$, $\mathbf{X}$ and $[R_{\mathbf{X}\mathbf{X}}]$ are written $\mathcal{P}(\Omega)$, X and R_{XX} respectively for the sake of simplicity.
- $\mathcal{F}^{(M)}$ denotes a subset of $\mathbb{H}$ that gathers M functions with values in $\mathbb{R}^Q$ that are defined on Ω .
- $\widehat{\mathbf{X}}^{\mathcal{F}^{(M)}}$ refers to the projection of $\mathbf{X}$ on the subspace spanned by $\mathcal{F}^{(M)}$.

Chapter 1

Short review of the methods for modeling random fields

1.1 Introduction

As presented in Introduction, the goal of this work is to quantify the influence of the track geometry variability on the train dynamical responses. A good approach to take into account this input variability is to consider the track geometry as a multivariate random field. It has moreover been shown that the most appropriate method to propagate the track variability through the mechanical model is the Monte Carlo (MC) method. For such a method to be implemented, one has therefore to be able to generate independent realizations of this track-geometry random field. Due to the specific interactions between the train and the track, this random field is neither Gaussian nor stationary. In this prospect, several existing methods to identify and generate non-Gaussian random fields are addressed in this chapter. More precisely, this chapter describes in detail the method on which the stochastic modeling of the track geometry will be based in the next chapters, which is based on the coupling of a Karhunen-Loève expansion and a polynomial chaos expansion.

1.2 Classical methods to generate random fields

For the last decades, the random fields analysis has been used in an increasing number of scientific fields, such as uncertainties quantification, material sciences, seismology, geophysics, quantitative finance, signal processing, control engineering etc. It is indeed a very interesting tool for stochastic modeling, forecasting, classification, signal detection and estimation. Let

$$X = \{X(s), s \in \Omega \subset \mathbb{R}\}, \quad (1.1)$$

be a random field for which we want to generate sample paths. For the sake of simplicity, and without any loss of generality, only centered random fields X are considered in this work:

$$E[X(s)] = 0, \quad \forall s \in \Omega, \quad (1.2)$$

where $E[\cdot]$ is the mathematical expectation.

The Gaussian case is a well-posed problem, as the Gaussian random fields are completely characterized only by their mean function and their autocorrelation function. It exists therefore many effective methods to simulate Gaussian random fields. In particular, when $\Omega = \mathbb{R}$, AutoRegressive-Moving-Average (ARMA) models, that were first introduced by Whittle for

time series [17, 18] and popularized by Box and Jenkins [19], allow the description of Gaussian stationary random fields as a parameterized integral of a Gaussian white noise random field. Based on limited knowledge of random field X , these models can therefore be used to emphasize particular properties of X and to extrapolate its value.

On the contrary, the random field simulation problem is an ill-posed problem. To characterize a non-Gaussian random field, we need to know the entire family of joint probability distributions $\{(X(s_1), \dots, X(s_n)), n \geq 1, (s_1, \dots, s_n) \in \Omega^n\}$. As this information is most of the time not accessible, only partial description of non-Gaussian random field can be given.

Two classes of methods are generally used to characterize such non-Gaussian random fields. On the first hand, translation methods allow the identification and the generation of a non-Gaussian random field from a memoryless nonlinear transformation of a known Gaussian random field (see for instance [20]).

On the other hand, in the general case, spectral methods ([21, 22]) based on a two-step approach have given very promising results to identify the distribution of *a priori* non-Gaussian and non-stationary random fields. The first step of these methods is generally the approximation of the random field, X , by its projection $\widehat{X}^{\mathcal{B}^{(M)}}$ on a M -dimension set of deterministic functions, $\mathcal{B}^{(M)} = \{b_m(s), s \in \Omega\}_{1 \leq m \leq M}$, that are supposed to be square integrable on Ω and orthonormal such that:

$$\widehat{X}^{\mathcal{B}^{(M)}} = \sum_{m=1}^M C_m b_m, \quad \int_{\Omega} b_m(s) b_p(s) ds = \delta_{mp}, \quad C_m = \int_{\Omega} X(s) b_m(s) ds, \quad (1.3)$$

where δ_{mp} is the kronecker symbol. The vector $\mathbf{C} = (C_1, \dots, C_M)$ is thus a M -dimension random vector, for which components are *a priori* dependent. The second step is then the identification of the multidimensional distribution of $\mathbf{C}$.

When the knowledge of the random field is limited to a set of independent realizations, as it is the case for the modeling of the track geometry, such spectral methods present many advantages. First, no hypothesis on the random field is required to implement these methods. Then, by proposing a discretized description of the random field, they take advantage of all the developments that have been done in the characterization of the multidimensional distribution of non-Gaussian random vectors.

1.3 The optimality of the Karhunen-Loève expansion to generate approximated realizations of random fields

1.3.1 Definition of the Karhunen-Loève expansion

Mathematically, the Karhunen-Loève (KL) expansion corresponds to the orthogonal projection theorem in separable Hilbert spaces. In this case, the Hilbertian basis, $\{k_m, m \geq 1\}$, is constructed as the eigenfunctions of the covariance operator of X , defined by the covariance function, R_{XX} , which is assumed, for instance, to be square integrable on $\Omega \times \Omega$. Therefore, for all (s, s') in $\Omega \times \Omega$ and $m \geq 1$ and $p \geq 1$, we get:

$$R_{XX}(s, s') \stackrel{\text{def}}{=} E [X(s)X(s')] = \sum_{m \geq 1} \lambda_m k_m(s) k_m(s'), \quad (1.4)$$

$$\int_{\Omega} R_{XX}(s, s') k_m(s') ds' = \lambda_m k_m(s), \quad (1.5)$$

$$(k_m, k_p) = \delta_{mp}, \quad \lambda_1 \geq \lambda_2 \geq \dots \rightarrow 0, \quad \sum_{m \geq 1} \lambda_m^2 < +\infty. \quad (1.6)$$

1.3.2 Optimality of the KL expansion

In order to represent the field X with a small number of vectors M , it is important to choose a relevant basis regarding X . Indeed, the more relevant the projection basis $\mathcal{B}^{(M)}$ is, the lower the dimension M has to be, to guarantee that the amplitude of the residue,

$$\mathcal{N}^2(X - \widehat{X}^{\mathcal{B}^{(M)}}), \quad (1.7)$$

is lower than a given threshold, and so the easier and the more precise the identification of the distribution of $\mathbf{C}$ will be. $\mathcal{N}^2$ is a norm that has to be adapted to the studied problem. If

$$\mathcal{N}^2(\cdot) = E \left[\int_{\Omega} (\cdot)^2 \right], \quad (1.8)$$

due to the orthogonal projection theorem in Hilbert spaces, for any integer M , the M -dimension family $\mathcal{K}^{(M)} = \{k_m, 1 \leq k \leq M\}$, which gathers the M first elements of the KL basis associated with X , minimizes the amplitude $\mathcal{N}^2(X - \widehat{X}^{\mathcal{F}^{(M)}})$ among all the M -dimension families $\mathcal{F}^{(M)}$, where $\widehat{X}^{\mathcal{F}^{(M)}}$ is the projection of X on $\mathcal{F}^{(M)}$. In other words, for any $M \geq 1$, it can be shown that:

$$\mathcal{N}^2(X - \widehat{X}^{\mathcal{K}^{(M)}}) \leq \mathcal{N}^2(X - \widehat{X}^{\mathcal{F}^{(M)}}), \quad (1.9)$$

where $\widehat{X}^{\mathcal{K}^{(M)}}$ is the projection of X on $\mathcal{K}^{(M)}$.

Due to this optimality property, the Karhunen-Loève (KL) basis has played, for the last decades, a major role and has been applied in many works (see for instance [23, 24, 25, 26, 27, 28, 29, 30, 31, 32, 33, 34, 35, 36, 37, 38, 39, 40, 41, 42, 43]).

1.3.3 Practical solving of the Fredholm equation

Equation (1.5) is commonly referred to as Fredholm equation, and issues concerning the solving of this integral eigenvalue problem can be found in [21, 44, 45]. The idea of this section is to describe the different steps to solve the Fredholm problem thanks to a finite element approach when $\Omega = [0, S]$. To this end, the functions k_m , $1 \leq m \leq M$, are searched as their finite element estimator k_m^{FE} , such that, for all s in Ω :

$$k_m(s) \approx k_m^{\text{FE}}(s) = \sum_{j=1}^{N_S} d_j^m h_j(s), \quad (1.10)$$

$$\mathbf{d}^m = (d_1^m, \dots, d_{N_S}^m), \quad \mathbf{h}(s) = (h_1(s), \dots, h_{N_S}(s)), \quad (1.11)$$

where $\mathbf{d}^m$ is the unknown vector to be identified, and $\{s \mapsto h_j(s), 1 \leq j \leq N_S\}$ are shape functions such that:

$$\begin{cases} s_1 = 0, \quad s_{N_S} = S, \quad s_{j+q} - s_j = qh, \\ h_j(s_k) = \delta_{jk}, \quad 1 \leq j, k \leq N_S, \\ \sum_{j=1}^{N_S} h_j(s) = 1, \quad s \in \Omega = [0, S], \end{cases} \quad (1.12)$$

with $h = S / (N_S - 1)$ the finite element discretization length. The finite element discretization of Eq. (1.5) yields:

$$([K] - \lambda_m[M]) \mathbf{d}^m = \mathbf{0}, \quad (1.13)$$

in which the positive-definite symmetric $(N_S \times N_S)$ real matrices $[K]$ and $[M]$ are defined by

$$[K] = \int_{\Omega} \int_{\Omega} \mathbf{h}(s)^T [R_{XX}(s, s')] \mathbf{h}(s') ds' ds, \quad (1.14)$$

$$[M] = \int_{\Omega} \mathbf{h}(s)^T \mathbf{h}(s) ds. \quad (1.15)$$

This approach is particularly well adapted to the modeling of random fields, for which experimental values are recorded every $\tilde{h}$ meters. Spatial discretization step h is thus chosen equal to $\tilde{h}$ to limit the error introduced by the finite element approach. Moreover, it has to be noticed that the regularity of the shape functions has to be adapted to the regularity of random field X . In particular, if the first and second order spatial derivatives of the random field paths are *a priori* non zero, at least cubic shape functions will be needed.

1.3.4 Approximated KL expansion

As presented in Section 1.3.1, the KL expansion of a centered random field X is based on the knowledge of its autocovariance function, R_{XX} . When the maximal available information about X is a set of ν independent realizations, $\{X(\theta_1), \dots, X(\theta_\nu)\}$, this function is not exactly known, but can be approximated by its empirical estimation, $\hat{R}_{XX}(\nu)$, such that:

$$R_{XX}(s, s') \approx \hat{R}_{XX}(\nu, s, s') = \frac{1}{\nu} \sum_{n=1}^{\nu} X(\theta_n, s) X(\theta_n, s'), \quad (s, s') \in \Omega \times \Omega. \quad (1.16)$$

By solving the Fredholm problem associated with $\hat{R}_{XX}(\nu)$ instead of R_{XX} , it is therefore possible to identify a rather good approximation of the KL basis of X , which is denoted by $\{\hat{k}_m(\nu), 1 \leq m\}$, especially when ν is high, as:

$$\begin{aligned} \lim_{\nu \rightarrow +\infty} \hat{R}_{XX}(\nu) &= R_{XX}, \\ \lim_{\nu \rightarrow +\infty} \hat{k}_m(\nu) &= k_m. \end{aligned} \quad (1.17)$$

1.4 Direct and indirect methods for the identification of the distribution of random vectors and their generation

Once random field X has been projected on a chosen deterministic M -dimension family, $\mathcal{B}^{(M)} = \{b_m(s), s \in \Omega\}_{1 \leq m \leq M}$, such that

$$X \approx \hat{X}^{\mathcal{B}^{(M)}} = \sum_{m=1}^M C_m b_m, \quad (1.18)$$

identifying its statistical distribution amounts to identifying the multidimensional distribution of random vector $\mathbf{C} = (C_1, \dots, C_M)$, denoted by $P_{\mathbf{C}}$. The mean value and the covariance matrix of $\mathbf{C}$ are moreover denoted by $\boldsymbol{\mu}_{\mathbf{C}}$ and $[R_{\mathbf{C}\mathbf{C}}]$ respectively, such that:

$$\boldsymbol{\mu}_C = E[\mathbf{C}], \quad [R_{CC}] = E[(\mathbf{C} - \boldsymbol{\mu}_C) \otimes (\mathbf{C} - \boldsymbol{\mu}_C)]. \quad (1.19)$$

In this work, it is assumed that $P_C(d\mathbf{x}) = p_C(\mathbf{x})d\mathbf{x}$, in which the probability density function (PDF) p_C is a function in the set $\mathcal{F}(\mathcal{D}, \mathbb{R}^*)$ of all the positive-valued functions defined on any part $\mathcal{D}$ of $\mathbb{R}^M$ and for which integral over $\mathcal{D}$ is 1.

Two kinds of methods can be used to build such a PDF: the direct and the indirect methods. Among the direct methods, the Prior Algebraic Stochastic Modeling (PASM) methods postulate an algebraic representation $\mathbf{C} \approx t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w})$, with t_{alg} a prior transformation, $\boldsymbol{\Xi}$ a random vector and $\mathbf{w}$ a vector of parameters to be identified. For instance, we can suppose that $\mathbf{C}$ can be written under the form:

$$\mathbf{C} \approx t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w}) = \mathbf{w}_1 + [w_2]\boldsymbol{\Xi}, \quad \mathbf{w} = \{\mathbf{w}_1, [w_2]\}, \quad (1.20)$$

with $\boldsymbol{\Xi}$ a M -dimension random vector for which components are independent, normally distributed with zero mean and unit variance. It can directly be seen that:

$$E[t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w})] = \mathbf{w}_1, \quad E[(t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w}) - \mathbf{w}_1) \otimes (t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w}) - \mathbf{w}_1)] = [w_2][w_2]^T. \quad (1.21)$$

Hence, supposing that $\mathbf{C} \approx t_{\text{alg}}(\boldsymbol{\Xi}, \mathbf{w})$ amounts to supposing that $\mathbf{C}$ is a Gaussian random vector, such that the most accurate values for $\mathbf{w}_1$ and $[w_2]$ correspond to the mean value of $\mathbf{C}$ and to the Cholesky decomposition matrix of matrix $[R_{CC}]$. If $\mathbf{C}$ is actually not Gaussian, this transformation is not relevant, and another one has to be introduced to better represent the behavior of $\mathbf{C}$, such as for instance:

$$\mathbf{C} \approx t_{\text{alg}}^{(2)}(\boldsymbol{\Xi}, \mathbf{w}) = \mathbf{w}_1 + [w_2]\boldsymbol{\Xi} + (\boldsymbol{\Xi} \otimes \boldsymbol{\Xi})\mathbf{w}_3, \quad (1.22)$$

where $\mathbf{w} = \{\mathbf{w}_1, [w_2], \mathbf{w}_3\}$ has once again to be identified to represent as well as possible the behavior of $\mathbf{C}$.

In the same category, the methods based on the Information Theory and the Maximum Entropy Principle (MEP) have been developed (see [8] and [9]) to compute p_C from the only available statistical information of the random vector $\mathbf{C}$. This information can be seen as the admissible set $\mathcal{C}^{\text{ad}}$ for p_C :

$$\mathcal{C}^{\text{ad}} = \left\{ p_C \in \mathcal{F}(\mathcal{D}, \mathbb{R}^*) \mid \int_{\mathcal{D}} p_C(\mathbf{x})d\mathbf{x} = 1, \right. \\ \left. \forall 1 \leq n \leq N, \int_{\mathcal{D}} \mathbf{g}_n(\mathbf{x})p_C(\mathbf{x})d\mathbf{x} = \mathbf{f}_n \right\}, \quad (1.23)$$

where $\{\mathbf{f}_n, 1 \leq n \leq N\}$ gathers N vectors which are respectively associated with the vector-valued functions $\{\mathbf{g}_n, 1 \leq n \leq N\}$. Hence, the MPE allows building p_C as the solution of the optimization problem:

$$p_C = \arg \max_{p_C \in \mathcal{C}^{\text{ad}}} \left\{ - \int_{\mathcal{D}} p_C(\mathbf{x}) \log(p_C(\mathbf{x})) d\mathbf{x} \right\}. \quad (1.24)$$

As an example, if the maximum available information about $\mathbf{C}$ is the fact that its realizations are in the hypercube $[-1, 1]^M$, the admissible set $\mathcal{C}^{\text{ad}}$ for p_C becomes:

$$\mathcal{C}^{\text{ad}} = \left\{ p_{\mathbf{C}} \in \mathcal{F}([-1, 1]^M, \mathbb{R}^*), \mid \int_{[-1, 1]^M} p_{\mathbf{C}}(\mathbf{x}) d\mathbf{x} = 1 \right\}, \quad (1.25)$$

and it can be shown that the PDF $p_{\mathbf{C}}$ that maximizes the optimization problem defined by Eq. (1.24) is the uniform PDF over $[-1, 1]^M$:

$$p_{\mathbf{C}}(\mathbf{x}) = \frac{1}{2^M}. \quad (1.26)$$

On the other hand, the indirect methods allow the construction of the PDF $p_{\mathbf{C}}$ of the considered random vector $\mathbf{C}$ thanks to a transformation $\mathbb{T}$ of a known PDF $p_{\boldsymbol{\xi}}$ of a random vector $\boldsymbol{\xi} = (\xi_1, \dots, \xi_{N_g})$ of given dimension $N_g \leq M$:

$$\mathbf{C} = \mathbf{t}(\boldsymbol{\xi}), \quad (1.27)$$

$$p_{\mathbf{C}} = \mathbb{T}(p_{\boldsymbol{\xi}}). \quad (1.28)$$

The construction of the transformation $\mathbf{t}$ is thus the key point of these indirect methods. In this context, the isoprobabilistic transformations such as the Nataf transformation (see [46]) or the Rosenblatt transformation (see [47]) have allowed the development of interesting results in the second part of the twentieth century but are still limited to very small dimension cases and not to the high dimension case considered in this work. Nowadays, the most popular indirect methods are the polynomial chaos expansion (PCE) methods, which have been first introduced by Wiener [48] for stochastic processes, and pioneered by Ghanem and Spanos [49, 22] for the use of it in computational sciences. In the last decade, this very promising method has thus been applied in many works (see, for instance [50, 51, 52, 53, 54, 55, 56, 57, 58, 59, 21, 60, 32, 61, 62, 63, 64, 65, 66, 67, 68, 69, 70, 71, 72]). The PCE is based on a direct projection of the random vector $\mathbf{C}$ on a chosen Hilbertian basis $\mathcal{B}_{\text{orth}} = \{\psi_j(\boldsymbol{\xi}), 0 \leq j\}$ of all the second-order random vectors with values in $\mathbb{R}^M$:

$$\mathbf{C} = \sum_{j=0}^{+\infty} \mathbf{y}^{(j)} \psi_j(\boldsymbol{\xi}), \quad (1.29)$$

$$E[\psi_j(\boldsymbol{\xi}) \psi_k(\boldsymbol{\xi})] = \delta_{jk}. \quad (1.30)$$

In practical terms, the PCE of $\mathbf{C}$ has to be truncated to its $N + 1$ most influential terms:

$$\mathbf{C} \approx \sum_{j=0}^N \mathbf{y}^{(j)} \psi_j(\boldsymbol{\xi}). \quad (1.31)$$

In particular, in the following, it will be assumed that $\psi_0(\boldsymbol{\xi}) = 1$, such that:

$$\mathbf{y}^{(0)} = E[\mathbf{C}] = \boldsymbol{\mu}_{\mathbf{C}}. \quad (1.32)$$

A method to choose these N particular terms and to quantify the amplitude of the truncation residue, $\sum_{j=N+1}^{+\infty} \mathbf{y}^{(j)} \psi_j(\boldsymbol{\xi})$, has therefore to be defined. Building the transformation $\mathbf{t}$ requires at last the construction of N deterministic coefficients, $\{\mathbf{y}^{(j)}, 1 \leq j \leq N\}$, from the available information about $\mathbf{C}$.

It has to be noticed that in such an approach, any distribution for $\boldsymbol{\xi}$ can be chosen. For instance, if the components of $\boldsymbol{\xi}$ are independent and uniformly distributed between -1 and

1, the corresponding Hilbertian basis, $\{\psi_j(\boldsymbol{\xi}), 1 \leq j\}$, is the set of the normalized Legendre polynomials.

When trying to identify in inverse the multidimensional distribution of an *a priori* non-Gaussian random vector, the PCE method appears to be very efficient, even when the statistical dimension of $\mathbf{C}$ is high. Indeed, this method can be applied to any random vector, is not based on *a priori* formulations, and allows a very easy generation of independent realizations of $\mathbf{C}$, once the projection coefficients are identified. Indeed each independent realization of germ $\boldsymbol{\xi}$ leads to an independent realization of $\mathbf{C}$.

1.5 PCE identification of random vectors

In this section, a description of the PCE identification with respect to an arbitrary measure is given. The objective is to summarize the different key steps of the PCE identification method and the way they can be practically implemented.

After having defined the theoretical frame of the PCE identification, the cost-functions that lead us to the computation of the PCE coefficients $\{\mathbf{y}^{(1)}, \dots, \mathbf{y}^{(N)}\}$ are presented, for a given truncation parameter N . Two cases are distinguished: the direct case, for which the PCE germ $\boldsymbol{\xi}$ is known, and the indirect case, for which the PCE germ is unknown. At last, to justify the choice of this truncation parameter, a method to perform the convergence analysis is introduced.

1.5.1 Theoretical frame

Let $\mathbf{C} = (C_1, \dots, C_M)$ be an element of the space $L^2_{\mathcal{P}}(\Theta, \mathbb{R}^M)$ of all the second-order M -dimension random vectors defined on the probability space $(\Theta, \mathcal{T}, \mathcal{P})$ with values in $\mathbb{R}^M$, equipped with the inner product $\langle \cdot, \cdot \rangle$. It is assumed that ν independent realizations, $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$, of $\mathbf{C}$ are known and gathered in the $(M \times \nu)$ real matrix $[\mathbf{C}^{\text{exp}}(\nu)]$:

$$[\mathbf{C}^{\text{exp}}(\nu)] = [\mathbf{C}(\theta_1) \ \dots \ \mathbf{C}(\theta_\nu)]. \quad (1.33)$$

Equation (1.31) can be rewritten as:

$$\mathbf{C} - \boldsymbol{\mu}_{\mathbf{C}} \approx \mathbf{C}^{\text{chaos}}(N) = [\mathbf{y}]\boldsymbol{\Psi}(\boldsymbol{\xi}), \quad (1.34)$$

$$[\mathbf{y}] = [\mathbf{y}^{(1)} \ \dots \ \mathbf{y}^{(N)}], \quad \boldsymbol{\Psi}(\boldsymbol{\xi}) = (\psi_1(\boldsymbol{\xi}), \dots, \psi_N(\boldsymbol{\xi})). \quad (1.35)$$

The orthonormality property of the projection basis $\{\psi_j(\boldsymbol{\xi}), 1 \leq j \leq N\}$ yields the condition:

$$E[\boldsymbol{\Psi}(\boldsymbol{\xi}, p) \otimes \boldsymbol{\Psi}(\boldsymbol{\xi}, p)] = [I_N], \quad (1.36)$$

where $[I_N]$ is the $(N \times N)$ identity matrix. Let $[R_{\mathbf{C}\mathbf{C}}^{\text{chaos}}(N)]$ be the covariance matrix of centered random vector $\mathbf{C}^{\text{chaos}}(N)$:

$$[R_{\mathbf{C}\mathbf{C}}^{\text{chaos}}(N)] = E[\mathbf{C}^{\text{chaos}}(N) \otimes \mathbf{C}^{\text{chaos}}(N)] = [\mathbf{y}]E[\boldsymbol{\Psi}(\boldsymbol{\xi}, p) \otimes \boldsymbol{\Psi}(\boldsymbol{\xi}, p)][\mathbf{y}]^T = [\mathbf{y}][\mathbf{y}]^T. \quad (1.37)$$

To simplify the notations, it is supposed in the following that $\mathbf{C}$ is a centered random vector, such that:

$$\mathbf{y}^{(0)} = \boldsymbol{\mu}_{\mathbf{C}} = \mathbf{0}. \quad (1.38)$$

No distinction is therefore made between the covariance and the autocorrelation matrices of $\mathbf{C}$ in the next sections.

1.5.2 Identification of the polynomial chaos expansion coefficients

In this section, a particular choice for the N_g -dimension PCE germ, $\boldsymbol{\xi} = (\xi_1, \dots, \xi_{N_g})$, and a particular value of the truncation parameter N are considered. Let $[\Psi(\nu^{\text{chaos}})]$ be the $(N \times \nu^{\text{chaos}})$ real matrix of independent realizations of the truncated PCE basis $\Psi(\boldsymbol{\xi})$:

$$[\Psi(\nu^{\text{chaos}})] = [\Psi(\boldsymbol{\xi}(\Theta_1)) \cdots \Psi(\boldsymbol{\xi}(\Theta_{\nu^{\text{chaos}}}))], \quad (1.39)$$

where the set $\{\boldsymbol{\xi}(\Theta_1), \dots, \boldsymbol{\xi}(\Theta_{\nu^{\text{chaos}}})\}$ gathers ν^{chaos} independent realizations of random vector $\boldsymbol{\xi}$. As a direct consequence of the orthonormality of the PCE vector $\Psi(\boldsymbol{\xi})$, matrix $[\Psi(\nu^{\text{chaos}})]$ verifies the asymptotic property:

$$\lim_{\nu^{\text{chaos}} \rightarrow +\infty} \frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T = E[\Psi(\boldsymbol{\xi}) \otimes \Psi(\boldsymbol{\xi})] = [I_N]. \quad (1.40)$$

Direct identification

If the realizations of $\mathbf{C}$ are solutions of a mechanical system, and if $\boldsymbol{\xi}$ corresponds to the variable inputs of this system, then $\nu = \nu^{\text{chaos}}$ and both realizations of $\mathbf{C}$, $\{\mathbf{C}(\Theta_1), \dots, \mathbf{C}(\Theta_{\nu^{\text{chaos}}})\}$, and $\Psi(\boldsymbol{\xi})$, $\{\Psi(\boldsymbol{\xi}(\Theta_1)), \dots, \Psi(\boldsymbol{\xi}(\Theta_{\nu^{\text{chaos}}}))\}$, are known at the same time. They verify:

$$[C^{\nu^{\text{chaos}}}] = [\mathbf{C}(\Theta_1) \cdots \mathbf{C}(\Theta_{\nu^{\text{chaos}}})] \approx [C^{\text{chaos}}(N)] = [y][\Psi(\nu^{\text{chaos}})]. \quad (1.41)$$

In this case, two classical methods are generally used to identify such coefficient matrix $[y]$:

- **Methods based on the empirical estimation of the mean function.** From Eq. (1.31), as family $\{\psi_j(\boldsymbol{\xi}), 1 \leq j\}$ is orthonormal, it can be seen that for all $1 \leq j \leq N$:

$$\begin{aligned} [y] &= E[\mathbf{C} \otimes \Psi(\boldsymbol{\xi})] \\ &\approx [y_1^{\text{opt}}(\nu^{\text{chaos}})] = \frac{1}{\nu^{\text{chaos}}} \sum_{p=1}^{\nu^{\text{chaos}}} \mathbf{C}(\Theta_p) \otimes \Psi(\boldsymbol{\xi}(\Theta_p)) = \frac{1}{\nu^{\text{chaos}}} [C^{\text{chaos}}(N)][\Psi(\nu^{\text{chaos}})]^T. \end{aligned} \quad (1.42)$$

- **Regression-based methods.** Let $\mathcal{C}([y], \nu^{\text{chaos}})$ be the cost function that quantifies the mean-square distance between $\mathbf{C}$ and its PCE approximation, $C^{\text{chaos}}(N)$, defined by:

$$\begin{aligned} \mathcal{C}([y], \nu^{\text{chaos}}) &= \left\| [C^{\text{chaos}}(N)] - [y][\Psi(\nu^{\text{chaos}})] \right\|^2 \\ &\stackrel{\text{def}}{=} \text{Tr} \left[\left([C^{\text{chaos}}(N)] - [y][\Psi(\nu^{\text{chaos}})] \right) \left([C^{\text{chaos}}(N)] - [y][\Psi(\nu^{\text{chaos}})] \right)^T \right], \end{aligned} \quad (1.43)$$

with $\text{Tr}[\cdot]$ the trace operator. PCE matrix $[y]$ can therefore be searched as the argument that minimizes $\mathcal{C}([y], \nu^{\text{chaos}})$. The cost function $\mathcal{C}([y], \nu^{\text{chaos}})$ being convex, it admits a minimum, $[y_2^{\text{opt}}(\nu^{\text{chaos}})]$, which verifies:

$$[y] \approx [y_2^{opt}(\nu^{\text{chaos}})] = \arg \min_{[y]} \left\{ \mathcal{C}([y], \nu^{\text{chaos}}) \right\}, \quad (1.44)$$

$$[y_2^{opt}(\nu^{\text{chaos}})] = [C^{\text{chaos}}(N)][\Psi(\nu^{\text{chaos}})]^T \left([\Psi(\nu^{\text{chaos}})][\Psi(\nu^{\text{chaos}})]^T \right)^{-1}. \quad (1.45)$$

From Eqs (1.40), (1.42) and (1.45), it can be directly verified that the two former methods give asymptotically the same results:

$$\lim_{\nu^{\text{chaos}} \rightarrow +\infty} [y_1^{opt}(\nu^{\text{chaos}})] = \lim_{\nu^{\text{chaos}} \rightarrow +\infty} [y_2^{opt}(\nu^{\text{chaos}})]. \quad (1.46)$$

Indirect identification

If $\mathbf{C}$ is a random vector that gathers the projection coefficients of a random field X on a particular basis, as it is the case in this thesis, the realizations of $\mathbf{C}$ are deduced from the available realizations of X , such that there is *a priori* no direct link between the two sets of realizations of $\boldsymbol{\xi}$ and $\mathbf{C}$. Alternative methods have thus to be used to identify $[y]$.

To this end, let $\mathcal{M}_{M,N}$ be the space of all the $(M \times N)$ real matrices. For a given value of $[y^*]$ in $\mathcal{M}_{M,N}$, the random vector $\mathbf{U}([y^*]) = [y^*]\boldsymbol{\Psi}(\boldsymbol{\xi})$ is a centered M -dimension random vector, for which the autocorrelation is equal to $[y^*][y^*]^T$. Let $p_{\mathbf{U}([y^*])}$ be its multidimensional PDF.

When the only available information about $\mathbf{C}$ is limited to a set of ν independent realizations, the most general and relevant method to identify in inverse the optimal coefficients matrix $[y]$, is to search it as the argument that maximizes the log-likelihood $\mathcal{L}_{\mathbf{U}([y^*])}([C^{\text{exp}}(\nu)])$ of $\mathbf{U}([y^*])$ at the experimental points gathered in $[C^{\text{exp}}(\nu)]$:

$$[y] = \arg \max_{[y^*] \in \mathcal{M}_{M,N}} \mathcal{L}_{\mathbf{U}([y^*])}([C^{\text{exp}}(\nu)]), \quad (1.47)$$

$$\mathcal{L}_{\mathbf{U}([y^*])}([C^{\text{exp}}(\nu)]) = \sum_{n=1}^{\nu} \log p_{\mathbf{U}([y^*])}(\mathbf{C}(\theta_n)). \quad (1.48)$$

1.5.3 Practical solving of the log-likelihood maximization

Solving the optimization problem defined by Eq. (1.47) has required the development of specific algorithms, which are described in this section.

The need for statistical algorithms to maximize the log-likelihood

The log-likelihood $\mathcal{L}_{\mathbf{U}([y^*])}([C^{\text{exp}}(\nu)])$ being nonconvex, deterministic algorithms such as gradient algorithms cannot be applied to solve Eq. (1.47), and random search algorithms have to be used. Hence, the precision of the PCE has to be correlated to a numerical cost Z , which corresponds to a number of independent trials of $[y^*]$ in $\mathcal{M}_{M,N}$. The higher the value of Z is, the better the PCE identification should be. Therefore, this value has to be chosen as high as possible while respecting the computational resource limitation. Let $\mathcal{Y} = \{[y^*]^{(z)}, 1 \leq z \leq Z\}$ be a set of Z elements, which have been chosen randomly in $\mathcal{M}_{M,N}$. For a given numerical cost Z , the most accurate PCE coefficients matrix $[y]$ is approximated by:

$$[y] \approx [y_{\mathcal{Y}}] = \arg \max_{[y^*] \in \mathcal{Y}} \mathcal{L}_{\mathbf{U}([y^*])}([C^{\text{exp}}(\nu)]). \quad (1.49)$$

Restriction of the maximization domain

From the ν independent realizations $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$, the covariance matrix $[R_{CC}]$ of $\mathbf{C}$ can be estimated by:

$$[R_{CC}] \approx [\hat{R}_{CC}(\nu)] = \frac{1}{\nu} \sum_{n=1}^{\nu} \mathbf{C}(\theta_n) \otimes \mathbf{C}(\theta_n) = \frac{1}{\nu} [C^{\text{exp}}(\nu)][C^{\text{exp}}(\nu)]^T. \quad (1.50)$$

A good way to improve the efficiency of the numerical identification of $[y]$ is then to restrict the research set to $\mathcal{O}_C \subset \mathcal{M}_{M,N}$, with:

$$\mathcal{O}_C = \left\{ [y^*] = [\mathbf{y}^{*,(1)}, \dots, \mathbf{y}^{*,(N)}] \in \mathcal{M}_{M,N} \mid [y^*][y^*]^T = [\hat{R}_{CC}(\nu)] \right\}, \quad (1.51)$$

which, taking into account Eq. (1.37), guarantees by construction that:

$$[R_{CC}^{\text{chaos}}(N)] = [\hat{R}_{CC}(\nu)]. \quad (1.52)$$

Hence, the PCE coefficients matrix $[y]$ can be approximated as the argument in $\mathcal{O}_C$ that maximizes the log-likelihood $\mathcal{L}_{U([y^*])}([C^{\text{exp}}(\nu)])$. By defining $\mathcal{W}$ the set that gathers Z randomly raised elements of $\mathcal{O}_C$, $[y]$ can then be assessed as the solution of the new optimization problem:

$$[y] \approx [y_{\mathcal{W}}] = \arg \max_{[y^*] \in \mathcal{W}} \mathcal{L}_{U([y^*])}([C^{\text{exp}}(\nu)]). \quad (1.53)$$

Approximation of the log-likelihood function

From a particular matrix of realizations $[\Psi(\nu^{\text{chaos}})]$ (which is defined in Eq. (1.39)), if $[y^*]$ is an element of $\mathcal{O}_C$, ν^{chaos} independent realizations $\{\mathbf{U}([y^*], \Theta_p) = [y^*]\Psi(\boldsymbol{\xi}(\Theta_p)), 1 \leq p \leq \nu^{\text{chaos}}\}$ of random vector $\mathbf{U}([y^*])$ can be computed and gathered in the matrix $[U]$:

$$[U] = [\mathbf{U}([y^*], \Theta_1) \quad \dots \quad \mathbf{U}([y^*], \Theta_{\nu^{\text{chaos}}})] = [y^*][\Psi(\nu^{\text{chaos}})]. \quad (1.54)$$

Hence, using Gaussian Kernels, the PDF $p_{U([y^*])}$ of $\mathbf{U}([y^*])$ can be directly estimated by its non parametric estimator $\hat{p}_U$:

$$\forall \mathbf{x} \in \mathbb{R}^M, \quad p_{U([y^*])}(\mathbf{x}) \approx \hat{p}_U(\mathbf{x}) = \frac{1}{(2\pi)^{M/2} \nu^{\text{chaos}} \prod_{m=1}^M h_m} \sum_{p=1}^{\nu^{\text{chaos}}} \exp \left(-\frac{1}{2} \sum_{m=1}^M \left(\frac{x_m - U_m([y^*], \Theta_p)}{h_m} \right)^2 \right), \quad (1.55)$$

where $\mathbf{h} = (h_1, \dots, h_M)$ is the multidimensional optimal Silverman bandwidth vector (see [2]) of the Kernel smoothing estimation of $p_{U([y^*])}$:

$$\forall 1 \leq m \leq M, \quad h_m = \hat{\sigma}_{U_m} \left(\frac{4}{(2+M)\nu^{\text{chaos}}} \right)^{1/(M+4)}, \quad (1.56)$$

where $\hat{\sigma}_{U_m}$ is the empirical estimation of the standard deviation of each component U_m of $\mathbf{U}$. It has to be noticed that $\hat{p}_U$ only depends on the bandwidth vector $\mathbf{h}$, and the two matrices $[y^*]$ and $[\Psi(\nu^{\text{chaos}})]$. Hence, according to the Eqs. (1.48), (1.54) and (1.55), for a given value of ν^{chaos} , the maximization of the log-likelihood function $\mathcal{L}_{U([y^*])}$ can be replaced by the maximization of the cost-function $\mathcal{C}([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})])$ such that:

$$[y] \approx [y_{\mathcal{O}_C}] = \arg \max_{[y^*] \in \mathcal{O}_C} \mathcal{C}([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})]), \quad (1.57)$$

where:

$$\mathcal{C}([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})]) = \mathcal{C}_C + \mathcal{C}_V([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})]), \quad (1.58)$$

$$\mathcal{C}_C = -\nu \ln \left((2\pi)^{M/2} \nu^{\text{chaos}} \prod_{m=1}^M h_m \right), \quad (1.59)$$

$$\mathcal{C}_V([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})]) = \sum_{n=1}^{\nu} \ln \left(\sum_{p=1}^{\nu^{\text{chaos}}} \exp \left(-\frac{1}{2} \sum_{m=1}^M \left(\frac{C_m(\theta_n) - U_m([y^*], \Theta_p)}{h_m} \right)^2 \right) \right). \quad (1.60)$$

Hence, the optimization problem defined by Eq. (1.53) can finally be estimated by:

$$[y] \approx [y_{\mathcal{O}_C}^Z] = \arg \max_{[y^*] \in \mathcal{W}} \mathcal{C}([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})]). \quad (1.61)$$

Accuracy of the PCE identification

For a given computation cost Z and a given value for the truncation parameter N , let $[y_{\mathcal{O}_C}^Z]$ be an optimal solution of Eq. (1.61). $[y_{\mathcal{O}_C}^Z]$ is a numerical estimation of the PCE coefficients matrix $[y]$. For a new $(N \times \nu^{\text{chaos},*})$ real matrix $[\Psi^*(\nu^{\text{chaos},*})]$ of independent realizations ($\nu^{\text{chaos},*}$ can be higher than ν^{chaos}), the robustness of $[y_{\mathcal{O}_C}^Z]$ regarding the choice of $[\Psi(\nu^{\text{chaos}})]$ can then be estimated by comparing $\mathcal{C}([C^{\text{exp}}(\nu)], [y_{\mathcal{O}_C}^Z], [\Psi(\nu^{\text{chaos}})])$ and $\mathcal{C}([C^{\text{exp}}(\nu)], [y_{\mathcal{O}_C}^Z], [\Psi^*(\nu^{\text{chaos},*})])$. If ν new independent realizations of $\mathbf{C}$ were available and gathered in the matrix $[C^{\text{exp,new}}(\nu)]$, the over-learning of the method could be measured by comparing $\mathcal{C}([C^{\text{exp}}(\nu)], [y_{\mathcal{O}_C}^Z], [\Psi(\nu^{\text{chaos}})])$ and $\mathcal{C}([C^{\text{exp,new}}(\nu)], [y_{\mathcal{O}_C}^Z], [\Psi(\nu^{\text{chaos}})])$. At last, for the same value for Z , if $[y_{\mathcal{O}_C}^{Z,\text{new}}]$ is a new optimal solution of Eq. (1.61), the global accuracy of the identification stems from the comparison between $\mathcal{C}([C^{\text{exp,new}}(\nu)], [y_{\mathcal{O}_C}^Z], [\Psi^*(\nu^{\text{chaos},*})])$ and $\mathcal{C}([C^{\text{exp,new}}(\nu)], [y_{\mathcal{O}_C}^{Z,\text{new}}], [\Psi^*(\nu^{\text{chaos},*})])$.

1.5.4 Identification of the PCE truncation parameters

As shown in Section 1.4, two truncation parameters, N_g and N , appear in the truncated PCE, $\mathbf{C}^{\text{chaos}}(N)$, of $\mathbf{C}$. A method to choose the size N_g and these N elements from basis $\mathcal{B}_{\text{orth}}$ as well as a method to quantify the relevance of such a N -dimension basis have thus to be defined.

Restriction of the admissible projection basis

In this work, only polynomial basis are addressed, such that for $1 \leq j$, a particular element $\psi_j(\boldsymbol{\xi})$ in $\mathcal{B}_{\text{orth}}$ can be written under the form:

$$\psi_j(\boldsymbol{\xi}) = \sum_{q=1}^{+\infty} c_q^{(j)} \xi_1^{\alpha_1^{(q)}} \times \cdots \times \xi_{N_g}^{\alpha_{N_g}^{(q)}}, \quad (\alpha_1^{(q)}, \dots, \alpha_{N_g}^{(q)}) \in \mathbb{N}^{N_g}. \quad (1.62)$$

In addition, the classical assumption that the most influential elements of $\mathcal{B}_{\text{orth}}$ correspond to the elements of lowest total polynomial order is introduced in this work. Let p be the maximal polynomial order of the projection basis, such that for $1 \leq j \leq N$, we choose:

$$\psi_j(\boldsymbol{\xi}) = \sum_{q=1}^N c_q^{(j)} \xi_1^{\alpha_1^{(q)}} \times \cdots \times \xi_{N_g}^{\alpha_{N_g}^{(q)}}, \quad \sum_{\ell=1}^{N_g} \alpha_{\ell}^{(q)} \leq p. \quad (1.63)$$

Given this choice for the extraction of N elements in $\mathcal{B}_{\text{orth}}$, it can be seen that N increases very quickly with N_g and p , as:

$$N = (N_g + p)! / (N_g! p!). \quad (1.64)$$

Definition of a log error function

For each component $C_m^{\text{chaos}}(N)$ of the truncated PCE, $\mathbf{C}^{\text{chaos}}(N) = [y]\boldsymbol{\Psi}(\boldsymbol{\xi})$, of $\mathbf{C}$, the L^1 -log error function err_m is introduced as described in [38]:

$$\forall 1 \leq m \leq M, \quad err_m(N_g, p) = \int_{\text{BI}_m} |\log_{10}(p_{C_m}(x_m)) - \log_{10}(p_{C_m^{\text{chaos}}}(x_m))| dx_m, \quad (1.65)$$

where:

- BI_m is the support of the kernel estimator of p_{C_m} . This bounding domain has thus to be adapted to the available realizations of $\mathbf{C}$, which are gathered in $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_{\nu})\}$;
- p_{C_m} and $p_{C_m^{\text{chaos}}}$ are the PDF of C_m and C_m^{chaos} respectively.

The multidimensional error function $err(N_g, p)$ is then deduced from the unidimensional L^1 -log error function as:

$$err(N_g, p) = \sum_{m=1}^M err_m(N_g, p). \quad (1.66)$$

The parameters N_g and p have thus to be determined to minimize the multidimensional L^1 -log error function $err(N_g, p)$.

For given values of truncation parameters N_g and p , it is reminded that PCE coefficients matrix $[y]$ is searched in order to maximize the multidimensional log-likelihood function, which allows us to consider *a priori* strongly correlated problems. Once this matrix $[y]$ is identified, it is possible to generate as many independent realizations of truncated PCE $\mathbf{C}^{\text{chaos}}(N)$ as needed to estimate as precisely as possible the non parametric estimator $\hat{p}_{\mathbf{U}}$ of its multidimensional PDF. The number ν of available experimental realizations of $\mathbf{C}$ is however limited. This number is generally too small for the non parametric estimator of multidimensional PDF $p_{\mathbf{C}}$ of $\mathbf{C}$ to be relevant, whereas it is most of the time large enough to define the estimators of the marginals of $p_{\mathbf{C}}$. Therefore, the log-error functions defined by Eqs. (1.65) and (1.66) only consider the marginals of the PDF of $p_{\mathbf{C}}$ and $p_{\mathbf{C}^{\text{chaos}}}$. In addition, the logarithm function has been introduced in order to measure the errors on the tails of the probability density function.

Definition of an admissible set for the truncation parameters

As it exists an isoprobabilistic transformation between $\mathbf{C}$ and $(\Xi_1, \dots, \Xi_M)$, where the set $\{\Xi_m, 1 \leq m \leq M\}$ gathers M independent centered normalized Gaussian random variables, the convergence analysis can be restricted to the values of N_g which verify:

$$N_g \leq M. \quad (1.67)$$

Moreover, imposing the $(M \times N)$ real matrix $[y]$ to be in $\mathcal{O}_C$ amounts to imposing $\frac{M(M+1)}{2}$ constraints on $[y]$, which implies:

$$MN \geq \frac{M(M+1)}{2} \Leftrightarrow N \geq \frac{M+1}{2}. \quad (1.68)$$

The set $\mathcal{Q}(M)$ of the admissible values for p and N_g is thus:

$$\mathcal{Q}(M) = \{(p, N_g) \in \mathbb{N}^2, \mid N_g \leq M, N = (N_g + p)! / (N_g! p!) \geq (M+1)/2\}. \quad (1.69)$$

Theoretically, increasing p and N_g adds terms in the PCE of the considered random vector, and therefore should induce the decrease of the error function:

$$\forall p^* \geq p, N_g^* \geq N_g, \quad err(N_g, p) \geq \max\{err(N_g^*, p), err(N_g, p^*)\} \geq \min\{err(N_g^*, p), err(N_g, p^*)\} \geq err(N_g^*, p^*). \quad (1.70)$$

However, the higher the values of p and N_g are, the bigger the PCE coefficients matrix is, the harder the numerical identification is. Hence, introducing ε as an error threshold, which has to be adapted to the problem, let $\mathcal{P}(\varepsilon, M)$ be the set:

$$\mathcal{P}(\varepsilon, M) = \{(p, N_g) \in \mathcal{Q}(M) \mid err(N_g, p) \leq \varepsilon\}. \quad (1.71)$$

Finally, given the error threshold ε , rather than directly minimizing the L^1 -log error function $err(N_g, p)$, it appears to be more accurate to look for the optimal values of p and N_g that minimize the size of the projection basis $N = (N_g + p)! / (N_g! p!)$:

$$(p, N_g) = \arg \min_{(p^*, N_g^*) \in \mathcal{P}(\varepsilon, M)} (N_g^* + p^*)! / (N_g^*! p^*!). \quad (1.72)$$

If the polynomial order (which is a priori unknown) of the non truncated PCE of $\mathbf{C}$ is infinite, it may not exist values of p and N_g in $\mathcal{P}(\varepsilon, M)$ for error function $err(N_g, p)$ to be inferior to small values of ε . In this case, the former algorithms can nevertheless be used to find the most accurate values of p and N_g with respect to an available computational cost.

1.6 Conclusions

In this chapter, it has been shown that nowadays most promising methods to identify in inverse the statistical properties of a non-Gaussian and non-stationary random field X , when the available information about X is a set of ν independent realizations, are based on a double decomposition. First, thanks to a KL expansion, the statistical properties of X can be condensed through its projection on a particularly well-adapted orthonormal reduced basis, $\{k_m, 1 \leq m \leq M\}$, such that

$$X \approx \sum_{m=1}^M C_m k_m. \quad (1.73)$$

Secondly, from the ν independent realizations of X , it has been shown that ν independent realizations of random vector $\mathbf{C} = (C_1, \dots, C_M)$ can be deduced. Based on this available

information, it has been emphasized that a PCE-based approach allows the identification of $\mathbf{C}$, such that:

$$X \approx \sum_{m=1}^M \sum_{j=1}^N k_m y_m^{(j)} \psi_j(\boldsymbol{\xi}), \quad (1.74)$$

where $\boldsymbol{\xi} = (\xi_1, \dots, \xi_{N_g})$ is a random vector for which distribution is known and chosen. In such two-step approach, three truncation parameters, M , N and N_g , have been introduced, which have to be identified from convergence analysis. Advanced methods and algorithms to identify the projection basis k_m and coefficients $y_m^{(j)}$ for $1 \leq m \leq M$ and $1 \leq j \leq N$ from a set of ν independent realizations of X have been presented in detail.

At last, from each realization of $\boldsymbol{\xi}$, such a method gives access to a realistic realization of X that is representative of the set of its available realizations.

Chapter 2

Optimal reduced basis for random fields defined by a set of realizations

2.1 Introduction

The use of reduced basis has spread to many scientific fields for the last decades to condense the statistical properties of the random fields, which are written $X = \{X(s), s \in \Omega \subset \mathbb{R}\}$ in this work, and for which mean value is assumed to be zero. Among these basis, the classical Karhunen-Loève basis associated with X , $\{k_m, 1 \leq m\}$, which has been introduced in Chapter 1, corresponds to the Hilbertian basis that is constructed as the eigenfunctions of the covariance operator of X , R_{XX} . The importance of this basis stems from its optimality in the sense that it minimizes the total mean square error. In most of the applications based on random fields, the knowledge of these random fields is however limited. Indeed, their statistical properties are generally known through a set of ν independent realizations, $\{X(\theta_1), \dots, X(\theta_\nu)\}$, which stem from experimental measurements. In these cases, the covariance operator is not perfectly known but can only be estimated. If we define $\widehat{R}_{XX}$ as the empirical estimator of R_{XX} , there is however no reason for the eigenfunctions of $\widehat{R}_{XX}$ to be still optimal. In reply to this concern, this chapter presents an adaptation of the Karhunen-Loève expansion to identify, in inverse, projection families that are as relevant as possible for X , even if the number of available realizations, ν , is relatively small. This method is first based on an innovative technique to *a posteriori* evaluate the projection errors for X , and secondly, on an original optimization problem that can be seen as an extension of the classical Fredholm equation.

In Section 2.2, the theoretical frame of this chapter is described. Section 2.3 introduces then the method we propose to identify optimized projection basis from a set of independent realizations. At last, Section 2.4 illustrates the possibilities of such a method on an application based on simulated data.

2.2 Theoretical frame

2.2.1 Quantification of the relevance of a projection basis

Let $(\Theta, \mathcal{C}, P)$ be a probability space and $\mathcal{P}(\Omega)$ be the space of all the second-order $\mathbb{R}$ -valued random fields, indexed by the compact interval $\Omega = [0, S]$, where $S < +\infty$. The space $\mathbb{H} = L^2(\Omega, \mathbb{R})$ denotes moreover the space of square integrable functions on Ω , with values in $\mathbb{R}$, equipped with the inner product $(\cdot, \cdot)$, such that for all u and v in $\mathbb{H}$,

$$(u, v) = \int_{\Omega} u(s)v(s)ds. \quad (2.1)$$

Let $X = \{X(s), s \in \Omega\}$ be an element of $\mathcal{P}(\Omega)$, for which ν independent realizations, $\{X(\theta_1), \dots, X(\theta_\nu)\}$, are supposed to be known. Without loss of generality, it is once again supposed that the mean value of X is equal to zero:

$$E[X(s)] = 0, \quad \forall s \in \Omega. \quad (2.2)$$

It is assumed that the covariance function, R_{XX} , of centered random field X is square integrable on $\Omega \times \Omega$,

$$\int_{\Omega} \int_{\Omega} R_{XX}(s, s')^2 ds ds' < +\infty. \quad (2.3)$$

Let $\mathcal{B} = \{b_m(s), s \in \Omega\}_{m \geq 1}$, be a Hilbertian basis of $\mathbb{H}$, such that:

$$X = \sum_{m \geq 1} C_m b_m, \quad (2.4)$$

$$(b_m, b_p) = \delta_{mp}, \quad C_m = (X, b_m), \quad (2.5)$$

where the projection coefficients, $\{C_m, m \geq 1\}$, are centered random variables that are statistically dependent and *a priori* correlated. For practical purposes, this basis has to be truncated. For all $M \geq 1$, $\hat{X}^{\mathcal{B}^{(M)}}$ is thus introduced as the projection of X on the subset $\mathcal{B}^{(M)} = \{b_m, 1 \leq m \leq M\} \subset \mathcal{B}$:

$$\hat{X}^{\mathcal{B}^{(M)}} = \sum_{m=1}^M C_m b_m = \mathbf{b}^T \mathbf{C}, \quad (2.6)$$

$$\mathbf{b} = (b_1, \dots, b_M), \quad \mathbf{C} = (C_1, \dots, C_M). \quad (2.7)$$

The relevance of $\mathcal{B}^{(M)}$ to characterize X is analyzed with respect to the normalized L^2 -error, that is denoted by ε^2 , such that:

$$\begin{aligned} \varepsilon^2(\mathcal{B}^{(M)}) &= \left\| X - \hat{X}^{\mathcal{B}^{(M)}} \right\|_{\mathcal{P}(\Omega)}^2 / \|X\|_{\mathcal{P}(\Omega)}^2 \\ &= 1 - \frac{1}{\|X\|_{\mathcal{P}(\Omega)}^2} \sum_{m \leq M} E[C_m^2], \end{aligned} \quad (2.8)$$

where $\|\cdot\|_{\mathcal{P}(\Omega)}$ is the L_2 norm in $\mathcal{P}(\Omega)$, such that:

$$\|Y\|_{\mathcal{P}(\Omega)}^2 = E \left[\int_{\Omega} Y^2(s) ds \right], \quad Y \in \mathcal{P}(\Omega). \quad (2.9)$$

Therefore, if $\mathcal{B}_1 = \{b_m^{(1)}, m \geq 1\}$ and $\mathcal{B}_2 = \{b_m^{(2)}, m \geq 1\}$ are two distinct Hilbertian basis of $\mathbb{H}$, the family $\mathcal{B}_1^{(M_1)} = \{b_m^{(1)}, 1 \leq m \leq M_1\}$ is said to be more relevant than the family $\mathcal{B}_2^{(M_2)} = \{b_m^{(2)}, 1 \leq m \leq M_2\}$ (M_1 can be greater or smaller than M_2) to characterize X if and only if:

$$\varepsilon^2(\mathcal{B}_1^{(M_1)}) \leq \varepsilon^2(\mathcal{B}_2^{(M_2)}). \quad (2.10)$$

2.2.2 Optimality of the Karhunen-Loève expansion

As presented in Section 1.3.2, due to the orthogonal projection theorem in Hilbert space, the Karhunen-Loève basis associated with X , that was denoted by $\mathcal{K} = \{k_m, m \geq 1\}$, is optimal in the sense that, for all $M \geq 1$, $\mathcal{K}^{(M)} = \{k_m, 1 \leq m \leq M\}$ minimizes error ε^2 among the M -dimension families of $\mathbb{H}$:

$$\mathcal{K}^{(M)} = \arg \min_{\mathcal{B}^{(M)} \in \mathbb{H}^M} \left\{ \varepsilon^2(\mathcal{B}^{(M)}) \right\}. \quad (2.11)$$

Hence, when dealing with correlated random fields, for which the covariance function R_{XX} is known, minimizing error ε^2 amounts to identifying the KL basis associated with X . Once these functions $\{k_m, m \geq 1\}$ have been identified, the projection of X on $\mathcal{K}$ can be written as:

$$X = \sum_{m \geq 1} A_m k_m, \quad (2.12)$$

where, by construction of the KL basis, it can be noticed that, for all $m \geq 1$ and $p \geq 1$:

$$E[A_m A_p] = \delta_{mp} \lambda_m. \quad (2.13)$$

The KL basis associated with X allows therefore the uncorrelation of the projection coefficients, $\{A_m, m \geq 1\}$. Reciprocally, it can directly be shown that if $\mathcal{B}^* = \{b_m^*, m \geq 1\}$ is a basis, for which the projection coefficients, $\{C_m^*, m \geq 1\}$, of X on $\mathcal{B}^*$ are uncorrelated, then functions b_m^* have to be solution of the Fredholm eigenvalue problem defined by Eq. (1.5), such that $\mathcal{B}^* = \mathcal{K}$. Hence, even if R_{XX} is unknown, the uncorrelation of the projection coefficients is a **sufficient** condition for the identification of the Karhunen-Loève basis.

2.2.3 Difficulties concerning the identification of the Karhunen-Loève expansion from independent realizations

Random field X is now supposed to be only known through a set of ν independent realizations, $\{X(\theta_1), \dots, X(\theta_\nu)\}$.

Let $\mathcal{B} = \{b_m, m \geq 1\}$ be a Hilbertian basis of $\mathbb{H}$, such that the projection of X on $\mathcal{B}$ is given by:

$$X = \sum_{m \geq 1} C_m b_m. \quad (2.14)$$

From a theoretical point of view, according to Section 2.2.2, it can be *a posteriori* said that it can be extracted from $\mathcal{B}$ the projection families that minimize error ε^2 , defined by Eq. (2.8), if and only if, for all $m \geq 1$, it exists $\lambda_m \geq 0$, such that one of the two following equivalent conditions is verified:

$$\begin{aligned} 1. & E[C_m C_p] = \delta_{mp} \lambda_m, \quad \forall p \geq 1, \\ 2. & \int_{\Omega} R_{XX}(s, s') b_m(s') ds' = \lambda_m b_m(s), \quad \forall s \in \Omega. \end{aligned} \quad (2.15)$$

From a numerical point of view, when only ν independent realizations, $\{X(\theta_1), \dots, X(\theta_\nu)\}$, of X are available, the *a priori* best evaluations of the covariance function of X and of the mean values $E[C_m C_p]$ are given by the following empirical estimators:

$$\begin{aligned}
R_{XX}(s, s') &\approx \widehat{R}_{XX}(s, s') \stackrel{\text{def}}{=} \frac{1}{\nu} \sum_{n=1}^{\nu} X(\theta_n, s) X(\theta_n, s'), \\
E[C_m C_p] &\approx \frac{1}{\nu} \sum_{n=1}^{\nu} C_m(\theta_n) C_p(\theta_n),
\end{aligned} \tag{2.16}$$

where, for all $m \geq 1$, the ν independent realizations, $\{C_m(\theta_1), \dots, C_m(\theta_\nu)\}$, of C_m can be deduced from the ν available independent realizations of X as:

$$C_m(\theta_n) = (X(\theta_n), b_m), \quad 1 \leq n \leq \nu. \tag{2.17}$$

Given these two estimators, a direct translation of the two conditions given by Eq. (2.15) would therefore be based on the existence of $\widehat{\lambda}_m$, such that for all $m \geq 1$:

$$\frac{1}{\nu} \sum_{n=1}^{\nu} C_m(\theta_n) C_p(\theta_n) = \delta_{mp} \widehat{\lambda}_m, \tag{2.18}$$

$$\int_{\Omega} \widehat{R}_{XX}(s, s') b_m(s') ds' = \widehat{\lambda}_m b_m(s). \tag{2.19}$$

Equations (2.18) and (2.19) are however no more equivalent, and there is no reason anymore for a basis that respects one of these conditions to be still optimal with respect to error ε^2 .

- On the first hand, for any subset $\mathcal{B}^{(M)} = \{b_m, 1 \leq m \leq M\} \subset \mathcal{B}$, such that $M \leq \nu$, if we define $[C]$ as the following matrix of independent realizations:

$$[C] = \begin{bmatrix} C_1(\theta_1) & \cdots & C_1(\theta_\nu) \\ \vdots & \ddots & \vdots \\ C_M(\theta_1) & \cdots & C_M(\theta_\nu) \end{bmatrix}, \tag{2.20}$$

in which for all $1 \leq m \leq M$ and $1 \leq n \leq \nu$, $C_m(\theta_n) = (X(\theta_n), b_m)$, matrix $[R_{CC}] = \frac{1}{\nu} [C][C]^T$ is real and symmetrical, and can be rewritten as:

$$[R_{CC}] = [D][\ell][D]^T, \tag{2.21}$$

with:

$$[\ell] = \begin{bmatrix} \ell_1 & 0 & \cdots & 0 \\ 0 & \ell_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \ell_M \end{bmatrix}, \tag{2.22}$$

a diagonal matrix and $[D]$ an orthogonal matrix, such that $[D]^T[D]$ is equal to the $(M \times M)$ real unit matrix. From Eq. (2.21), it can be seen that the family $\mathcal{G}^{(M)} = \left\{ g_k = \sum_{m=1}^M [D]_{mk} b_m, 1 \leq k \leq M \right\}$ verifies the conditions given by Eq. (2.18). Indeed, for all $1 \leq k, m \leq M$, we have:

$$\begin{aligned}
\frac{1}{\nu} \sum_{n=1}^{\nu} (X(\theta_n), g_k) (X(\theta_n), g_m) &= \frac{1}{\nu} \sum_{n=1}^{\nu} \left\{ \sum_{i=1}^M [D]_{ik} (X(\theta_n), b_i) \right\} \left\{ \sum_{j=1}^M [D]_{jm} (X(\theta_n), b_j) \right\} \\
&= \sum_{i=1}^M \sum_{j=1}^M [D]_{ik} [D]_{jm} \left\{ \frac{1}{\nu} \sum_{n=1}^{\nu} C_i(\theta_n) C_j(\theta_n) \right\} \\
&= ([D]^T [R_{CC}] [D])_{km} \\
&= \delta_{km} \ell_k.
\end{aligned} \tag{2.23}$$

By construction, families $\mathcal{G}^{(M)}$ and $\mathcal{B}^{(M)}$ span the same space, such that for all $M \geq 1$:

$$\varepsilon^2(\mathcal{B}^{(M)}) = \varepsilon^2(\mathcal{G}^{(M)}). \tag{2.24}$$

Hence, whereas the uncorrelation of the projection coefficients of X on $\mathcal{K}^{(M)}$ implies the optimality of $\mathcal{K}^{(M)}$ regarding error ε^2 , there is no reason for a basis that verifies Eq. (2.18) to be optimal.

- On the other hand, let $\{\hat{k}_m, m \geq 1\}$ be the eigenfunctions of $\hat{R}_{XX}$, defined by Eq. (2.16), such that for all (s, s') in $\Omega \times \Omega$:

$$\int_{\Omega} \hat{R}_{XX}(s, s') \hat{k}_m(s') ds' = \hat{\lambda}_m \hat{k}_m(s), \quad (s, m) \in \Omega \times \mathbb{N}^*. \tag{2.25}$$

As the rank of the linear operator defined by the kernel $\hat{R}_{XX}$ is by construction lower or equal to ν , the number of elements of the basis $\{\hat{k}_m, m \geq 1\}$, for which the eigenvalues $\hat{\lambda}_m$ are non zero, is also lower or equal to ν . Hence, if eigenvalues $\{\hat{\lambda}_m, m \geq 1\}$ are sorted in a decreasing order, such that for all $m \geq 1$, $\hat{\lambda}_m \geq \hat{\lambda}_{m+1}$, the set, $\{X(\theta_1), \dots, X(\theta_{\nu})\}$, of available realizations of X is orthogonal to the subset $\{\hat{k}_m, m > \nu\}$:

$$\int_{\Omega} X(\theta_n, s) \hat{k}_m(s) ds = 0, \quad 1 \leq n \leq \nu, \quad m > \nu. \tag{2.26}$$

Therefore, the eigenfunctions of $\hat{R}_{XX}$, $\{\hat{k}_m, m \geq 1\}$, cannot be seen as an optimal basis. Only a subset of this set will be adapted to X .

Finally, whereas the M -dimension truncated KL basis, $\mathcal{K}^{(M)}$, is well characterized in theory, its numerical identification can be difficult when covariance function R_{XX} is not perfectly known. The idea of the following sections is therefore to present an innovative method to optimize the approximation of $\mathcal{K}^{(M)}$ when X is only known through a finite set of independent realizations.

2.3 Identification of optimal basis from a finite set of independent realizations

From Section 2.2, the solving of the Fredholm problem associated with any square integrable kernel function $(s, s') \mapsto A(s, s')$ on $\Omega \times \Omega$,

$$\int_{\Omega} A(s, s') b_m^A(s') ds' = \lambda_m^A b_m^A(s), \quad s \in \Omega, \quad (2.27)$$

can be seen as a generator of a particular family $\{b_m^A, m \geq 1\}$. If A is equal to the covariance function of X , the solving of this problem allows us to identify the optimal projection basis for X , that minimizes L^2 -error ε^2 . When X is only known through a finite set of independent realizations, it will first be shown in this section that, for all $M \geq 1$, the minimization of ε^2 over the M -dimension sets of functions in $\mathbb{H}$ can be replaced by an optimization problem over the kernel function on which the Fredholm problem is based. It will then be pointed out that such an optimization problem asks for a method to *a posteriori* evaluate the representativeness error associated with projection families that depends on the available realizations, that is to say when no assessment set is available. This motivates the introduction of the Leave-One-Out error, that will be presented in the second part of this section.

2.3.1 Reformulation of the projection error minimization

In Section 1.3.1, for all $M \geq 1$, the KL basis $\mathcal{K} = \{k_m, m \geq 1\}$, has been introduced as the set gathering the solutions of the Fredholm problem, defined by Eq. (1.5), associated with R_{XX} . For any function A in $\mathcal{S}(\mathbb{R})$, such that:

$$\mathcal{S}(\mathbb{R}) = \{A \in L^2(\Omega \times \Omega, \mathbb{R}), \mid A(s, s') = A(s', s) \in \mathbb{R}, \quad (s, s') \in \Omega \times \Omega\}, \quad (2.28)$$

and for any $M \geq 1$, let $\mathcal{B}(A) = \{b_m^A, m \geq 1\}$ be the set that gathers the solutions in $\mathbb{H}$ of the Fredholm problem associated with A , that is to say such that for all s in Ω and for all $m \geq 1$ and $p \geq 1$:

$$\int_{\Omega} A(s, s') b_m^A(s') ds' = \lambda_m^A b_m^A(s), \quad \lambda_1^A \geq \lambda_2^A \geq \dots \rightarrow 0, \quad (b_m^A, b_p^A) = \delta_{mp}. \quad (2.29)$$

For any $M \geq 1$ and for any function A in $\mathcal{S}(\mathbb{R})$, the set $\mathcal{B}_A^{(M)} = \{b_m^A, 1 \leq m \leq M\}$ is then introduced as the family gathering the eigenfunctions of highest eigenvalues of the Fredholm problem associated with A . The Karhunen-Loève expansion being optimal for X with respect to error ε^2 , it can be deduced that for all $M \geq 1$:

$$R_{XX} = \arg \min_{A \in \mathcal{S}(\mathbb{R})} \left\{ \varepsilon^2(\mathcal{B}_A^{(M)}) \right\}. \quad (2.30)$$

Hence, for all $M \geq 1$, the M -dimension optimal family, $\mathcal{K}^{(M)}$, which was first introduced as the solution of the optimization problem that is defined by Eq. (2.11), can equivalently be searched as the solution of the following optimization problem,

$$\mathcal{K}^{(M)} = \arg \min_{\mathcal{B}_A^{(M)}, A \in \mathcal{S}(\mathbb{R})} \left\{ \varepsilon^2(\mathcal{B}_A^{(M)}) \right\}. \quad (2.31)$$

2.3.2 Restriction of the search space

From Eqs. (2.30) and (2.31), it can directly be seen that $A = R_{XX}$ is the optimal choice for A in $\mathcal{S}(\mathbb{R})$, such that $\mathcal{K}^{(M)} = \mathcal{B}_{R_{XX}}^{(M)}$. Hence, when random field X is only characterized by a finite set of ν independent realizations, the best approximation for R_{XX} from these realizations, the most relevant for X the corresponding M -dimension family.

In this prospect, adopting the same notations than in Section 2.2.3, two classical estimators for R_{XX} , that are denoted by $\widehat{R}_{XX}$ and $\widetilde{R}_{XX}$, are introduced, such that for all (s, s') in $\Omega \times \Omega$:

$$\widehat{R}_{XX}(s, s') = \frac{1}{\nu} \sum_{n=1}^{\nu} X(\theta_n, s)X(\theta_n, s'), \quad (2.32)$$

$$\widetilde{R}_{XX}(s, s') = \begin{cases} \frac{1}{S-(s'-s)} \int_0^{S-(s'-s)} \widehat{R}_{XX}(x, x+(s'-s))dx & \text{if } S > s'-s \geq 0, \\ \frac{1}{S-(s-s')} \int_0^{S-(s-s')} \widehat{R}_{XX}(x+(s-s'), x)dx & \text{if } S > s-s' > 0, \\ \widehat{R}(s, s') & \text{otherwise.} \end{cases} \quad (2.33)$$

Recall that function $\widehat{R}_{XX}$ is the empirical estimator of R_{XX} , which converges towards R_{XX} at the convergence rate of $1/\sqrt{\nu}$. Moreover, if random field X is the restriction to Ω of a mean-square stationary random field indexed by $\mathbb{R}$, that is to say if $R_{XX}(s, s')$ only depends on the difference $|s - s'|$, function $\widetilde{R}_{XX}$ is the classical stationary estimator of R_{XX} .

From the realizations of X , statistical tests can be achieved to evaluate the relevance of a stationary hypothesis for X , in order to help us to choose the best estimator. Nevertheless, in many cases, these tests do not give a clear-cut answer. From the point of view of the minimization of Eq. (2.31), even if X is actually mean-square stationary, $\widehat{R}_{XX}$ is however considered as a better function than $\widetilde{R}_{XX}$ if:

$$\varepsilon^2(\mathcal{B}_{\widehat{R}_{XX}}^{(M)}) \leq \varepsilon^2(\mathcal{B}_{\widetilde{R}_{XX}}^{(M)}). \quad (2.34)$$

From a more general point of view, for α in $[0, 1]$, let $A(\alpha)$ be the following function:

$$A(\alpha) = \alpha \widehat{R}_{XX} + (1 - \alpha) \widetilde{R}_{XX}. \quad (2.35)$$

By construction, for any α in $[0, 1]$, functions $A(\alpha)$ are symmetrical and have the same L^2 -norm:

$$\int_{\Omega \times \Omega} A(\alpha, s, s')^2 ds ds' = \int_{\Omega \times \Omega} \widehat{R}_{XX}(s, s')^2 ds ds' = \int_{\Omega \times \Omega} \widetilde{R}_{XX}(s, s')^2 ds ds'. \quad (2.36)$$

In this work, when random field X is only characterized by a set of ν independent realizations, it is proposed to search the optimal projection basis as the solution, $\mathcal{K}^{(M)}(\alpha^*)$, of an optimization problem with respect to α in $[0, 1]$:

$$\begin{cases} \mathcal{K}^{(M)}(\alpha^*) = \mathcal{B}_{A(\alpha^*)}^{(M)} \\ \alpha^* = \arg \min_{\alpha \in [0, 1]} \left\{ \varepsilon^2(\mathcal{B}_{A(\alpha)}^{(M)}) \right\}. \end{cases} \quad (2.37)$$

Such an approach appears to be very efficient when the number of available realizations, ν , is small. Indeed, it has been shown in Section 2.2.3 that for $\nu < M$, family $\mathcal{B}_{\widehat{R}_{XX}}^{(M)}$ can be decomposed as:

$$\mathcal{B}_{\widehat{R}_{XX}}^{(M)} = \mathcal{B}^{\text{Im}}(\widehat{R}_{XX}) \cup \mathcal{B}^{\text{Ker}}(\widehat{R}_{XX}), \quad (2.38)$$

$$\begin{cases} \mathcal{B}^{\text{lm}}(\widehat{R}_{XX}) \stackrel{\text{def}}{=} \left\{ b_m, \mid \int_{\Omega} \widehat{R}_{XX}(\cdot, s') b_m(s') ds' = \widehat{\lambda}_m b_m, \widehat{\lambda}_m > 0 \right\}_{1 \leq m \leq \nu}, \\ \mathcal{B}^{\text{Ker}}(\widehat{R}_{XX}) \stackrel{\text{def}}{=} \left\{ b_m, \mid \int_{\Omega} \widehat{R}_{XX}(\cdot, s') b_m(s') ds' = 0 \right\}_{\nu < m \leq M}. \end{cases} \quad (2.39)$$

Therefore, whereas family $\mathcal{B}^{\text{lm}}(\widehat{R}_{XX})$ is likely to be particularly well adapted to X , family $\mathcal{B}^{\text{Ker}}(\widehat{R}_{XX})$ has no reason to be adapted to X as it is orthogonal to the set of available realizations. On the contrary, by construction, for $\alpha > 0$, the rank of $\mathcal{B}_{A(\alpha)}^{(M)}$ is higher than ν . Using the same notations than in Eq. (2.39), the number of elements of $\mathcal{B}^{\text{Ker}}(A(\alpha))$, which are by definition orthogonal to each available realization of X , will be smaller than $M - \nu$, such that the L^2 -error associated with $\mathcal{B}_{A(\alpha)}^{(M)}$ is likely to be smaller than the one associated with $\mathcal{B}_{\widehat{R}_{XX}}^{(M)}$.

All these considerations can directly be extended to the case when X is a $\mathbb{R}^Q$ -valued random field, $Q \geq 1$. Indeed, let $[R_{\mathbf{Y}\mathbf{Y}}]$ be the $(Q \times Q)$ matrix-valued covariance function of the $\mathbb{R}^Q$ -valued stochastic process, $\mathbf{Y} = (Y_1, \dots, Y_Q)$, for which ν independent realizations, $\{\mathbf{Y}(\theta_1), \dots, \mathbf{Y}(\theta_\nu)\}$, are available. We can thus define $[A(\alpha)]$, such that for all (s, s') in $\Omega \times \Omega$ and $1 \leq p, q \leq P$:

$$[A(\alpha)] = [\alpha][\widehat{R}] + ([I_Q] - [\alpha])[\widetilde{R}], \quad (2.40)$$

$$[\alpha]_{pq} = \alpha_p \delta_{pq}, \quad (\alpha_1, \dots, \alpha_Q) \in [0, 1]^Q, \quad (2.41)$$

$$[\widehat{R}(s, s')]_{pq} = \widehat{R}_{pq}(s, s'), \quad (2.42)$$

$$[\widetilde{R}(s, s')]_{pq} = \widetilde{R}_{pq}(s, s'), \quad (2.43)$$

$$\widehat{R}_{pq}(s, s') = \frac{1}{\nu} \sum_{n=1}^{\nu} Y_p(\theta_n, s) Y_q(\theta_n, s'), \quad (2.44)$$

$$\widetilde{R}_{pq}(s, s') = \begin{cases} \frac{1}{S - (s' - s)} \int_0^{S - (s' - s)} \widehat{R}_{pq}(x, x + (s' - s)) dx & \text{if } S > s' - s \geq 0, \\ \frac{1}{S - (s - s')} \int_0^{S - (s - s')} \widehat{R}_{pq}(x + (s' - s), x) dx & \text{if } S > s - s' > 0, \\ \widehat{R}_{pq}(s, s') & \text{otherwise.} \end{cases} \quad (2.45)$$

The optimal value for the matrix-valued function used in the Fredholm problem, $[A(\alpha^*)]$, can finally be searched as the solution of the following minimization problem:

$$\begin{cases} [A(\alpha^*)] = [\alpha^*][\widehat{R}] + ([I_Q] - [\alpha^*])[\widetilde{R}], \\ [\alpha^*] = \arg \min_{[\alpha]} \left\{ \varepsilon^2(\mathcal{B}_{[A(\alpha)]}^{(M)}) \right\}. \end{cases} \quad (2.46)$$

2.3.3 A posteriori evaluation of the representativeness error

In order to solve the problem defined by Eq. (2.37), a method to *a posteriori* evaluate $\varepsilon^2(\mathcal{B}_{A(\alpha)}^{(M)})$ from the only ν available independent realizations, $\{X(\theta_n), 1 \leq n \leq \nu\}$, for all α in $[0, 1]$, is then required. Indeed, from the limited set $\{X(\theta_n), 1 \leq n \leq \nu\}$, the L^2 -error, $\varepsilon^2(\mathcal{B}^{(M)})$, corresponding to any M -dimension family $\mathcal{B}^{(M)}$ of $\mathbb{H}^M$, cannot be exactly calculated, but has to be evaluated as precisely as possible. Two cases can be distinguished:

- case 1: $\mathcal{B}^{(M)}$ is defined without any reference to $\{X(\theta_1), \dots, X(\theta_\nu)\}$.
- case 2: the knowledge of $\{X(\theta_1), \dots, X(\theta_\nu)\}$ is used to optimize the representativeness of $\mathcal{B}^{(M)}$. In this case, $\mathcal{B}^{(M)}$ depends on the available realizations of X .

Case 1: realizations and projection basis are independent

If $\mathcal{B}^{(M)}$ has been computed without any reference to the set $\{X(\theta_1), \dots, X(\theta_\nu)\}$, error $\varepsilon^2(\mathcal{B}^{(M)})$ can be evaluated from its empirical estimation, $\hat{\varepsilon}_\nu^2(\mathcal{B}^{(M)})$, such that:

$$\hat{\varepsilon}_\nu^2(\mathcal{B}^{(M)}) = \frac{1}{\nu} \sum_{n=1}^{\nu} \left(X(\theta_n) - \hat{X}^{\mathcal{B}^{(M)}}(\theta_n), X(\theta_n) - \hat{X}^{\mathcal{B}^{(M)}}(\theta_n) \right). \quad (2.47)$$

Indeed, according to the central limit theorem (see [15] for further details),

$$\lim_{\nu \rightarrow +\infty} P \left(\left| \varepsilon^2(\mathcal{B}^{(M)}) - \hat{\varepsilon}_\nu^2(\mathcal{B}^{(M)}) \right| \leq z(p) \sqrt{\frac{\text{Var} \left\{ \left(X - \hat{X}^{\mathcal{B}^{(M)}}, X - \hat{X}^{\mathcal{B}^{(M)}} \right) \right\}}{\nu}} \right) = 1 - p, \quad (2.48)$$

$$1 - \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{z(p)} \exp \left(-\frac{x^2}{2} \right) dx = \frac{p}{2}, \quad (2.49)$$

such that for sufficiently high values of ν :

$$\varepsilon^2(\mathcal{B}^{(M)}) \approx \hat{\varepsilon}_\nu^2(\mathcal{B}^{(M)}). \quad (2.50)$$

Case 2: the projection basis depends on the available realizations

In order to make projection family $\mathcal{B}^{(M)}$ be particularly adapted to random field X , it can be interesting to exploit as much as possible the information about X that is gathered in independent realizations $\{X(\theta_1), \dots, X(\theta_\nu)\}$. In this case, $\mathcal{B}^{(M)}$ is dependent on $\{X(\theta_1), \dots, X(\theta_\nu)\}$, and error $\hat{\varepsilon}_\nu^2$ strongly underestimates ε^2 . This phenomenon is generally called *overlearning*. For instance, if we define $\mathcal{B}^{(M)} = \{b_m, 1 \leq m \leq M\}$ as the Gram-Schmidt orthogonalization of the deterministic family of available independent realizations $\{X(\theta_1), \dots, X(\theta_\nu)\}$:

$$\left[\begin{array}{l} b_1 = X(\theta_1) / (X(\theta_1), X(\theta_1)), \quad K = 1, \\ \text{for } 2 \leq m \leq M : \\ \quad b_m^* = X(\theta_m) - \sum_{k=1}^{m-1} (X(\theta_m), b_k) b_k \\ \quad \text{if } (b_m^*, b_m^*) > 0 : \\ \quad \quad K = K + 1, \quad b_K = b_m^* / \sqrt{(b_m^*, b_m^*)} \\ \quad \text{end if} \\ \text{end for} \\ M = K, \end{array} \right. \quad (2.51)$$

then, the M -dimension projection,

$$\widehat{X}^{\mathcal{B}^{(M)}} = \sum_{m=1}^M C_m b_m, \quad C_m = (X, b_m), \quad (2.52)$$

of random field X on $\mathcal{B}^{(M)}$ verifies:

$$\widehat{X}^{\mathcal{B}^{(M)}}(\theta_n) = X(\theta_n), \quad 1 \leq n \leq \nu. \quad (2.53)$$

By construction, error $\widehat{\varepsilon}_\nu^2(\mathcal{B}^{(M)})$ is always equal to zero, whereas $\varepsilon^2(\mathcal{B}^{(M)})$ should be in general strictly greater than 0, as the number of available realization, ν , and the dimension of the projection basis, M , are limited.

In order to correctly evaluate error ε^2 , a separation in two sets of the available realizations is generally performed:

- the first set, $\{X(\theta_1), \dots, X(\theta_{\nu^*})\}$, is a learning set, on which the definition of $\mathcal{B}^{(M)}$ is based,
- the second set, $\{X(\theta_{\nu^*+1}), \dots, X(\theta_\nu)\}$, is an assessment set, on which the computation of $\widehat{\varepsilon}_{\nu-\nu^*}^2$ is achieved to evaluate ε^2 .

With such a method, it can be noticed that the higher ν^* , the less precise the evaluation of ε^2 . This limits strongly the scope of such approaches when number of available realizations ν is small, compared to the number of functions that are needed to characterize X . Indeed, in such cases, we would be interested in taking into account most of the available realizations of X , that is to say in making ν^* tends to ν , which leads us nevertheless to a very bad evaluation of ε^2 .

To this end, for all set of ν^* indices, $\mathbb{J}(\nu^*)$, such that:

$$\mathbb{J}(\nu^*) = \{j_1 \neq \dots \neq j_{\nu^*}\} \in \{1, \dots, \nu\}^{\nu^*}, \quad (2.54)$$

let $\mathcal{B}^{(M)}(\mathbb{J}(\nu^*))$ be the M -dimension family that has been computed from the ν^* -dimension set $\{X(\theta_{j_1}), \dots, X(\theta_{j_{\nu^*}})\}$ (ν^* can vary) of independent realizations of X (the family that stems from the Gram-Schmidt orthogonalization defined by Eq. (2.51) is an example of such a family). The two following hypotheses are then assumed.

1. First, it is supposed that error $\varepsilon^2(\mathcal{B}^{(M)}(\mathbb{J}(\nu^*)))$ decreases when ν^* increases.
2. Then, given two sets $\mathbb{J}^{(1)}(\nu^*)$ and $\mathbb{J}^{(2)}(\nu^*)$ that have been randomly chosen in $\{1, \dots, \nu\}^{\nu^*}$, it is assumed that:

$$P_{\widehat{X}^{(1)}} \approx P_{\widehat{X}^{(2)}}, \quad (2.55)$$

where $P_{\widehat{X}^{(1)}}$ and $P_{\widehat{X}^{(2)}}$ are the distributions of the projected random fields, $\widehat{X}^{(1)}$ and $\widehat{X}^{(2)}$, on $\mathcal{B}^{(M)}(\mathbb{J}^{(1)}(\nu^*))$ and $\mathcal{B}^{(M)}(\mathbb{J}^{(2)}(\nu^*))$ respectively.

In other words, the first hypothesis means that the modeling errors are expected to decrease when the information is increasing, whereas the second hypothesis asks, for the application, that computing the projection basis from a limited set of realizations yields robust results.

Let $\mathcal{J}(\nu - 1)$ be the random variable, whose distribution is discrete, such that for all $1 \leq n \leq \nu$, $\mathcal{J}(\nu - 1)$ takes set value $\mathcal{J}(\nu - 1, \theta_n) = \{1, \dots, n - 1, n + 1, \dots, \nu\}$ with probability $1/\nu$. For all $1 \leq n \leq \nu$, $\mathcal{J}(\nu - 1, \theta_n)$ corresponds thus to a particular set of $\nu - 1$ indices. With such a formalism, the projection basis that is only based on the knowledge of the $\nu - 1$ realizations, $\{X(\theta_1), \dots, X(\theta_{n-1}), X(\theta_{n+1}), \dots, X(\theta_\nu)\}$, of X , which is denoted by $\mathcal{B}^{(M)}(\mathcal{J}(\nu - 1, \theta_n))$, is independent of $X(\theta_n)$. Therefore, under the two former hypotheses, if $\hat{X}(\mathcal{J}(\nu - 1, \theta_n), \theta_n)$ is the projection of $X(\theta_n)$ on $\mathcal{B}^{(M)}(\mathcal{J}(\nu - 1, \theta_n))$, the set $\{e^2(\theta_n), 1 \leq n \leq \nu\}$, where:

$$e^2(\theta_n) = \left(X(\theta_n) - \hat{X}(\mathcal{J}(\nu - 1, \theta_n), \theta_n), X(\theta_n) - \hat{X}(\mathcal{J}(\nu - 1, \theta_n), \theta_n) \right), \quad (2.56)$$

can be seen as a set of ν independent realizations of the random variable:

$$e^2 = \left(X - \hat{X}(\mathcal{J}(\nu - 1)), X - \hat{X}(\mathcal{J}(\nu - 1)) \right), \quad (2.57)$$

such that:

$$\varepsilon^2 \left(\mathcal{B}^{(M)}(\mathbb{J}(\nu - 1)) \right) \approx \varepsilon_{LOO}^2(\mathcal{B}^{(M)}) \stackrel{\text{def}}{=} \frac{1}{\nu} \sum_{i=1}^{\nu} e^2(\theta_i). \quad (2.58)$$

According to the central limit theorem, this error converges to $\varepsilon^2(\mathcal{B}^{(M)}(\mathbb{J}(\nu - 1)))$ at the convergence rate of $1/\sqrt{\nu}$. For all projection family $\mathcal{B}^{(M)}$, the estimation $\varepsilon_{LOO}^2(\mathcal{B}^{(M)})$ is called Leave-One-Out (LOO) error, and it can be seen as a good approximation of $\varepsilon^2(\mathcal{B}^{(M)})$ for ν sufficiently high. This LOO error can be considered as the application of the jackknife theory (see [73, 74, 75] for further details) to the evaluation of projection errors. Hence, contrary to the two-sets approach, the Leave-One-Out method allows us to compute projection basis $\mathcal{B}^{(M)}$ from all the available realizations of X , while still giving access to an accurate estimation of its corresponding representativeness error, $\varepsilon^2(\mathcal{B}^{(M)})$, when ν is sufficiently high.

Finally, the L^2 -error, ε^2 , in the optimization problem, defined by Eq. (2.37), can be replaced by the LOO error, such that the M -dimension optimal projection family, $\mathcal{K}^{(M)}$, can be approximated by the following optimization problem:

$$\begin{cases} \mathcal{K}^{(M)} \approx \mathcal{B}_{A(\alpha^*)}^{(M)}, \\ \alpha^* = \arg \min_{\alpha \in [0,1]} \left\{ \varepsilon_{LOO}^2(\mathcal{B}_{A(\alpha)}^{(M)}) \right\}. \end{cases} \quad (2.59)$$

2.4 Applications

In order to illustrate the benefits that stem from the optimization problem defined by Eq. (2.59), an application based on simulated data is presented in this section. This application aims at justifying the relevance of the Leave-One-Out error, and at emphasizing the difficulties of the classical Karhunen-Loève expansion-based methods to identify optimal basis when the number of available realizations is low, while the generalized Karhunen-Loève expansion, characterized by Eq. (2.59) gives very promising results.

2.4.1 Generation of independent realizations of the random field

Let $\Omega = [0, 1]$, and X be a random field of $\mathcal{P}(\Omega)$, for which the covariance function R_{XX} is represented in Figure 2.1 and the mean value is zero. This random field has been chosen on purpose **non stationary**.

From this choice for R_{XX} , we can numerically identify the Karhunen-Loève basis, $\mathcal{K} = \{k_m, m \geq 1\}$, associated with X by solving the Fredholm problem associated with R_{XX} , defined by Eq. (1.5). For any value of $M \geq 1$, it is reminded that $\mathcal{K}^{(M)} = \{k_m, 1 \leq m \leq M\}$ is optimal in the sense that it minimizes ε^2 :

$$\varepsilon^2(\mathcal{K}^{(M)}) = \min_{\mathcal{B}^{(M)} \in \mathbb{H}^M} \varepsilon^2(\mathcal{B}^{(M)}). \quad (2.60)$$

From a numerical point of view, the Fredholm equation is solved using a Galerkin-type approximation, as presented in Section 1.3.3. Interval Ω is discretized with the spatial step $h = 0.005$, and random field X is approximated by its N -dimension projection, $X^{(N)}$, with $N = 1/h + 1 = 201$, such that:

$$X(s) \approx X^{(N)}(s) = \sum_{i=1}^N h_i(s) X((i-1)h), \quad s \in \Omega, \quad (2.61)$$

$$h_i(s) = \begin{cases} (ih - s)/h, & \text{if } (i-1)h \leq s \leq ih, \\ (s - (i-2)h)/h, & \text{if } (i-2)h \leq s \leq (i-1)h, \\ 0 & \text{otherwise.} \end{cases} \quad (2.62)$$

The covariance function R_{XX} , of centered random field X , is also approximated by its Galerkin projection, R_{XX}^h , such that, for all (s, s') in $\Omega \times \Omega$:

$$R_{XX}(s, s') \approx R_{XX}^h(s, s') = \sum_{i=1}^N \sum_{j=1}^N h_i(s) h_j(s') E[X((i-1)h) X((j-1)h)]. \quad (2.63)$$

The eigenvalue problem, defined by Eq. (1.5), associated with kernel R_{XX}^h , leads us to the definition of N functions, that are denoted by $\{k_m^h, 1 \leq m \leq N\}$. Let $\{\lambda_m^h, 1 \leq m \leq N\}$ be the corresponding eigenvalues, such that:

$$R_{XX}^h(s, s') = \sum_{m=1}^N \lambda_m^h k_m^h(s) k_m^h(s'), \quad (s, s') \in \Omega \times \Omega, \quad (2.64)$$

$$X^{(N)}(s) = \sum_{m=1}^N \sqrt{\lambda_m^h} k_m^h(s) \xi_m, \quad s \in \Omega, \quad (2.65)$$

with $\boldsymbol{\xi} = (\xi_1, \dots, \xi_N)$ a N -dimension random vector of uncorrelated random variables. In this application, X is supposed to be a Gaussian random field (the Gaussian hypothesis is just introduced in order to simplify the generation of independent realizations of X but the following conclusions would be exactly the same for a non-Gaussian case). Consequently, the components of $\boldsymbol{\xi}$ are independent normalized Gaussian random variables.

Two sets, $\mathcal{X}^{\text{exp}} = \{X(\theta_1), \dots, X(\theta_\nu)\}$ and $\mathcal{X}^{\text{valid}} = \{X(\Theta_1), \dots, X(\Theta_{\nu^{\text{valid}}})\}$, of independent realizations of X are then generated from the KL decomposition defined by Eq. (2.65). Four particular independent realizations of X are represented in Figure 2.2. Set $\mathcal{X}^{\text{exp}}$ represents the available information about X , whereas $\mathcal{X}^{\text{valid}}$ is the assessment set, which will only be used according to Section 2.3.3 to evaluate the projection error, $\varepsilon^2(\mathcal{B}^{(M)})$, corresponding to any projection family $\mathcal{B}^{(M)}$ in $\mathbb{H}^M$.

In the following, ν^{valid} is chosen equal to 4,000 for the convergence of $\varepsilon_{\nu^{\text{valid}}}^2(\mathcal{B}^{(M)})$, defined by Eq. (2.47), towards $\varepsilon^2(\mathcal{B}^{(M)})$ to be achieved.

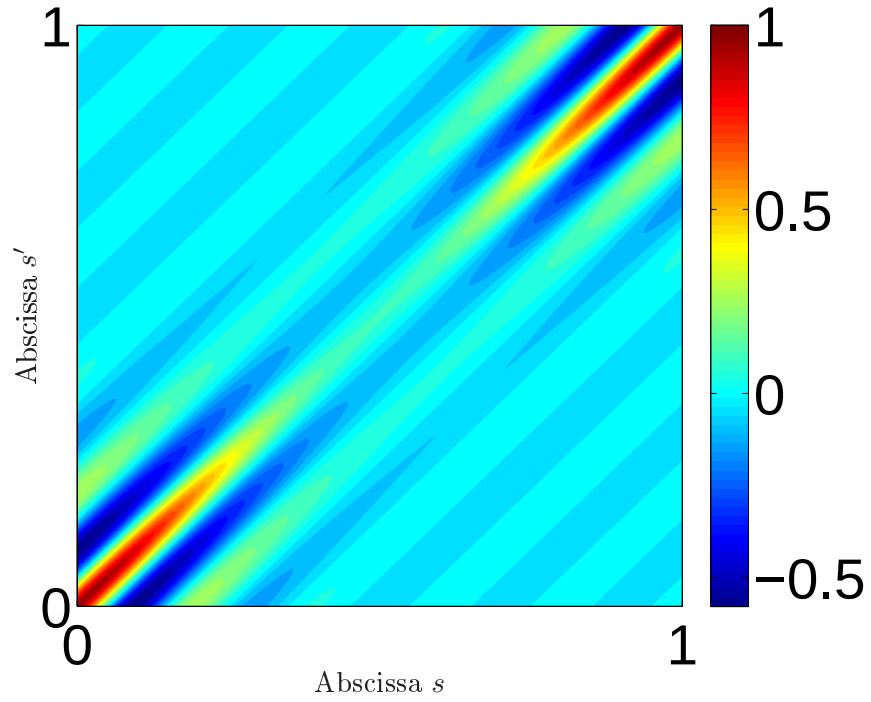


Figure 2.1: Representation of the covariance function $(s, s') \mapsto R_{XX}(s, s')$.

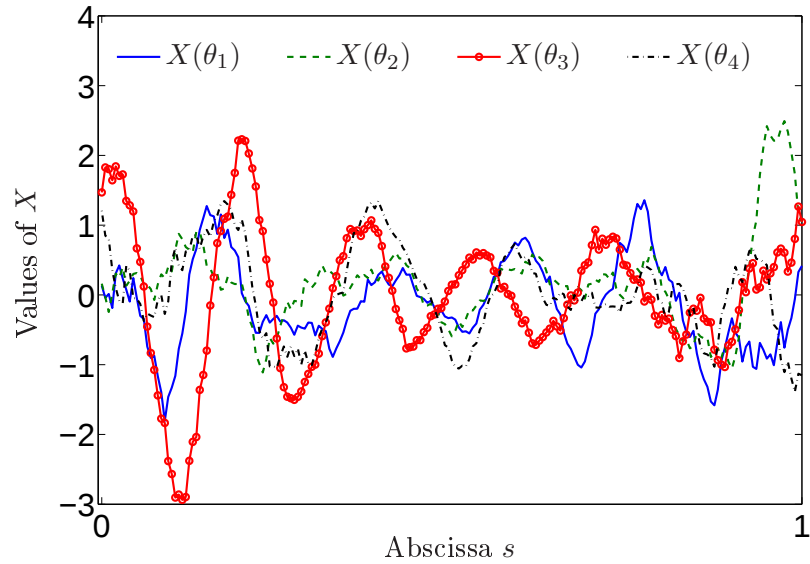


Figure 2.2: Representation of four independent realizations of X .

2.4.2 Improvement of the projection basis with respect to the available information

The number of available realizations, ν , is now supposed to be in the set $\{0, 10, 20, 50, 200\}$. The case $\nu = 0$ corresponds to a limit case when no realization of X is available. For the other cases, the empirical estimator of covariance function R_{XX} , which is denoted by $\widehat{R}(\nu)$, such that:

$$\widehat{R}(\nu, s, s') = \frac{1}{\nu} \sum_{n=1}^{\nu} X(\theta_n, s) X(\theta_n, s'), \quad (s, s') \in \Omega \times \Omega, \quad (2.66)$$

is compared in Figure 2.3 for different values of ν . In these figures, it can be verified that the higher ν , the more relevant $\widehat{R}(\nu)$. Using the same notations than in Section 2.3.2, for all $1 \leq M \leq N$, $\mathcal{B}_{\widehat{R}(\nu)}^{(M)} = \{b_1^\nu, \dots, b_M^\nu\}$ is introduced as the M -dimension family such that:

$$\int_{\Omega} \widehat{R}(\nu)(\cdot, s') b_m^\nu(s') ds' = \lambda_m^\nu b_m^\nu, \quad 1 \leq m \leq M, \quad (2.67)$$

$$\lambda_1^\nu \geq \lambda_2^\nu \geq \dots \geq \lambda_M^\nu \geq 0. \quad (2.68)$$

By construction, for $\nu < M$, the rank of $\widehat{R}(\nu)$ is equal to ν , and then,

$$\lambda_{\nu+1}^\nu = \lambda_{\nu+2}^\nu = \dots = \lambda_M^\nu = 0. \quad (2.69)$$

Therefore, as described in Section 2.3.1, the elements of $\{b_{\nu+1}^\nu, \dots, b_M^\nu\}$ are orthogonal to the available realizations of X : their characterization does not take into account any information about X . In particular, in the case $\nu = 0$, $\mathcal{B}_{\widehat{R}(\nu=0)}^{(M)}$ corresponds to any M -dimension set of orthonormal functions of $\mathbb{H}$. In Figure 2.3 are thus compared the evolutions of the error functions, $\varepsilon^2(\mathcal{B}_{\widehat{R}(\nu)}^{(M)})$, with respect to M , for four considered values of ν . First, in this figure, it can be noticed that $\varepsilon^2(\mathcal{B}_{\widehat{R}(\nu=0)}^{(M)})$ decreases linearly with respect to M , which means that the relevance of each element of $\mathcal{B}_{\widehat{R}(\nu=0)}^{(M)}$ to describe X is approximatively the same. This is a direct and natural consequence of the fact that all these elements have been defined without information about X . Then, two phases can clearly be identified in the evolution of $\varepsilon^2(\mathcal{B}_{\widehat{R}(\nu)}^{(M)})$ with respect to M , for $\nu = 10, 20, 50, 200$: the decrease of $\varepsilon^2(\mathcal{B}_{\widehat{R}(\nu)}^{(M)})$ is indeed much faster for $M \leq \nu$ than for $M > \nu$, where a quasi-linear decrease is found again. This behavior can be justified by the fact that the ν first elements of $\mathcal{B}_{\widehat{R}(\nu)}^{(M)}$ are based on the available realizations of X , whereas the $M - \nu$ last elements are not.

2.4.3 Optimized basis when few realizations are available

In the former section, it has been shown that the basis that stems from the solving of Eq. (2.67) appears to be relevant to characterize X , especially when the number of available realizations, ν , is high. This section aims at illustrating the benefits of the approach introduced in Section 2.3.1, in cases when $\nu \ll N$. For ν in $\{10, 20, 50, 200\}$ and $1 \leq M \leq N$, let $\alpha^*(\nu)$ be the solution of the optimization problem, defined by Eq. (2.37) (in this application, Eq. (2.37) has been solved using an algorithm based on a dichotomy), and let $A(\alpha^*(\nu))$ be the corresponding function in $\mathcal{S}(\mathbb{R})$. The relevance of $\mathcal{B}_{\widehat{R}(\nu)}^{(M)}$ and $\mathcal{B}_{A(\alpha^*(\nu))}^{(M)}$ is then compared in Figure 2.5. The optimal value of the projection error, $\varepsilon^2(\mathcal{K}^{(M)})$, has been added in these figures as a limit state. In each case, it can thus be seen that $\varepsilon^2(\mathcal{B}_{A(\alpha^*(\nu))}^{(M)}) \leq \varepsilon^2(\mathcal{B}_{\widehat{R}(\nu)}^{(M)})$. For these four choices for ν ,

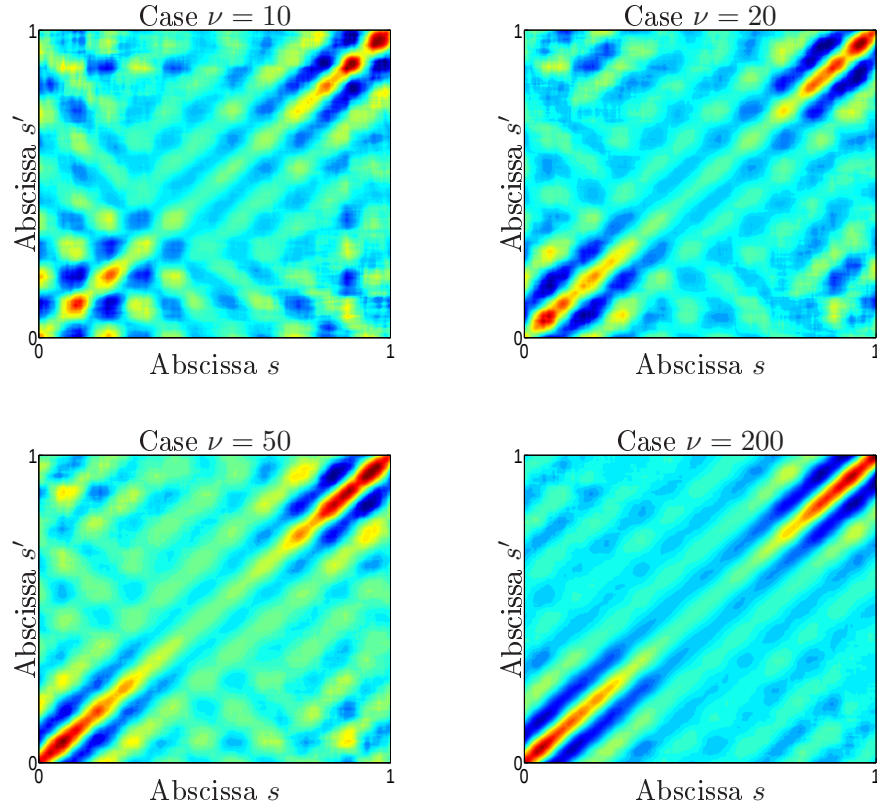


Figure 2.3: Empirical estimators $\hat{R}(\nu)$ for four values of ν .

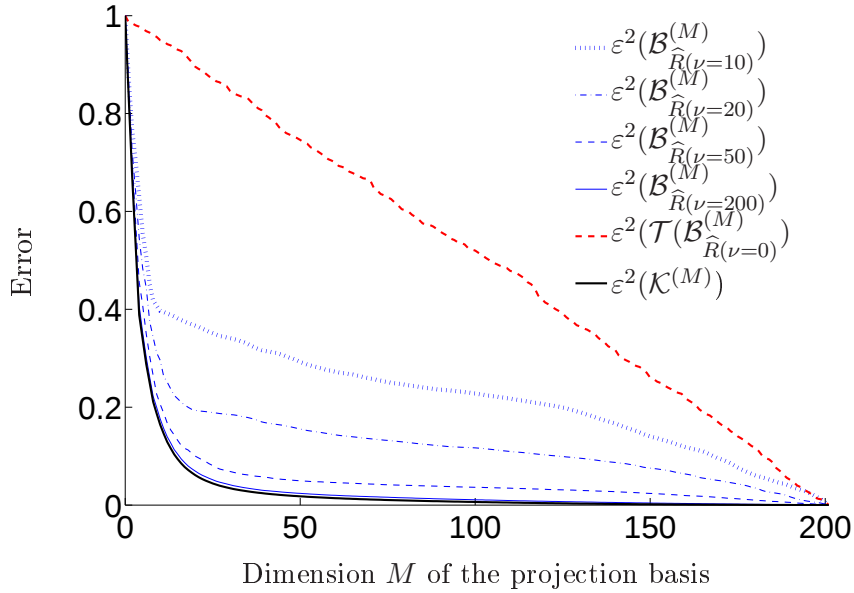


Figure 2.4: Improvement of the projection basis with respect to the number of available realizations, ν .

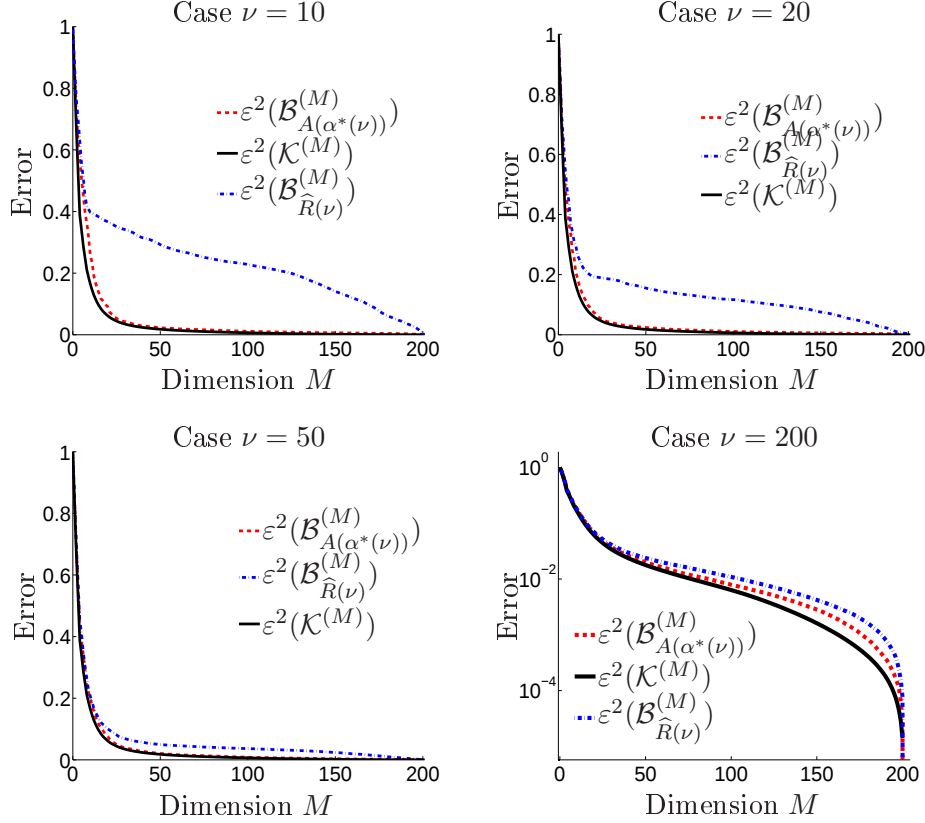


Figure 2.5: Improvement of the projection basis.

this figure underlines the great benefits of the formulation defined by Eq. (2.37), especially when $\nu \ll N$. For instance, for $\nu = 10 \ll N = 201$ and $M = 50$, whereas $\varepsilon^2(\mathcal{B}_{\hat{R}(\nu)}^{(50)}) = 29.2\%$, $\varepsilon^2(\mathcal{B}_{A(\alpha^*(\nu))}^{(50)}) = 2.27\%$. Indeed, whereas the rank of $\hat{R}(\nu = 10)$ is 10, the rank of $A(\alpha^*(\nu = 10))$ is by construction much higher than ν . Therefore, more elements of $\mathcal{B}_{A(\alpha^*(\nu))}^{(50)}$ are based on the knowledge of X than the elements of $\mathcal{B}_{\hat{R}(\nu)}^{(50)}$, which explains such an improvement of the projection basis, even if X is non stationary.

2.4.4 Relevance of the LOO error

As presented in Section 2.3.3, when the assessment set, $\mathcal{X}^{\text{valid}}$, is not available, which is the general case, the LOO error allows us to evaluate the projection error from the only set $\mathcal{X}^{\text{exp}}$. For $M = 50$, the relevance of the LOO error is illustrated in Figure 2.6. It can be seen in this figure that for even low values of ν , LOO error $\varepsilon_{LOO}^2(\mathcal{B}_{\hat{R}(\nu)}^{(M)})$, which is only based on the ν available realizations of X , is very close to the validation error $\hat{\varepsilon}_{\nu^{\text{valid}}}^2(\mathcal{B}_{\hat{R}(\nu)}^{(M)})$, defined by Eq. (2.47), which is based on the ν^{valid} realizations gathered in $\mathcal{X}^{\text{valid}}$. Error bars have been added in these two graphs for several values of ν . These bars correspond to the 95% confidence intervals and emphasize the convergence in $1/\sqrt{\nu}$ of the LOO error towards $\hat{\varepsilon}_{\nu^{\text{valid}}}^2(\mathcal{B}_{\hat{R}(\nu)}^{(M)})$.

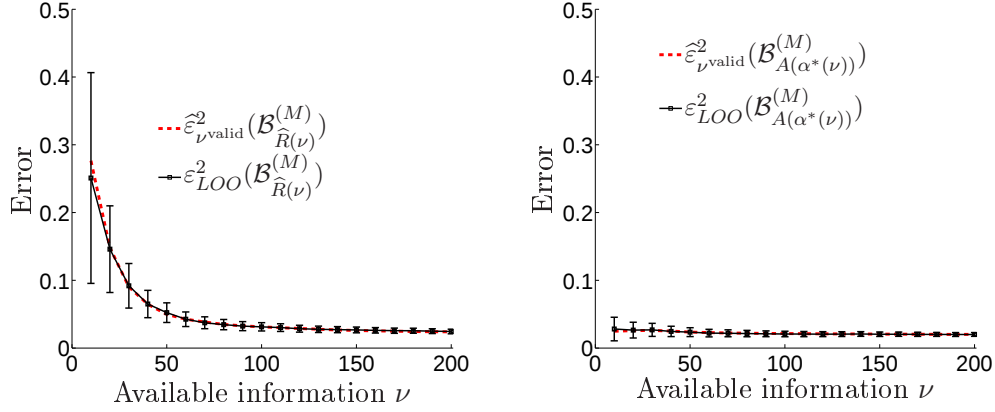


Figure 2.6: Relevance of the LOO estimator.

2.5 Conclusions

For the last decades, the increasing computational power has encouraged many scientific fields to take into account random field in their modeling. The development of reduced basis that could condense at best the statistical properties of these random fields is therefore of great interest. In most of these applications, the knowledge of these random fields is however limited to a finite set of independent realizations. In this context, this chapter emphasized the efficiency of a method based on an adaptation of the Karhunen-Loève expansion, in order to construct optimized basis from a relatively small set of independent realizations. First, this method defined an original optimization problem, and secondly, required a way to *a posteriori* evaluate projection errors. Finally, when interested in studying complex systems that are excited by random fields that are only known through a set of limited independent realizations, the method proposed opens new opportunities to optimize the projection basis with respect to the available information.

Chapter 3

PCE identification in high dimension from a set of realizations

3.1 Introduction

In Section 1.4, it has been shown to what extent the Polynomial Chaos Expansion (PCE) method allows us to identify in inverse the multidimensional distribution of the M -dimension random vector $\mathbf{C} = (C_1, \dots, C_M)$ when the available information about this random vector is a finite set of ν independent realizations, which are written $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$. Without loss of generality, $\mathbf{C}$ is supposed to be a centered random vector in this chapter.

As presented in Section 1.5, this method is based on a direct projection of $\mathbf{C}$ on a known and chosen N -dimension orthonormal basis $\{\psi_1(\boldsymbol{\xi}), \dots, \psi_N(\boldsymbol{\xi})\}$, such that:

$$\mathbf{C} \approx \mathbf{C}^{\text{chaos}}(N) = [\mathbf{y}]\boldsymbol{\Psi}(\boldsymbol{\xi}), \quad \boldsymbol{\Psi}(\boldsymbol{\xi}) = (\psi_1(\boldsymbol{\xi}), \dots, \psi_N(\boldsymbol{\xi})), \quad (3.1)$$

with $\boldsymbol{\xi}$ a N_g -dimension random vector ($N_g \leq M$) whose distribution is known.

To identify the multidimensional distribution of $\mathbf{C}$, the $(M \times N)$ projection matrix $[\mathbf{y}]$ has then to be calculated from the available information about $\mathbf{C}$, and the values of the truncation parameters N_g and N have to be justified according to convergence analysis. To this end, an error function has been defined in Section 1.5.4 to quantify the amplitude of the PCE residue, $\mathbf{C} - \mathbf{C}^{\text{chaos}}(N)$, whereas a random search algorithm has been introduced in Section 1.5.3 to allow the computation of $[\mathbf{y}]$ from the realizations $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$ of $\mathbf{C}$.

Dealing with high dimensional problems, that is to say when M and N are very high, raises however at least two major difficulties.

- First, when M and N are high, the dimension of the admissible set $\mathcal{O}_C$ becomes huge, such that the convergence of the random search algorithms that are based on independent and uniformly distributed generations of $[\mathbf{y}^*]$ in $\mathcal{O}_C$ to solve the optimization problem defined by Eq. (1.61), is very low. A method to optimize the generation of elements in $\mathcal{O}_C$ is therefore needed for such PCE inverse identification to give relevant results.
- Secondly, the optimization problem defined by Eq. (1.61) is based on the generation of the matrix $[\Psi(\nu^{\text{chaos}})]$ of independent realizations of projection vector $\boldsymbol{\Psi}(\boldsymbol{\xi})$. Recurrence formula or algebraic explicit representations are generally used to compute such matrix $[\Psi(\nu^{\text{chaos}})]$, which are supposed to verify the asymptotic property:

$$\lim_{\nu^{\text{chaos}} \rightarrow +\infty} \frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T = [I_N], \quad (3.2)$$

as a direct consequence of the orthonormality of the PCE basis. However, for numerically admissible values of ν^{chaos} (between 1000 and 10000), it has been shown in [76] that the difference $\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N]$ can be very significant when N is high. This difference induces a detrimental bias in the PCE identification, which makes the convergence of the classical PCE in high dimension very difficult. Innovative methods to generate matrices $[\Psi(\nu^{\text{chaos}})]$ that numerically verify $\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T = [I_N]$ are thus expected to allow this convergence in high dimension.

Solutions to these two difficulties are therefore proposed in this chapter. First, Section 3.2 presents an original method to optimize the trials in the admissible set

$$\mathcal{O}_C = \left\{ [y] = [\mathbf{y}^{(1)}, \dots, \mathbf{y}^{(N)}] \in \mathcal{M}_{M,N} \mid [y][y]^T = [\hat{R}_{CC}(\nu)] \right\}, \quad (3.3)$$

$$[\hat{R}_{CC}(\nu)] = \frac{1}{\nu} \sum_{n=1}^{\nu} \mathbf{C}(\theta_n) \otimes \mathbf{C}(\theta_n), \quad (3.4)$$

even if ν is very small compared to M and N , such that relevant values for $[y]$ can be computed at a reasonable computational cost.

Then, Section 3.3 addresses the stabilization of the matrix of realizations $[\Psi(\nu^{\text{chaos}})]$ in high dimension. At last, two applications are presented in Section 3.4 to emphasize the benefits of such improvements in the PCE identification process.

3.2 Optimized trials of independent realizations of random matrices under correlation constraints

In this chapter, we use the same notations than in Section 1.5.

To solve Eq. (1.61) with a random search algorithm, as an extension of the work described in [38], this section aims at proposing two methods to optimize the trials in $\mathcal{O}_C$. Adaptations of these methods are then presented when $\nu < M$, that is to say when the available information is very limited compared to the size of $\mathbf{C}$.

3.2.1 Reformulation of the correlation constraints

From Eqs. (1.33) and (1.50), for a given value of ν , it is recall that the matrices $[C^{\text{exp}}(\nu)]$ and $[\hat{R}_{CC}(\nu)]$ are defined such that:

$$[C^{\text{exp}}(\nu)] = [\mathbf{C}(\theta_1) \dots \mathbf{C}(\theta_\nu)], \quad [\hat{R}_{CC}(\nu)] = \frac{1}{\nu} [C^{\text{exp}}(\nu)] [C^{\text{exp}}(\nu)]^T. \quad (3.5)$$

For any $(M \times M)$ invertible real matrix $[G]$, the random vector $\tilde{\mathbf{C}}$ is introduced as:

$$\tilde{\mathbf{C}} = [G]\mathbf{C}. \quad (3.6)$$

As $[G]$ is invertible, $\tilde{\mathbf{C}}$ and $\mathbf{C}$ belong to the same statistical space. Therefore, searching the PCE $[y]\Psi(\boldsymbol{\xi})$ describing $\mathbf{C}$ under the constraint $[y][y]^T = [\hat{R}_{CC}(\nu)]$ is equivalent to searching the PCE $[u]\Psi(\boldsymbol{\xi})$ describing $\tilde{\mathbf{C}}$, where matrix $[u] = [G][y]$ has to verify:

$$[u][u]^T = [G][\widehat{R}_{CC}(\nu)][G]^T. \quad (3.7)$$

Matrix $[\widehat{R}_{CC}(\nu)]$ being a symmetrical and real matrix, it exists an orthogonal matrix $[\widehat{V}(\nu)]$ and a diagonal matrix $[\widehat{\lambda}(\nu)]$ in $\mathcal{M}_{M,M}$ such that:

$$[\widehat{R}_{CC}(\nu)] = \left([\widehat{V}(\nu)][\widehat{\lambda}(\nu)]^{1/2}\right) \left([\widehat{V}(\nu)][\widehat{\lambda}(\nu)]^{1/2}\right)^T, \quad [\widehat{V}(\nu)]^T [\widehat{V}(\nu)] = [I_M]. \quad (3.8)$$

$$[\widehat{\lambda}(\nu)]^{1/2} = \begin{bmatrix} \sqrt{\widehat{\lambda}_1(\nu)} & 0 & \dots & 0 \\ 0 & \sqrt{\widehat{\lambda}_2(\nu)} & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \sqrt{\widehat{\lambda}_M(\nu)} \end{bmatrix}, \quad \widehat{\lambda}_1(\nu) \geq \dots \geq \widehat{\lambda}_M(\nu) \geq 0. \quad (3.9)$$

The idea is thus to find a particular matrix $[G]$ that could simplify the orthogonality constraints for $[y]$. If $\widehat{\lambda}_M(\nu) > 0$, $[\widehat{V}(\nu)][\widehat{\lambda}(\nu)]^{1/2}$ is invertible. The particular choice $[G] = [\widehat{\lambda}(\nu)]^{-1/2}[\widehat{V}(\nu)]^T$ imposes therefore on $[u]$ to belong to the Stiefel manifold $V_{N,M}$ (see [77] for further details about the Stiefel manifold), such that:

$$V_{N,M} = \{[u] \in \mathcal{M}_{M,N} \mid [u][u]^T = [I_M]\}. \quad (3.10)$$

3.2.2 Notations and definitions

In this section, a series of notations are introduced, on which the next sections will be based. For $1 \leq z \leq Z$:

- $\mathcal{J}_z$ is a random permutation from $\{1, 2, \dots, M\}$ to $\{1, 2, \dots, M\}$, such that:

$$\mathcal{J}_z = (j_1^{(z)}, \dots, j_M^{(z)}) \in \{1, 2, \dots, M\}^M, \quad j_1^{(z)} \neq \dots \neq j_M^{(z)}; \quad (3.11)$$

- the set $\{\mathbf{v}_{(z),m}^T, 1 \leq m \leq M\}$ gathers the M rows of matrix $[y^{(z)}]$, such that:

$$[y^{(z)}] = \begin{bmatrix} \mathbf{v}_{(z),1}^T \\ \vdots \\ \mathbf{v}_{(z),M}^T \end{bmatrix}; \quad (3.12)$$

- For $1 \leq m \leq M$, PDF $\widehat{p}_{(U_{j_1}^{(z)}, \dots, U_{j_m}^{(z)})}$ refer to the kernel estimators of the multidimensional PDF $p_{(U_{j_1}^{(z)}, \dots, U_{j_m}^{(z)})}$ of the random vector

$$(U_{j_1}^{(z)}, \dots, U_{j_m}^{(z)}) = \begin{bmatrix} \mathbf{v}_{(z),j_1}^T \\ \vdots \\ \mathbf{v}_{(z),j_m}^T \end{bmatrix} \Psi(\boldsymbol{\xi}). \quad (3.13)$$

- In the same manner, the set $\{\widetilde{\mathbf{v}}_{(z),m}^T, 1 \leq m \leq M\}$ gathers the M rows of matrix $[u^{(z)}]$, such that:

$$[u^{(z)}] = \begin{bmatrix} \tilde{\mathbf{v}}_{(z),1}^T \\ \vdots \\ \tilde{\mathbf{v}}_{(z),M}^T \end{bmatrix}; \quad (3.14)$$

- For $1 \leq m \leq M$, PDF $\hat{p}_{(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})}$ refer to the kernel estimators of the multidimensional PDF $p_{(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})}$ of the random vector

$$(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)}) = \begin{bmatrix} \tilde{\mathbf{v}}_{(z),j_1}^T \\ \vdots \\ \tilde{\mathbf{v}}_{(z),j_m}^T \end{bmatrix} \Psi(\xi). \quad (3.15)$$

- For $1 \leq m \leq M$, $\hat{\mathcal{L}}_{(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})} \left(\left\{ (\tilde{C}_{j_1}(\theta_n), \dots, \tilde{C}_{j_m}(\theta_n)), 1 \leq n \leq \nu \right\} \right)$ is the estimation of the multidimensional log-likelihood of random vector $(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})$ that is evaluated at the experimental points $\left\{ (\tilde{C}_{j_1}(\theta_n), \dots, \tilde{C}_{j_m}(\theta_n)), 1 \leq n \leq \nu \right\}$, such that:

$$\begin{aligned} \hat{\mathcal{L}}_{(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})} \left(\left\{ (\tilde{C}_{j_1}(\theta_n), \dots, \tilde{C}_{j_m}(\theta_n)), 1 \leq n \leq \nu \right\} \right) \\ = \sum_{n=1}^{\nu} \ln \hat{p}_{(\tilde{U}_{j_1}^{(z)}, \dots, \tilde{U}_{j_m}^{(z)})} \left((\tilde{C}_{j_1}(\theta_n), \dots, \tilde{C}_{j_m}(\theta_n)) \right). \end{aligned} \quad (3.16)$$

- for $1 \leq P \leq M$, if the set $\mathcal{B} = \{\mathbf{b}_1, \dots, \mathbf{b}_P\}$ gathers P vectors that are in $\mathbb{R}^M$, $\text{Ker}(\mathcal{B}) = \{\mathbf{ker}_1^{\mathcal{B}}, \dots, \mathbf{ker}_Q^{\mathcal{B}}\}$ is an orthonormal basis of the null space of $\mathcal{B}$, such that:

$$\langle \mathbf{ker}_q^{\mathcal{B}}, \mathbf{ker}_{q'}^{\mathcal{B}} \rangle = \delta_{qq'}, \quad \langle \mathbf{ker}_q^{\mathcal{B}}, \mathbf{b}_p \rangle = 0, \quad 1 \leq p \leq P, \quad 1 \leq q, q' \leq Q. \quad (3.17)$$

- for $1 \leq m \leq M$, $\mathcal{S}^{(m)}(1)$ correspond to the m -dimension unit hypersphere, such that:

$$\mathcal{S}^{(m)}(1) = \{\mathbf{s} \in \mathbb{R}^m, \|\mathbf{s}\| = 1\}. \quad (3.18)$$

In addition, we denote by $\mathbf{S}^{(m)}$ the m -dimension random vector that is uniformly distributed on $\mathcal{S}^{(m)}(1)$. If Ξ is a m -dimension random vector whose components are centered, independent, normally distributed of variance equal to 1, as the distribution of Ξ is invariant by rotation, it can be seen that if $\{\Xi(\Theta_1), \dots, \Xi(\Theta_Q)\}$ are Q independent realizations of Ξ , the set

$$\left\{ \mathbf{S}^{(m)}(\Theta_1) = \Xi(\Theta_1) / \|\Xi(\Theta_1)\|, \dots, \mathbf{S}^{(m)}(\Theta_Q) = \Xi(\Theta_Q) / \|\Xi(\Theta_Q)\| \right\} \quad (3.19)$$

gathers Q independent realizations of random vector $\mathbf{S}^{(m)}$.

3.2.3 Theoretical frame

For $m \geq 1$, let

- $[R]$ be a $(m+1 \times m+1)$ real matrix such that

$$[R] = \begin{bmatrix} [R_{m,m}] & \mathbf{R}_* \\ \mathbf{R}_*^T & R_{**} \end{bmatrix}, \quad (3.20)$$

where $[R_{m,m}]$ is a $(m \times m)$ real and symmetrical matrix, $\mathbf{R}_*$ is a m -dimension real vector and $R_{**} \geq 0$.

- $[z]$ be a $(m \times N)$ real matrix such that $[z][z]^T = [R_{m,m}]$.
- $\mathbf{v}$ be a N -dimension real vector.

Proposition 1 *The matrix*

$$[Z] = \begin{bmatrix} [z] \\ \mathbf{v}^T \end{bmatrix} \quad (3.21)$$

fulfills the orthogonality constraint $[Z][Z]^T = [R]$ if and only if vector $\mathbf{v}$ verifies:

$$\mathbf{v} = [V] \begin{pmatrix} \boldsymbol{\alpha} \\ \boldsymbol{\beta} \end{pmatrix}, \quad \boldsymbol{\alpha} = [\ell]^{-1}[U]^T \mathbf{R}_*, \quad (3.22)$$

where $\boldsymbol{\beta}$ is any $(N-m)$ -dimension vector, for which norm is given by

$$\|\boldsymbol{\beta}\| = \sqrt{R_{**} - \|\boldsymbol{\alpha}\|^2}, \quad (3.23)$$

$[U]$ is a $(m \times m)$ real orthogonal matrix, $[V]$ is a $(N \times N)$ real orthogonal matrix, and $[\ell]$ is a $(m \times m)$ real and strictly positive-definite diagonal matrix, such that:

$$[z] = [U] \begin{bmatrix} [\ell] & \begin{bmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{bmatrix} \end{bmatrix} [V]^T, \quad [V]^T[V] = [I_N], \quad [U]^T[U] = [I_m]. \quad (3.24)$$

□ **Proof:** We have the following equivalences:

$$\begin{aligned} & [Z][Z]^T = [R] \\ & \Leftrightarrow [z]\mathbf{v} = \mathbf{R}_*, \quad \|\mathbf{v}\|^2 = R_{**}, \text{ by definition of } [R] \\ & \Leftrightarrow [U] \begin{bmatrix} [\ell] & \begin{bmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{bmatrix} \end{bmatrix} [V]^T \mathbf{v} = \mathbf{R}_*, \quad \|\mathbf{v}\|^2 = R_{**} \\ & \Leftrightarrow \mathbf{v} = [V] \begin{pmatrix} \boldsymbol{\alpha} \\ \boldsymbol{\beta} \end{pmatrix}, \quad \boldsymbol{\alpha} = [\ell]^{-1}[U]^T \mathbf{R}_*, \quad \|\boldsymbol{\beta}\| = \sqrt{R_{**} - \|\boldsymbol{\alpha}\|^2}. \end{aligned} \quad (3.25)$$

□

Corrolary 1 If $[R] = [I_{m+1}]$, the matrix

$$[Z] = \begin{bmatrix} [z] \\ \mathbf{v}^T \end{bmatrix} \quad (3.26)$$

fulfills the orthogonality constraint $[Z][Z]^T = [I_{m+1}]$ if and only if vector $\mathbf{v}$ verifies:

$$\mathbf{v} = \sum_{q=1}^{N-m} \mathbf{ker}_q^{\mathcal{B}^{(m)}} \beta_q, \quad \|\beta\| = 1, \quad (3.27)$$

where $\mathcal{B}^{(m)} = \{z_1, \dots, z_m\}$, $[z] = \begin{bmatrix} z_1^T \\ \vdots \\ z_m^T \end{bmatrix}$.

□ **Proof:** If $[R] = [I_{m+1}]$, then $R_{**} = 1$ and $\mathbf{R}_* = \mathbf{0}$, which leads directly to the result.
□

3.2.4 Iterative algorithms

From Section 1.5, it has been shown that a good approach to numerically identify the PCE projection matrix $[y]$ of $\mathbf{C}$, such that

$$\mathbf{C} \approx [y]\Psi(\xi), \quad (3.28)$$

is to search $[y]$ as the *direct* solution of the following optimization problem:

$$[y] \approx [y_{\mathcal{O}_C}^Z] = \arg \max_{[y^*] \in \mathcal{W}} \mathcal{C} \left([C^{\text{exp}}(\nu)], [y^*], [\Psi(\nu^{\text{chaos}})] \right), \quad (3.29)$$

where the cost function, $\mathcal{C}$, is defined by Eqs. (1.58), (1.59) and (1.59) and:

$$\mathcal{W} = \left\{ [y^{(1)}], \dots, [y^{(Z)}] \right\}, \quad \left\{ [y^{(1)}], \dots, [y^{(Z)}] \right\} \subset \mathcal{O}_C. \quad (3.30)$$

Alternatively, it has been underlined in Section 3.2.1 that matrix $[y]$ can be equivalently searched as the *indirect* solution of the following problem:

$$[y] \approx [\widehat{V}(\nu)][\widehat{\lambda}(\nu)]^{1/2}[u_{\mathcal{O}_C}^Z], \quad (3.31)$$

$$[u_{\mathcal{O}_C}^Z] = \arg \max_{[u^*] \in \widetilde{\mathcal{W}}} \mathcal{C} \left([\widetilde{C}^{\text{exp}}(\nu)], [u^*], [\Psi(\nu^{\text{chaos}})] \right), \quad (3.32)$$

where

$$[\widetilde{C}^{\text{exp}}(\nu)] = [\widehat{\lambda}(\nu)]^{-1/2}[\widehat{V}(\nu)]^T[C^{\text{exp}}(\nu)], \quad (3.33)$$

$$\widetilde{\mathcal{W}} = \left\{ [u^{(1)}], \dots, [u^{(Z)}] \right\}, \quad \left\{ [u^{(1)}], \dots, [u^{(Z)}] \right\} \subset V_{N,M}. \quad (3.34)$$

Hence, for the solving of these two optimization problems, the more adapted to the realizations of $\mathbf{C}$ and $\widetilde{\mathbf{C}}$ these sets $\mathcal{W}$ and $\widetilde{\mathcal{W}}$ are, the more relevant the identification of $[y]$ is likely to be. As an extension of the work achieved in [38], this section presents two iterative algorithms that stem from the theoretical developments of Section 3.2.3 to compute line by line matrices $[y^{(z)}]$ or $[u^{(z)}]$, $1 \leq z \leq Z$, that are particularly well adapted to the available information about $\mathbf{C}$ and $\widetilde{\mathbf{C}}$ respectively.

Direct method

In the following, z and Q verify $1 \leq z \leq Z$ and $Q \geq 1$.

Initialization. The $j_1^{(z)th}$ line of $[y^{(z)}]$, $\mathbf{v}_{(z),j_1}^T$, is first searched in order to maximize the unidimensional likelihood:

$$\mathbf{v}_{(z),j_1}^T = \arg \min_{\mathbf{v}_{(z),j_1}^T \in \mathcal{T}^{(1)}(Q)} \widehat{\mathcal{L}}_{U_{j_1}^{(z)}}(\{C_{j_1}(\theta_n), 1 \leq n \leq \nu\}), \quad (3.35)$$

where $\{\mathbf{S}^{(N)}(\Theta_1), \dots, \mathbf{S}^{(N)}(\Theta_Q)\}$ gathers Q independent realizations of $\mathbf{S}^{(N)}$ and $\mathcal{T}^{(1)}(Q) = \left\{ \sqrt{[\widehat{R}_{CC}]_{j_1 j_1}} \mathbf{S}^{(N)}(\Theta_1), \dots, \sqrt{[\widehat{R}_{CC}]_{j_1 j_1}} \mathbf{S}^{(N)}(\Theta_Q) \right\}$, such that if $\mathbf{v}_{(z),j_1}^T$ is in $\mathcal{T}^{(1)}(Q)$, $\|\mathbf{v}_{(z),j_1}^T\|^2 = [\widehat{R}_{CC}]_{j_1 j_1}$.

Iteration. For $2 \leq i \leq M$, the $j_i^{(z)th}$ line of $[y^{(z)}]$, $\mathbf{v}_{(z),j_i}^T$, is then searched such that:

$$\tilde{\mathbf{v}}_{(z),j_i}^T = \arg \min_{\tilde{\mathbf{v}}_{(z),j_i}^T \in \mathcal{T}^{(i)}(Q)} \widehat{\mathcal{L}}_{(U_{j_1}^{(z)}, \dots, U_{j_i}^{(z)})}(\{(C_{j_1}(\theta_n), \dots, C_{j_i}(\theta_n)), 1 \leq n \leq \nu\}), \quad (3.36)$$

where $\mathcal{T}^{(i)}(Q) = \{\mathbf{t}^1, \dots, \mathbf{t}^Q\}$ gathers Q real vectors with values in $\mathbb{R}^M$ such that:

$$\mathbf{t}^q = [V^{j_i}] \left(\frac{\boldsymbol{\alpha}^{j_i}}{\sqrt{R_{**}^{j_i} - \|\boldsymbol{\alpha}^{j_i}\|^2}} \mathbf{S}^{(N-i+1)}(\Theta_q) \right), \quad 1 \leq q \leq Q, \quad (3.37)$$

$$\boldsymbol{\alpha}^{j_i} = [\ell^{j_i}]^{-1} [U^{j_i}]^T \mathbf{R}_*^{j_i}, \quad (3.38)$$

$$\begin{bmatrix} \mathbf{v}_{(z),j_1}^T \\ \vdots \\ \mathbf{v}_{(z),j_{i-1}}^T \end{bmatrix} = [U^{j_i}] \left[\begin{bmatrix} \ell_1^{j_i} & 0 & \cdots & 0 \\ 0 & \ell_2^{j_i} & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \ell_{i-1}^{j_i} \end{bmatrix} \begin{bmatrix} 0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & 0 \end{bmatrix} \right] [V^{j_i}]^T, \quad (3.39)$$

$$[V^{j_i}]^T [V^{j_i}] = [I_N], \quad [U^{j_i}]^T [U^{j_i}] = [I_{i-1}]. \quad (3.40)$$

and $\{\mathbf{S}^{(N-i+1)}(\Theta_q), 1 \leq q \leq Q\}$ gathers Q independent realizations of random vector $\mathbf{S}^{(N-i+1)}$.

Indirect method

For $1 \leq z \leq Z$, and $Q \geq 1$, $[u^{(z)}]$ is defined according to the following iterative algorithm.

Initialization. The $j_1^{(z)th}$ line of $[u^{(z)}]$, $\tilde{\mathbf{v}}_{(z),j_1}^T$, is first searched in order to maximize the unidimensional likelihood:

$$\tilde{\mathbf{v}}_{(z),j_1}^T = \arg \min_{\tilde{\mathbf{v}}_{(z),j_1}^T \in \tilde{\mathcal{T}}^{(1)}(Q)} \widehat{\mathcal{L}}_{\tilde{U}_{j_1}^{(z)}}(\{\tilde{C}_{j_1}(\theta_n), 1 \leq n \leq \nu\}), \quad (3.41)$$

where $\tilde{\mathcal{T}}^{(1)}(Q)$ gathers Q independent realizations of random vector $\mathbf{S}^{(N)}$.

Iteration. For $2 \leq i \leq M$, the $j_i^{(z)}$ th line of $[u^{(z)}]$, $\tilde{\mathbf{v}}_{(z),j_i}^T$, is then searched such that:

$$\tilde{\mathbf{v}}_{(z),j_i}^T = \arg \min_{\tilde{\mathbf{v}}_{(z),j_i}^T \in \tilde{\mathcal{T}}^{(i)}(Q)} \hat{\mathcal{L}}_{(\tilde{v}_{j_1}^{(z)}, \dots, \tilde{v}_{j_i}^{(z)})} \left(\left\{ \left(\tilde{C}_{j_1}(\theta_n), \dots, \tilde{C}_{j_i}(\theta_n) \right), 1 \leq n \leq \nu \right\} \right), \quad (3.42)$$

where $\tilde{\mathcal{T}}^{(i)}(Q) = \{\tilde{\mathbf{t}}^1, \dots, \tilde{\mathbf{t}}^Q\}$ gathers Q real vectors with values in $\mathbb{R}^M$ such that:

$$\tilde{\mathbf{t}}^q = \sum_{k=1}^{M-i+1} \mathbf{ker}_k^{\mathcal{B}^{(i)}} S_k^{(N-i+1)}(\Theta_q), \quad 1 \leq q \leq Q, \quad (3.43)$$

and $\{\mathbf{S}^{(M-i+1)}(\Theta_q), 1 \leq q \leq Q\}$ gathers Q independent realizations of random vector $\mathbf{S}^{(M-i+1)}$, and $\mathcal{B}^{(i)} = \{\tilde{\mathbf{v}}_{(z),j_1}, \dots, \tilde{\mathbf{v}}_{(z),j_{i-1}}\}$.

Comments on the iterative algorithms

- First, from Section 3.2.3, such algorithms allows us to generate matrices $[y^{(z)}]$ and $[u^{(z)}]$ that verify the orthogonality constraints $[y^{(z)}][y^{(z)}]^T = [\hat{R}_{CC}(\nu)]$ and $[u^{(z)}][u^{(z)}]^T = [I_M]$.
- Thanks to the iterative construction, each line of these matrices are moreover defined to be adapted at most to the available realizations of $\mathbf{C}$ and $\tilde{\mathbf{C}}$.
- By imposing on the vectors $\mathbf{t}^q$ (for the direct method) and $\tilde{\mathbf{t}}^q$ (for the indirect method) to be uniformly distributed on their definition domains, we try to explore as objectively as possible the sets $\mathcal{O}_C$ and $V_{N,M}$.
- These algorithms are indexed by integer Q . To solve Eqs. (3.29) or (3.32), the total cost is therefore globally proportional to $Q \times Z$. From a numerical point of view, it appears that the accuracy of the results is however much more dependent on Q than Z . For a limited computational cost, it is thus advised to choose Q as high as possible, even if that forces Z to be inferior to 10.

3.2.5 Adaptations to the case $\nu < M$

Motivations

When the information about $\mathbf{C}$ is limited, and more precisely, when ν is lower than M , the direct and indirect iterative algorithms formerly introduced are no more accurate. Indeed, the rank r of matrix $[\hat{R}_{CC}(\nu)]$ is by construction lower than ν , such that

$$\hat{\lambda}_{M-\nu+1}(\nu) = \hat{\lambda}_{M-\nu+2}(\nu) = \dots = \hat{\lambda}_M(\nu) = 0. \quad (3.44)$$

Matrix $[\hat{V}(\nu)][\hat{\lambda}(\nu)]^{1/2}$ is no more invertible, and imposing $[y]$ to verify the constraint $[y][y]^T = [\hat{R}_{CC}(\nu)]$ is therefore equivalent to project $\mathbf{C}$ in the (r) -dimension image space of $[\hat{R}_{CC}(\nu)]$. If we want $\mathbf{C}$ to remain in its M -dimension space, the former constraint on $[y]$ has to be loosened a little, such that $[y][y]^T \approx [\hat{R}_{CC}(\nu)]$. Keeping in mind that the rank of each $(r \times r)$ sub-matrix of $[\hat{R}_{CC}(\nu)]$ is also r , the possibility we propose to loosen this constraint is based on a block by block adaptation of the direct iterative algorithm defined in Section 3.2.4 to generate matrices $[y^{(z)}]$, $1 \leq z \leq Z$, such that $[y^{(z)}][y^{(z)}]^T \approx [\hat{R}_{CC}(\nu)]$.

Adapted PCE identification iterative algorithm

For $1 \leq z \leq Z$, and $Q \geq 1$.

Initialization. The $j_1^{(z)th}$ line of $[y^{(z)}]$, $\mathbf{v}_{(z),j_1}^T$, is searched in order to maximize the unidimensional likelihood:

$$\mathbf{v}_{(z),j_1}^T = \arg \min_{\mathbf{v}_{(z),j_1}^T \in \mathcal{T}^{(1)}(Q)} \widehat{\mathcal{L}}_{U_{j_1}^{(z)}}(\{C_{j_1}(\theta_n), 1 \leq n \leq \nu\}), \quad (3.45)$$

where $\mathcal{T}^{(1)}(Q) = \left\{ \sqrt{[\widehat{R}_{CC}]_{j_1 j_1}} \mathbf{S}^{(N)}(\Theta_1), \dots, \sqrt{[\widehat{R}_{CC}]_{j_1 j_1}} \mathbf{S}^{(N)}(\Theta_Q) \right\}$ is the space that has already been introduced in Section 3.2.4.

Iteration - first part. For $2 \leq i \leq r$, the $j_i^{(z)th}$ line of $[y^{(z)}]$, $\mathbf{v}_{(z),j_i}^T$, is searched such that:

$$\tilde{\mathbf{v}}_{(z),j_i}^T = \arg \min_{\tilde{\mathbf{v}}_{(z),j_i}^T \in \mathcal{T}^{(i)}(Q)} \widehat{\mathcal{L}}_{(U_{j_1}^{(z)}, \dots, U_{j_i}^{(z)})}(\{(C_{j_1}(\theta_n), \dots, C_{j_i}(\theta_n)), 1 \leq n \leq \nu\}), \quad (3.46)$$

where $\mathcal{T}^{(i)}(Q)$ has also been defined in Section 3.2.4.

Iteration - second part. For $r+1 \leq i \leq M$, we define

$$\mathbf{R}_* = ([\widehat{R}_{CC}]_{j_1 j_i}, [\widehat{R}_{CC}]_{j_2 j_i}, \dots, [\widehat{R}_{CC}]_{j_{i-1} j_i}), \quad (3.47)$$

$$\begin{cases} \mathcal{I}^{(r)} = \{i_1, \dots, i_r\} \subset \{j_1, \dots, j_{i-1}\}, \\ |[\widehat{R}_{CC}]_{i_1 j_i}| \geq \dots \geq |[\widehat{R}_{CC}]_{i_r j_i}| \geq \max_{j \notin \mathcal{I}^{(r)}} |[\widehat{R}_{CC}]_{j j_i}|. \end{cases} \quad (3.48)$$

The $j_i^{(z)th}$ line of $[y^{(z)}]$, $\mathbf{v}_{(z),j_i}^T$, is then searched such that:

$$\tilde{\mathbf{v}}_{(z),j_i}^T = \arg \min_{\tilde{\mathbf{v}}_{(z),j_i}^T \in \mathbb{T}^{(i)}(Q)} \widehat{\mathcal{L}}_{(U_{j_1}^{(z)}, \dots, U_{j_i}^{(z)})}(\{(C_{j_1}(\theta_n), \dots, C_{j_i}(\theta_n)), 1 \leq n \leq \nu\}), \quad (3.49)$$

where $\mathbb{T}^{(i)}(Q) = \{\mathbf{t}^1, \dots, \mathbf{t}^Q\}$ gathers Q real vectors with values in $\mathbb{R}^M$ such that:

$$\begin{bmatrix} \mathbf{v}_{(z),i_1}^T \\ \vdots \\ \mathbf{v}_{(z),i_r}^T \\ \mathbf{t}^q \end{bmatrix} \begin{bmatrix} \mathbf{v}_{(z),i_1}^T \\ \vdots \\ \mathbf{v}_{(z),i_r}^T \\ \mathbf{t}^q \end{bmatrix}^T = \begin{bmatrix} [\widehat{R}_{CC}]_{i_1 i_1} & \cdots & [\widehat{R}_{CC}]_{i_1 i_r} & [\widehat{R}_{CC}]_{i_1 j_i} \\ \vdots & \ddots & \vdots & \vdots \\ [\widehat{R}_{CC}]_{i_r i_1} & \cdots & [\widehat{R}_{CC}]_{i_r i_r} & [\widehat{R}_{CC}]_{i_r j_i} \\ [\widehat{R}_{CC}]_{j_i i_1} & \cdots & [\widehat{R}_{CC}]_{j_i i_r} & [\widehat{R}_{CC}]_{j_i j_i} \end{bmatrix}, \quad 1 \leq q \leq Q, \quad (3.50)$$

which have been randomly generated using the same developments than in Section 3.2.4.

Therefore, such an algorithm allows us to build matrices $[y^{(z)}]$ such that the highest terms in absolute value of $[\widehat{R}_{CC}]$ are exactly reproduced in $[y^{(z)}][y^{(z)}]^T$.

3.3 Numerical stabilization of the polynomial basis in high dimension

As it has been presented in the former sections, the $(N \times \nu^{\text{chaos}})$ real matrix $[\Psi(\nu^{\text{chaos}})]$ gathers ν^{chaos} independent realizations of the N -dimension random vector $\Psi(\xi)$. Moreover, the numerical identification of the PCE coefficients $[y]$ can be seen as the minimization of a cost function involving the elements of the $(M \times \nu^{\text{chaos}})$ real matrix of independent realizations $[U] = [y][\Psi(\nu^{\text{chaos}})]$ of random vector $\mathbf{U} = [y]\Psi(\xi)$ and the elements of the $(M \times \nu)$ real matrix of independent realizations $[C^{\text{exp}}(\nu)] = [\mathbf{C}(\theta_1) \ \cdots \ \mathbf{C}(\theta_\nu)]$ of $\mathbf{C}$. In theoretical terms, this cost function should be minimum when the multidimensional PDF p_U of $\mathbf{U}$ is as near as possible to the multidimensional PDF p_C of $\mathbf{C}$. In practical terms, this cost function is however minimum when $\hat{p}_U$ is as near as possible to $\hat{p}_C$, where $\hat{p}_U$ and $\hat{p}_C$ are the multidimensional non parametric estimators of p_U and p_C defined by Eq. (1.55). With respect to ν and ν^{chaos} , three bias are then introduced in the PCE identification:

- a bias due to a lack of information about $\mathbf{C}$:

$$b^{(1)}(\nu) = \int_{\mathbb{R}^M} |\hat{p}_C(\mathbf{x}) - p_C(\mathbf{x})| d\mathbf{x}, \quad (3.51)$$

- a bias due to a lack of information about $\mathbf{U}$:

$$b^{(2)}(\nu^{\text{chaos}}) = \int_{\mathbb{R}^M} |\hat{p}_U(\mathbf{x}) - p_U(\mathbf{x})| d\mathbf{x}, \quad (3.52)$$

- a bias due to the truncation and to the fact that the global maximum is not necessary reached:

$$b^{(3)}(\nu, \nu^{\text{chaos}}) = \int_{\mathbb{R}^M} |\hat{p}_C(\mathbf{x}) - \hat{p}_U(\mathbf{x})| d\mathbf{x}. \quad (3.53)$$

These three bias could also be expressed with respect to the statistical moments of $\mathbf{C}$ and $\mathbf{U}$. For instance, when focusing on the autocorrelation matrix, let err^1 , err^2 and err^3 be the autocorrelation errors corresponding respectively to the bias $b^{(1)}$, $b^{(2)}$ and $b^{(3)}$:

$$err^1(\nu) = \left\| [R_{CC}] - [\hat{R}_{CC}(\nu)] \right\|_F / \left\| [R_{CC}] \right\|_F, \quad (3.54)$$

$$err^2(\nu^{\text{chaos}}) = \left\| [R_{CC}^{\text{chaos}}(N)] - [\hat{R}_{UU}(\nu^{\text{chaos}})] \right\|_F / \left\| [R_{CC}^{\text{chaos}}(N)] \right\|_F, \quad (3.55)$$

$$err^3(\nu, \nu^{\text{chaos}}) = \left\| [\hat{R}_{UU}(\nu^{\text{chaos}})] - [\hat{R}_{CC}(\nu)] \right\|_F / \left\| [\hat{R}_{CC}(\nu)] \right\|_F, \quad (3.56)$$

where $\|\cdot\|_F$ is the Frobenius norm of matrices, and where it is reminded from Eqs. (1.37) and (1.50) that:

$$\left\{ \begin{array}{l} [\hat{R}_{CC}(\nu)] = \frac{1}{\nu} [C^{\text{exp}}(\nu)] [C^{\text{exp}}(\nu)]^T, \\ [\hat{R}_{UU}(\nu^{\text{chaos}})] = \frac{1}{\nu^{\text{chaos}}} [U] [U]^T = [y] \left(\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T \right) [y]^T, \\ [R_{CC}^{\text{chaos}}(N)] = [y] [y]^T. \end{array} \right. \quad (3.57)$$

Hence, the smaller these three errors are, the more precise the PCE identification is. The ν independent realizations $\{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$ being the maximum available information about $\mathbf{C}$, the bias $b^{(1)}$ and the autocorrelation error err^1 cannot be decreased, whereas the set $\mathcal{O}_C$, which was introduced to guarantee that $[R_{CC}^{\text{chaos}}(N)] = [\hat{R}_{CC}(\nu)]$, aims at reducing $b^{(2)}$, $b^{(3)}$, err^2 and err^3 . Therefore, imposing $[y]$ to be in $\mathcal{O}_C$ leads us to:

$$\begin{aligned} err^2(\nu^{\text{chaos}}) &= err^3(\nu, \nu^{\text{chaos}}) \\ &= \left\| [y] \left(\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N] \right) [y]^T \right\|_F / \|[y][I_N][y]^T\|_F. \end{aligned} \quad (3.58)$$

The following asymptotic property can thus be deduced from Eq. (1.40):

$$\lim_{\nu^{\text{chaos}} \rightarrow +\infty} err^2(\nu^{\text{chaos}}) = \lim_{\nu^{\text{chaos}} \rightarrow +\infty} err^3(\nu, \nu^{\text{chaos}}) = 0, \quad (3.59)$$

which is equivalent to say that the larger ν^{chaos} is, the more accurate the PCE identification should be. However, from a practical point of view, the value of ν^{chaos} is fixed by the available computation resources. As an extension of the work presented in [76], this section aims at quantifying the divergence of the ratio:

$$r = \left\| \frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N] \right\|_F / \|[I_N]\|_F, \quad (3.60)$$

when the truncation parameters N_g and p , which have been introduced in Section 1.5.4, increase for several statistical measures. From Eq. (3.58), r , defined by Eq. (3.60), can be seen as a general characterization of the autocorrelation errors err^2 and err^3 . This divergence being very detrimental to the PCE identification in high dimension, a new decomposition of the PCE coefficient matrix $[y]$ will be then presented in this section to make err^2 and err^3 be zero for any value of N_g and p .

3.3.1 Decomposition of the matrix of independent realizations

To better emphasize the influence of the truncation parameters on the ratio r , a rewriting of the matrix $[\Psi(\nu^{\text{chaos}})]$ is first presented.

Theoretical basis of the decomposition

From Eq. (1.39), matrix $[\Psi(\nu^{\text{chaos}})]$ gathers ν^{chaos} columns $\{\Psi(\xi(\theta_n)), 1 \leq n \leq \nu^{\text{chaos}}\}$, which are independent realizations of the N -dimension PCE random vector $\Psi(\xi) = (\psi_1(\xi), \dots, \psi_N(\xi))$. This basis being orthonormal leads us to the asymptotic condition on $[\Psi(\nu^{\text{chaos}})]$, defined by Eq. (1.40). Moreover, Eq. (1.62) implies that $[\Psi(\nu^{\text{chaos}})]$ can be expressed as:

$$[\Psi(\nu^{\text{chaos}})] = [A][M], \quad (3.61)$$

where $[A]$ is the $(N \times N)$ real matrix that gathers the coefficients of the orthonormal polynomials with respect to the probability measure of the N_g -dimension PCE germ, $\xi = (\xi_1, \dots, \xi_{N_g})$, and $[M]$ is a $(N \times \nu^{\text{chaos}})$ real matrix, which gathers ν^{chaos} independent realizations of the random vector $\mathcal{E}(\xi, p)$, such that:

$$[M] = [\mathcal{E}(\xi(\theta_1), p) \quad \dots \quad \mathcal{E}(\xi(\theta_{\nu^{\text{chaos}}}), p)], \quad (3.62)$$

$$\mathcal{E}(\xi, p) = (\mathcal{M}_{\alpha^{(1)}}(\xi), \dots, \mathcal{M}_{\alpha^{(N)}}(\xi)), \quad (3.63)$$

$$\mathcal{M}_{\alpha^{(q)}}(\xi) = \xi_1^{\alpha_1^{(q)}} \times \dots \times \xi_{N_g}^{\alpha_{N_g}^{(q)}}, \quad 1 \leq q \leq N, \quad (3.64)$$

where $\mathcal{A}_p = \{\alpha^{(1)}, \dots, \alpha^{(N)}\}$ is the set that gathers the N elements of $\mathbb{N}^{N_g}$ that verify the following constraint:

$$\sum_{\ell=1}^{N_g} \alpha_\ell^{(q)} \leq p, \quad 1 \leq q \leq N. \quad (3.65)$$

If $[A]$ is independent of $[M]$, Eq. (3.61) certifies that, if the columns of $[M]$ are independent, then the columns of $[\Psi(\nu^{\text{chaos}})]$ stay independent. Let $[R_{\mathcal{E}}]$ be the autocorrelation matrix of the random vector $\mathcal{E}(\xi, p)$:

$$[R_{\mathcal{E}}] = E \left[\mathcal{E}(\xi, p) \mathcal{E}(\xi, p)^T \right]. \quad (3.66)$$

It can be deduced from Eqs. (1.40), (3.61), (3.62) and (3.66) that:

$$[R_{\mathcal{E}}] = \lim_{\nu^{\text{chaos}} \rightarrow +\infty} \frac{1}{\nu^{\text{chaos}}} [M][M]^T = [A]^{-1} [A]^{-T}. \quad (3.67)$$

According to this decomposition, computing the classical Gram-Schmidt orthogonalization to identify the polynomial basis coefficients only requires the calculation of $[A]^{-T}$, which corresponds to the Cholesky decomposition matrix of the positive definite matrix $[R_{\mathcal{E}}]$. Hence, by construction, the matrix $[\Psi(\nu^{\text{chaos}})]$ can be written as the product of a lower triangular matrix $[A]$ and a matrix $[M]$ of independent realizations of a multi-index random vector $\mathcal{E}(\xi, p)$.

Practical computation of matrix $[\Psi(\nu^{\text{chaos}})]$

Thanks to Eq. (3.61), matrix $[\Psi(\nu^{\text{chaos}})]$ can be numerically computed without requiring computational recurrence formula nor algebraic explicit representation. An illustration of the method is presented hereinafter for a PCE based on a Gaussian measure. This development can be directly extended to any value of p and N_g , as well as to other statistical measures. Let ξ_1 and ξ_2 be two independent normalized Gaussian random variables, such that $\xi = (\xi_1, \xi_2)$, and $\alpha = (\alpha_1, \alpha_2)$. Choosing $p = 2$ and $N_g = 2$, which corresponds to $N = 6$, leads us to the following definition of $\mathcal{E}(\xi, p)$:

$$\mathcal{E}(\xi, 2) = (1, \xi_1, \xi_2, \xi_1 \xi_2, \xi_1^2, \xi_2^2). \quad (3.68)$$

According to this equation, matrix $[M]$ can thus be easily deduced from ν^{chaos} independent realizations of ξ . Moreover, let $[\alpha]$ be the $(N_g \times N)$ real matrix which gathers the admissible values for α in $\mathcal{A}_p$:

$$[\alpha] = \begin{bmatrix} 0 & 1 & 0 & 1 & 2 & 0 \\ 0 & 0 & 1 & 1 & 0 & 2 \end{bmatrix} \leftrightarrow \mathcal{A}_p = \{(0, 0), (1, 0), (0, 1), (1, 1), (2, 0), (0, 2)\}. \quad (3.69)$$

The random variables ξ_1 and ξ_2 being independent, normalized and Gaussian, the autocorrelation matrix $[R_{\mathcal{E}}]$ can thus be written as:

$$\begin{aligned} \forall i, j \in \{1, \dots, N\}, [R_{\mathcal{E}}]_{ij} &= E \left[\xi_1^{[\alpha]_{1i} + [\alpha]_{1j}} \times \dots \times \xi_{N_g}^{[\alpha]_{N_g i} + [\alpha]_{N_g j}} \right] \\ &= E \left[\xi_1^{[\alpha]_{1i} + [\alpha]_{1j}} \right] \times \dots \times E \left[\xi_{N_g}^{[\alpha]_{N_g i} + [\alpha]_{N_g j}} \right], \end{aligned} \quad (3.70)$$

where, for $1 \leq \ell \leq N_g$:

$$\begin{cases} E [\xi_{\ell}^q] = 0 & \text{if } q \text{ is not even,} \\ E [\xi_{\ell}^q] = \frac{q!}{(q/2)!2^{q/2}} & \text{if } q \text{ is even.} \end{cases} \quad (3.71)$$

Therefore, Eq. (3.67) allows us to numerically find back in $[A]$ the multidimensional Hermite polynomials $H_{\alpha_1} \times \dots \times H_{\alpha_{N_g}}$:

$$\forall x \in \mathbb{R}, \begin{cases} H_0(x_1) \times H_0(x_2) = 1 \\ H_1(x_1) \times H_0(x_2) = x_1 \\ H_0(x_1) \times H_1(x_2) = x_2 \\ H_1(x_1) \times H_1(x_2) = x_1 x_2 \\ H_2(x_1) \times H_0(x_2) = \frac{x_1^2 - 1}{\sqrt{2}} \\ H_0(x_1) \times H_2(x_2) = \frac{x_2^2 - 1}{\sqrt{2}} \end{cases} \leftrightarrow [A] = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ \frac{-1}{\sqrt{2}} & 0 & 0 & 0 & \frac{1}{\sqrt{2}} & 0 \\ \frac{-1}{\sqrt{2}} & 0 & 0 & 0 & 0 & \frac{1}{\sqrt{2}} \end{pmatrix}. \quad (3.72)$$

Noticing that:

- if ξ is a random variable uniformly distributed on $[-1, 1]$:

$$\begin{cases} E [\xi^q] = 0 & \text{if } q \text{ is not even,} \\ E [\xi^q] = \frac{1}{q+1} & \text{if } q \text{ is even,} \end{cases} \quad (3.73)$$

- if the random variable ξ is a random variable that is characterized by a normalized exponential distribution on $[0, +\infty[$:

$$E [\xi^q] = q!, \quad (3.74)$$

this method can directly be generalized to the uniform and exponential cases to compute the multidimensional Legendre and Laguerre polynomial coefficients, but also to an arbitrary probability measure for the germ ξ .

3.3.2 Influence of the truncation parameters and of the choice for the PCE probability measure

The convergence properties of ratio r when ν^{chaos} tends to infinity are strongly related to the statistical properties of germ ξ . This section aims therefore at emphasizing the dominant trends of this specific link, and to highlight the difficulties brought about by the divergence of ratio r , when trying to perform analysis of convergence in high dimension.

The definition of the Frobenius norm allows us to write that:

$$r = \left\| \frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N] \right\|_F / \| [I_N] \|_F = \sqrt{N} \Sigma(\nu^{\text{chaos}}), \quad (3.75)$$

where $\Sigma(\nu^{\text{chaos}})$ is such that:

$$\left\{\Sigma(\nu^{\text{chaos}})\right\}^2 = \frac{1}{N^2} \sum_{1 \leq i, j \leq N} \left(\left(\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N] \right)_{ij} \right)^2. \quad (3.76)$$

By construction, $\left\{\Sigma(\nu^{\text{chaos}})\right\}^2$ is an assessment of the mean value of the squared difference between the elements of $\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T$ and the elements of the unit matrix $[I_N]$. Hence, if $\left\{\Sigma(\nu^{\text{chaos}})\right\}^2$ remains constant when the size N of the polynomial basis increases, the ratio r should increase as $\sqrt{N}$. Moreover, Eqs. (3.61) and (3.67) yield,

$$\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N] = [A] \left(\frac{1}{\nu^{\text{chaos}}} [M] [M]^T - [R_{\mathcal{E}}] \right) [A]^T. \quad (3.77)$$

For all (i, j) in $\{1, \dots, N\}^2$, $[R_{\mathcal{E}}]_{ij}$ is such that:

$$[R_{\mathcal{E}}]_{ij} = E \left[\xi_1^{\alpha_1^{(i)} + \alpha_1^{(j)}} \times \dots \times \xi_{N_g}^{\alpha_{N_g}^{(i)} + \alpha_{N_g}^{(j)}} \right]. \quad (3.78)$$

Let $[\widehat{R}_{\mathcal{E}}]$ be the following estimator of $[R_{\mathcal{E}}]$:

$$[\widehat{R}_{\mathcal{E}}]_{ij} = \frac{1}{\nu^{\text{chaos}}} \sum_{n=1}^{\nu^{\text{chaos}}} \left(\Xi_1^{(n)} \right)^{\alpha_1^{(i)} + \alpha_1^{(j)}} \times \dots \times \left(\Xi_{N_g}^{(n)} \right)^{\alpha_{N_g}^{(i)} + \alpha_{N_g}^{(j)}}, \quad (3.79)$$

where $\left\{ \Xi^{(n)} = \left(\Xi_1^{(n)}, \dots, \Xi_{N_g}^{(n)} \right), 1 \leq n \leq \nu^{\text{chaos}} \right\}$ is a set of ν^{chaos} independent N_g -dimension random vectors, which have the same PDF than ξ . The central limit theorem yields that, for all (i, j) in $\{1, \dots, N\}^2$, we have:

$$\sqrt{\frac{\nu^{\text{chaos}}}{\text{Var} \left(\xi_1^{\alpha_1^{(i)} + \alpha_1^{(j)}} \times \dots \times \xi_{N_g}^{\alpha_{N_g}^{(i)} + \alpha_{N_g}^{(j)}} \right)}} \left([\widehat{R}_{\mathcal{E}}]_{ij} - [R_{\mathcal{E}}]_{ij} \right) \xrightarrow{\text{law}} \Xi, \quad (3.80)$$

where Ξ is a random variable that has a standard normal distribution, and $\text{Var}(\cdot)$ is the variance. Under this formalism, it can be noticed that $\frac{1}{\nu^{\text{chaos}}} [M] [M]^T$ is one particular realization of $[\widehat{R}_{\mathcal{E}}]$. Hence, from Eqs. (3.76), (3.77) and (3.80), we deduce that:

- if $\text{Var} \left(\xi_1^{\alpha_1^{(i)} + \alpha_1^{(j)}} \times \dots \times \xi_{N_g}^{\alpha_{N_g}^{(i)} + \alpha_{N_g}^{(j)}} \right) \leq \text{Var} \left(\xi_1^{\alpha_1^{(i)} + \alpha_1^{(j)}} \times \dots \times \xi_{N_g+1}^{\alpha_{N_g+1}^{(i)} + \alpha_{N_g+1}^{(j)}} \right)$, then $\Sigma(\nu^{\text{chaos}})$ potentially increases with respect to N_g .
- if $\text{Var} \left(\xi_{\ell}^{\alpha_{\ell}^{(i)}} \right) \leq \text{Var} \left(\xi_{\ell}^{\alpha_{\ell}^{(j)}} \right)$ for $\alpha_{\ell}^{(i)} \leq \alpha_{\ell}^{(j)}$, then $\Sigma(\nu^{\text{chaos}})$ potentially increases with respect to p .

As an illustration, for each couple (N_g, p) such that $1 \leq p \leq 10$ and $1 \leq N_g \leq 6$, three sets, $\{[\Psi_U^{(m)}(p, N_g)], 1 \leq m \leq 1000\}$, $\{[\Psi_G^{(m)}(p, N_g)], 1 \leq m \leq 1000\}$ and $\{[\Psi_E^{(m)}(p, N_g)], 1 \leq m \leq 1000\}$, are computed, such that $[\Psi_U^{(m)}(p, N_g)]$, $[\Psi_G^{(m)}(p, N_g)]$ and $[\Psi_E^{(m)}(p, N_g)]$ refer to particular $(N \times \nu^{\text{chaos}})$ real matrices of independent realizations of the basis $\{\psi_1(\xi), \dots, \psi_N(\xi)\}$, in the uniform, the Gaussian and the exponential cases, respectively. Hence, defining:

$$\begin{cases} r_U^m(\nu^{\text{chaos}}) = \left\| \frac{1}{\nu^{\text{chaos}}} [\Psi_U^{(m)}(p, N_g)] [\Psi_U^{(m)}(p, N_g)]^T - [I_N] \right\|_F / \| [I_N] \|_F, \\ r_G^m(\nu^{\text{chaos}}) = \left\| \frac{1}{\nu^{\text{chaos}}} [\Psi_G^{(m)}(p, N_g)] [\Psi_G^{(m)}(p, N_g)]^T - [I_N] \right\|_F / \| [I_N] \|_F, \\ r_E^m(\nu^{\text{chaos}}) = \left\| \frac{1}{\nu^{\text{chaos}}} [\Psi_E^{(m)}(p, N_g)] [\Psi_E^{(m)}(p, N_g)]^T - [I_N] \right\|_F / \| [I_N] \|_F, \end{cases} \quad (3.81)$$

allows us to compute, in each case, three approximations $err_U^{\text{ortho}}(p, N_g)$, $err_G^{\text{ortho}}(p, N_g)$ and $err_E^{\text{ortho}}(p, N_g)$ of the mean value of the ratio r , defined in Eq. (3.75), such that:

$$\begin{cases} err_U^{\text{ortho}}(p, N_g) = \frac{1}{1000} \sum_{m=1}^{1000} r_U^m(\nu^{\text{chaos}}), \\ err_G^{\text{ortho}}(p, N_g) = \frac{1}{1000} \sum_{m=1}^{1000} r_G^m(\nu^{\text{chaos}}), \\ err_E^{\text{ortho}}(p, N_g) = \frac{1}{1000} \sum_{m=1}^{1000} r_E^m(\nu^{\text{chaos}}). \end{cases} \quad (3.82)$$

For $\nu^{\text{chaos}} = 1000$, in Figure 3.1, the two factors which make the ratio r diverge with respect to p and N_g can therefore be emphasized. On the first hand, if increasing p or N_g does not increase the variance of the elements of $\mathcal{E}(\xi, p)$, which is the case if the PCE germ ξ is characterized by an uniform distribution (see Eq. (3.73)), the ratio r increases approximately as $\sqrt{N}$. On the other hand, if increasing p or N_g increases the variance of the element of $\mathcal{E}(\xi, p)$, as it is the case if the PCE germ ξ is characterized by a Gaussian or exponential distribution (see Eqs. (3.71) and (3.74)), the ratio r diverges very quickly with respect to the truncation parameters, and bias the PCE identification results.

As a conclusion, for a fixed value of ν^{chaos} , the difference $\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T - [I_N]$ increases when p and N_g increase. Therefore, imposing $[y]$ to be in $\mathcal{O}_C$ introduces a numerical bias in the PCE identification, which becomes very important when high values of p and N_g are needed. Such a phenomenon prevents thus to perform the analysis of convergence of the PCE in high dimension, especially when dealing with Gaussian and exponential PCE germs.

3.3.3 Adaptation of the optimization problem

In this section, fixed values for ν^{chaos} , p and N_g are considered. According to the notations of Section 3.3.1, a $(N \times \nu^{\text{chaos}})$ real matrix of independent realizations $[\Psi(\nu^{\text{chaos}})] = [A][M]$ can then be constructed. Under the condition $\nu^{\text{chaos}} \geq N$, $\frac{1}{\nu^{\text{chaos}}} [M][M]^T$ is positive definite by construction, which allows writing:

$$\frac{1}{\nu^{\text{chaos}}} [M][M]^T = [L][L]^T, \quad (3.83)$$

where $[L]$ is the Cholesky decomposition of $\frac{1}{\nu^{\text{chaos}}} [M][M]^T$, which yields:

$$\frac{1}{\nu^{\text{chaos}}} [\Psi(\nu^{\text{chaos}})] [\Psi(\nu^{\text{chaos}})]^T = [A][L][L]^T [A]^T = [B][B]^T, \quad (3.84)$$

$$[B] = [A][L]. \quad (3.85)$$

The matrix:

$$[\tilde{\Psi}] = [B]^{-1} [\Psi(\nu^{\text{chaos}})], \quad (3.86)$$

is then introduced, such that, by construction:

$$\frac{1}{\nu^{\text{chaos}}} [\tilde{\Psi}][\tilde{\Psi}]^T = [I_N]. \quad (3.87)$$

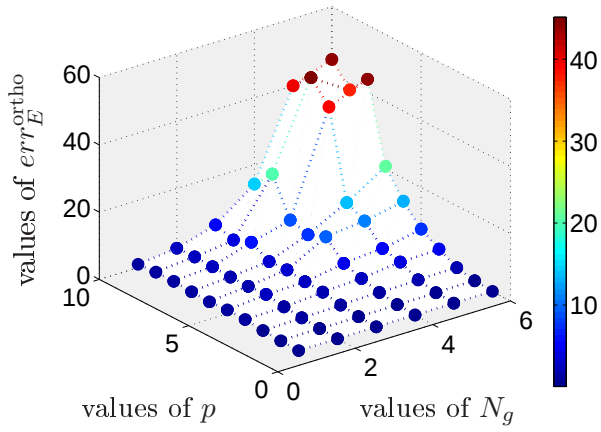
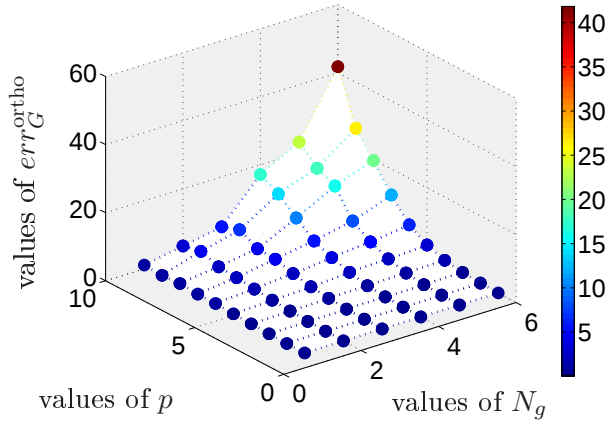
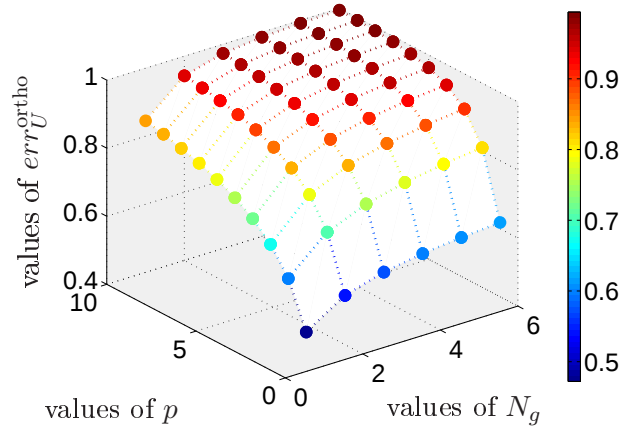


Figure 3.1: Graphs of the errors err_U^{ortho} , err_G^{ortho} and err_E^{ortho} with respect to the truncation parameters N_g and p .

Using the notations of Section 1.5, let $[y^*]$ be a $(M \times N)$ real matrix such that the random vector $\mathbf{U}$ is defined as:

$$\mathbf{U} = [y^*]\Psi(\xi). \quad (3.88)$$

Hence, ν^{chaos} independent realizations of $\mathbf{U}$ can be directly deduced from the matrix $[\Psi(\nu^{\text{chaos}})]$ and gathered in the matrix $[U] = [y^*][\Psi(\nu^{\text{chaos}})]$. Defining $[z]$ such that:

$$[z] = [y^*][B], \quad (3.89)$$

therefore yields the equality:

$$[U] = [y^*][\Psi(\nu^{\text{chaos}})] = ([z][B]^{-1}) \left([B][\tilde{\Psi}] \right) = [z][\tilde{\Psi}]. \quad (3.90)$$

If $[z]$ is in $\mathcal{O}_C$, $[z][z]^T = [\hat{R}_{CC}(\nu)]$, which implies that:

$$\begin{cases} [\hat{R}_{UU}(\nu^{\text{chaos}})] = \frac{1}{\nu^{\text{chaos}}} [U][U]^T = [z] \left(\frac{1}{\nu^{\text{chaos}}} [\tilde{\Psi}][\tilde{\Psi}]^T \right) [z]^T = [z][z]^T = [\hat{R}_{CC}(\nu)], \\ [R_{UU}] = E[\mathbf{U}\mathbf{U}^T] = \lim_{\nu^{\text{chaos}} \rightarrow \infty} [\hat{R}_{UU}(\nu^{\text{chaos}})] = [\hat{R}_{CC}(\nu)]. \end{cases} \quad (3.91)$$

From Eqs. (3.55) and (3.56), it can thus be deduced that imposing $[z]$ to be an element of $\mathcal{O}_C$ guarantees that, for any $\nu^{\text{chaos}} \geq N$, we have $\text{err}^2(\nu, \nu^{\text{chaos}}) = \text{err}^3(\nu, \nu^{\text{chaos}}) = 0$.

Hence, whereas the optimization problem defined by Eq. (1.61) is perturbed by autocorrelation errors, the new optimization problem:

$$\begin{cases} [y_{\mathcal{O}_C}^M] = [z_{\mathcal{O}_C}^M] [B]^{-1}, \\ [z_{\mathcal{O}_C}^M] = \arg \max_{[z^*] \in \mathcal{W}} \mathcal{C} \left([C^{\text{exp}}(\nu)], [z^*], [\tilde{\Psi}] \right), \end{cases} \quad (3.92)$$

is no more affected, which allows us to consider high values of the truncation parameters N_g and p . Equation (3.90) underlines that the two former optimization problems are equivalent, as the independent realizations of $\mathbf{U}$ have just been rewritten. Only the research set, for the PCE coefficient matrix, has been modified, which allows the numerical bias due to the finite dimension of $[\Psi(\nu^{\text{chaos}})]$ to be reduced.

Finally, if $[y]$ is the coefficients matrix of the truncated PCE, $\mathbf{C}^{\text{chaos}}(N)$, of random vector $\mathbf{C}$, such that $\mathbf{C}^{\text{chaos}}(N) = [y]\Psi(\xi)$, a good estimation of $[y]$ in high dimension can be computed by solving the optimization problem defined by Eq. (3.92).

3.3.4 Remarks on the new optimization problem

It has to be noticed that $[\tilde{\Psi}]$ is unique, and exactly keeps the same structure than $[\Psi(\nu^{\text{chaos}})]$. Indeed, let $[L^{\text{asym}}] = [A]^{-1}$ be the Cholesky decomposition matrix of the autocorrelation matrix $[R_{\mathcal{E}}]$, which is defined by Eq. (3.66). Hence, from Eq. (3.61), $[\Psi(\nu^{\text{chaos}})] = [L^{\text{asym}}]^{-1}[M]$, which has to be compared to $[\tilde{\Psi}] = [B]^{-1}[\Psi(\nu^{\text{chaos}})] = ([L]^{-1}[A]^{-1}) ([A][M]) = [L]^{-1}[M]$, where $[L]$ and $[L^{\text{asym}}]$ are two lower triangular matrices. Whereas $[L^{\text{asym}}]$ implies the asymptotic orthonormality, $[L]$ guarantees the numerical orthonormality. Moreover, from Eq. (3.92), the optimal PCE coefficients matrix $[y]$ is approximated as a product of two matrices:

$$[y] \approx [z_{\mathcal{O}_C}^M] [B]^{-1}. \quad (3.93)$$

For a fixed value of N , $[B]$ is strongly dependent on ν^{chaos} and $[\Psi(\nu^{\text{chaos}})]$. From Eq. (1.40), it also verifies the asymptotic property:

$$\lim_{\nu^{\text{chaos}} \rightarrow \infty} [B] = [I_N], \quad (3.94)$$

which implies that $[z_{\mathcal{O}_C}^M]$ converges towards $[y]$ if sufficiently high values of ν^{chaos} are considered. Hence, the less dependent on $[\Psi(\nu^{\text{chaos}})]$ the matrix $[z_{\mathcal{O}_C}^M]$ is, the more accurate the choice of ν^{chaos} is, and the better the PCE identification is.

If another $(N \times \nu^{\text{chaos},*})$ real matrix $[\Psi^*(\nu^{\text{chaos},*})]$ of independent realizations is considered, the matrices $[B^*]$ and $[\tilde{\Psi}^*] = [B^*]^{-1}[\Psi^*(\nu^{\text{chaos},*})]$ can be computed according to Eqs. (3.85) and (3.86). As it has previously been seen, $[\Psi^*(\nu^{\text{chaos},*})]$, $[\tilde{\Psi}]$ and $[\tilde{\Psi}^*]$ keep the same structure. The accuracy of $[z_{\mathcal{O}_C}^M]$ can thus be estimated by comparing $\mathcal{C}([C^{\text{exp}}(\nu)], [z_{\mathcal{O}_C}^M], [\Psi(\nu^{\text{chaos}})][B]^{-1})$ and $\mathcal{C}([C^{\text{exp}}(\nu)], [z_{\mathcal{O}_C}^M], [\Psi^*(\nu^{\text{chaos},*})][B^*]^{-1})$.

In particular, $\nu^{\text{chaos},*}$ and ν^{chaos} can be different. Finally, once the coefficient matrix $[z_{\mathcal{O}_C}^M]$ has been computed, the higher $\nu^{\text{chaos},*}$ is, the more accurate and general the validation is.

3.4 Application

In this section, we illustrate the efficiency of the methods proposed in the two former sections to identify in inverse the multidimensional distribution of a M -dimension random vector $\mathbf{C}$ characterized by a set of ν independent realizations. According to the notations of the former sections, these independent realizations are gathered in the $(M \times \nu)$ real matrix $[C^{\text{exp}}(\nu)]$. Three cases are therefore presented with respect to the values of M and ν :

- Case 1: $\nu = 1000 \gg M = 3$: first, a low dimension case with many available realizations is first introduced to underline the ability of the PCE method to identify in inverse complex and multidimensional distributions.
- Case 2: $\nu = 1000 \gg M = 50$: secondly, a high dimension case with many available realizations is addressed to illustrate the numerical convergence difficulties that arise when the size of the projection basis increases, and in what extent the proposed method allows us to overcome them.
- Case 3: $\nu = 100 < M = 150$: at last, we present a very high dimensional case with few available realizations. It will be shown that even in this case, the PCE method give very promising results.

In these three examples, another set of ν^{ref} independent realizations ($\nu^{\text{ref}} \gg \nu$) is used as a reference to validate the different modelings. Moreover, a distinction has to be made between the PDF modeling, achieved thanks to a PCE, and its estimation from PCE samples, computed thanks to nonparametric methods. In this context, let ν^{chaos} be the number of independent realizations used to carry out the PCE identification, and $\nu^{\text{chaos},*}$ the number of independent realizations of the identified PCE random vector, which will be used to draw graphical representations.

3.4.1 Application in low dimension

The objective of this section is to apply the whole PCE method to a $M = 3$ -dimension case. First, the statistical properties of the unknown random vector $\mathbf{C}$ are presented. Secondly, a convergence analysis is carried out in order to calculate the optimal truncation parameters

N_g and p of the PCE, $\mathbf{C}^{\text{chaos}}(N)$, of $\mathbf{C}$. Then, the PCE coefficients are identified from the ν independent realizations, $[\mathbf{C}^{\text{exp}}(\nu)]$, of $\mathbf{C}$. At last, the relevance of the PCE modeling is analyzed.

Generation of the random vector to identify. Let $[X]$ be a (3×6) real-valued random matrix whose coefficients are uniformly and independently chosen between -1 and 1, such that $\mathbf{C}$ is defined according to the notations of Section 3.3.1 as:

$$\mathbf{C} = [X] \mathcal{E}(\boldsymbol{\xi}^{\text{exp}}, 2), \quad (3.95)$$

where $\boldsymbol{\xi}^{\text{exp}} = (\xi_1^{\text{exp}}, \xi_2^{\text{exp}})$ is a normalized Gaussian random vector which components are independent. The components of $\mathbf{C}$ are however strongly dependent, and the PCE truncation parameters to be found back by the convergence analysis are $p^{\text{exp}} = 2$ and $N_g^{\text{exp}} = 2$.

Let $\{\boldsymbol{\xi}^{\text{exp}}(\theta_1), \dots, \boldsymbol{\xi}^{\text{exp}}(\theta_\nu)\}$ and $\{\boldsymbol{\xi}^{\text{exp}}(\theta_1), \dots, \boldsymbol{\xi}^{\text{exp}}(\theta_{\nu^{\text{ref}}})\}$ be ν and ν^{ref} independent realizations of the random vector $\boldsymbol{\xi}^{\text{exp}}$, such that the matrices of independent realizations $[\mathbf{C}^{\text{exp}}(\nu)]$ and $[\mathbf{C}^{\text{ref}}(\nu^{\text{ref}})]$ are given by:

$$[\mathbf{C}^{\text{exp}}(\nu)] = [X] [\mathcal{E}(\boldsymbol{\xi}^{\text{exp}}(\theta_1), 2) \quad \dots \quad \mathcal{E}(\boldsymbol{\xi}^{\text{exp}}(\theta_\nu), 2)], \quad (3.96)$$

$$[\mathbf{C}^{\text{ref}}(\nu^{\text{ref}})] = [X] [\mathcal{E}(\boldsymbol{\xi}^{\text{exp}}(\theta_1), 2) \quad \dots \quad \mathcal{E}(\boldsymbol{\xi}^{\text{exp}}(\theta_{\nu^{\text{ref}}}), 2)]. \quad (3.97)$$

Let $\{\hat{p}_C^{\text{ref},k}, 1 \leq k \leq 3\}$ be the Kernel smoothing estimations of the marginal PDFs of each component of $\mathbf{C}$, which are computed thanks to the ν^{ref} independent realizations of $\mathbf{C}$ gathered in $[\mathbf{C}^{\text{ref}}(\nu^{\text{ref}})]$. In this example, $\nu^{\text{ref}} = 2 \times 10^6 \gg \nu = 1000$. It is reminded that the PCE identification of $\mathbf{C}$ is only achieved thanks to the matrix of independent realizations $[\mathbf{C}^{\text{exp}}(\nu)]$, which is considered as the only available information. The PDFs $\{\hat{p}_C^{\text{ref},k}, 1 \leq k \leq 3\}$ are moreover supposed to build the marginal PDFs of the reference $\mathbf{C}$.

Identification of the PCE truncation parameters. Using the notations of Section 1.5.4, the boundary intervals BI_1 , BI_2 and BI_3 for which the convergence analysis is achieved, are chosen such that:

$$\forall 1 \leq k \leq 3, \text{BI}_k = \left\{ x \in \mathbb{R} \mid \hat{p}_C^{\text{ref},k}(x) \geq \frac{1}{\nu} \right\}. \quad (3.98)$$

Figure 3.2 displays the reference marginal PDFs of $\mathbf{C}$, as well as the marginal PDFs estimated from the ν independent realizations only, $\{\hat{p}_C^{\text{exp},k}, 1 \leq k \leq 3\}$. The $1/\nu$ tolerance is also plotted so that the boundary intervals can therefore be deduced from these graphs.

Figure 3.3 shows the values of $\text{err}(N_g, p)$, for nine pairs (N_g, p) in $\mathcal{Q}(3)$. On these graphs, the gradient break of $N \mapsto \text{err}(N)$ is observed at $N = 6$, which allows us to find back the initial solution $p^{\text{exp}} = 2$ and $N_g^{\text{exp}} = 2$. For this small dimension case, the optimal truncation parameters p and N_g given by the convergence analysis are equal to the parameters of the analytical reference PCE.

PCE Identification. The former convergence analysis leads us to the following PCE of $\mathbf{C}$:

$$\mathbf{C} \approx \mathbf{C}^{\text{chaos}}(6) = \sum_{j=1}^6 \mathbf{y}^j \Psi_j(\xi_1, \xi_2) = [\mathbf{y}] \boldsymbol{\Psi}(\xi_1, \xi_2), \quad (3.99)$$

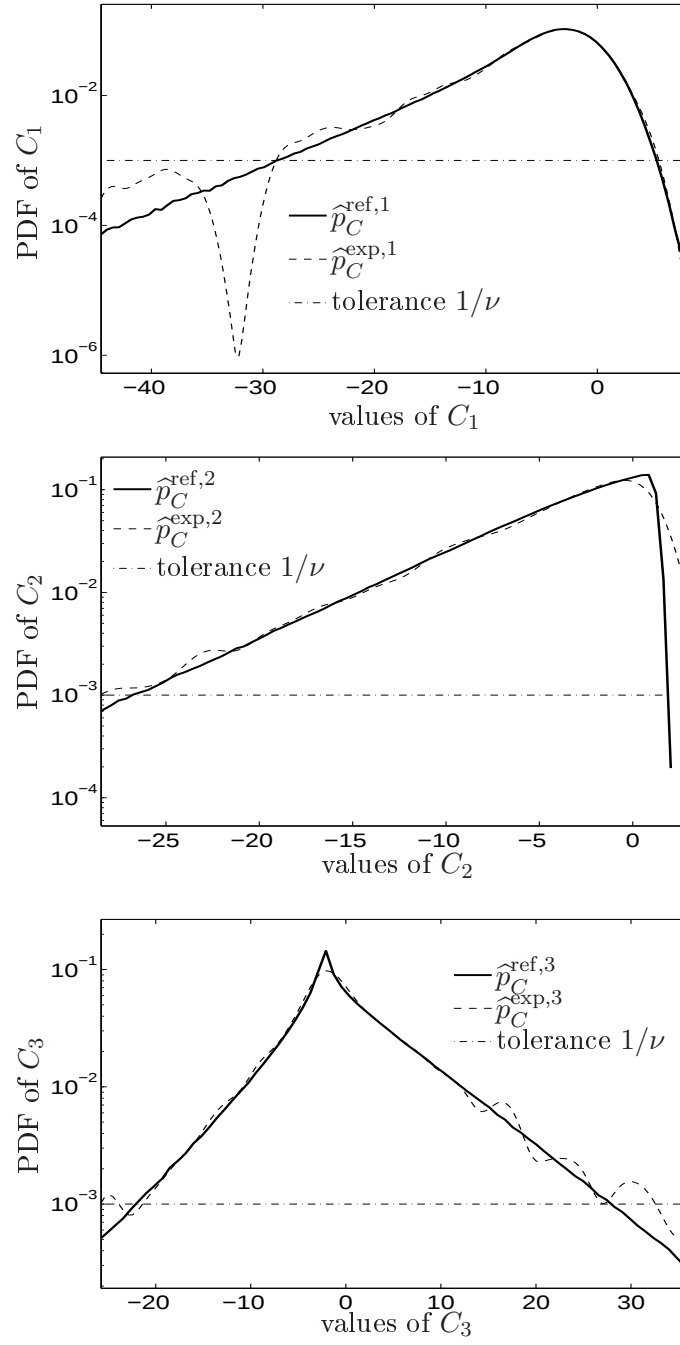


Figure 3.2: Graphs of the marginal PDFs of $\mathbf{C}$.

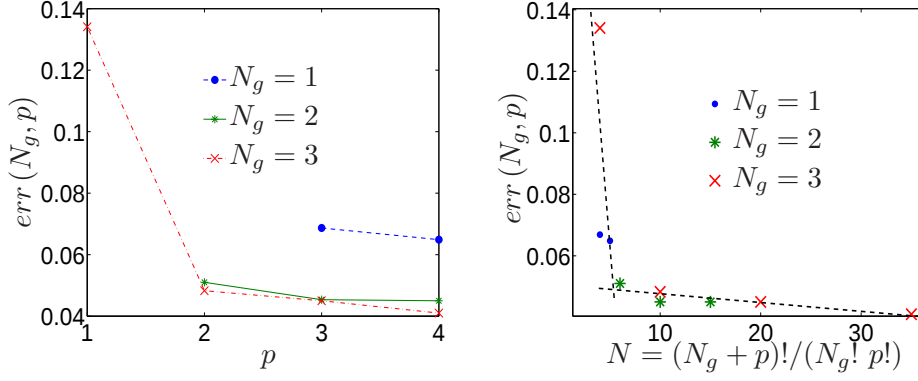


Figure 3.3: Convergence analysis of the PCE of $\mathbf{C}$.

where ξ_1 and ξ_2 are two independent normalized Gaussian random variables. We are now going to compare $[y^{\text{class}}]$ and $[y^{\text{new}}]$, where $[y^{\text{class}}]$ stems from the classical problem defined by Eq. (1.61), whereas $[y^{\text{new}}]$ comes from the maximization of the new formulation defined by Eq. (3.92). In this application, $\nu^{\text{chaos}} = 1000$, and the two PCE identifications have been computed thanks to the same indirect methods that are described in Section 3.2.4 to optimize the trials in the Stiefel manifold, with the same numerical cost ($Z = 10, Q = 10^4$). These values of Z and M have been chosen for the PCE error function $\text{err}(N_g, p)$ to be independent of them. Hence, for a new matrix of independent realizations, $[\Psi^*]$, of size $(6 \times \nu^{\text{chaos},*})$, independent realizations $[C^{\text{class}}(6)]$ and $[C^{\text{new}}(6)]$ of $\mathbf{C}^{\text{chaos}}(6)$ are deduced, with respect to the two optimization options:

$$[C^{\text{class}}(6)] = [y^{\text{class}}][\Psi^*], \quad (3.100)$$

$$[C^{\text{new}}(6)] = [y^{\text{new}}][\Psi^*]. \quad (3.101)$$

Let

$$[R_{CC}^{\text{exp}}] = \frac{1}{\nu} [C^{\text{exp}}(\nu)][C^{\text{exp}}(\nu)]^T, \quad (3.102)$$

$$[R_{CC}^{\text{ref}}] = \frac{1}{\nu^{\text{ref}}} [C^{\text{ref}}(\nu^{\text{ref}})][C^{\text{ref}}(\nu^{\text{ref}})]^T, \quad (3.103)$$

$$[R_{CC}^{\text{class}}] = \frac{1}{\nu^{\text{chaos},*}} [C^{\text{class}}(6)][C^{\text{class}}(6)]^T, \quad (3.104)$$

$$[R_{CC}^{\text{new}}] = \frac{1}{\nu^{\text{chaos},*}} [C^{\text{new}}(6)][C^{\text{new}}(6)]^T \quad (3.105)$$

be four estimations of the autocorrelation matrix $[R_{CC}]$ of $\mathbf{C}$. It is supposed that $[R_{CC}^{\text{ref}}]$ is the best approximation of $[R_{CC}]$ and will be considered as the reference. According to the Eqs. (3.54), (3.55) and (3.56), the autocorrelation errors $\text{err}^{1,\text{class}}$, $\text{err}^{2,\text{class}}$, $\text{err}^{3,\text{class}}$ and $\text{err}^{1,\text{new}}$, $\text{err}^{2,\text{new}}$, $\text{err}^{3,\text{new}}$ are then computed in each case. In figure 3.4, it can thus be verified that:

$$\forall \nu^{\text{chaos},*} \geq 6, \text{err}^{2,\text{new}}(\nu^{\text{chaos},*}) = \text{err}^{3,\text{new}}(\nu^{\text{chaos},*}, \nu) = 0, \quad (3.106)$$

$$\lim_{\nu^{\text{chaos},*} \rightarrow +\infty} \text{err}^{2,\text{class}}(\nu^{\text{chaos},*}) = \text{err}^{3,\text{class}}(\nu^{\text{chaos},*}, \nu) = 0. \quad (3.107)$$

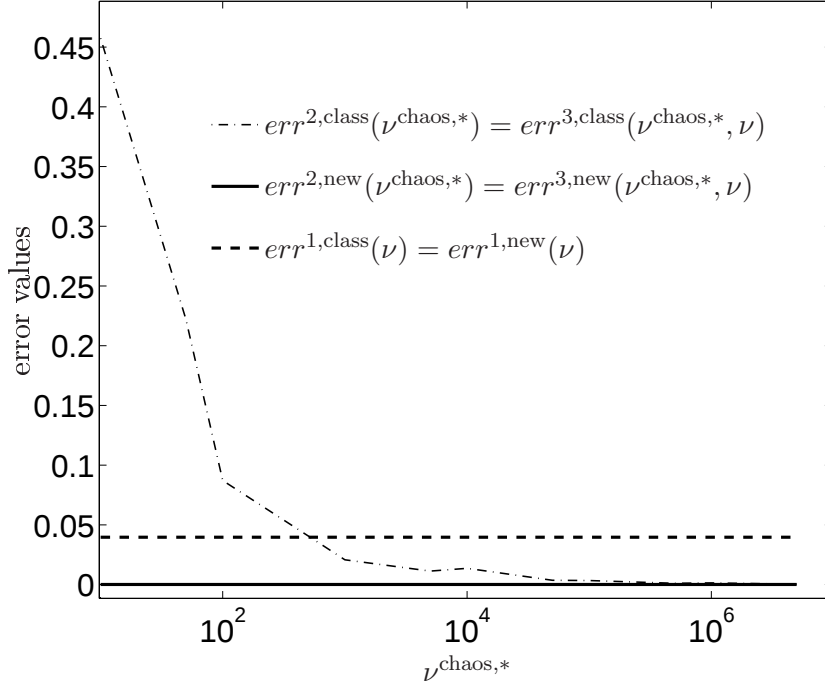


Figure 3.4: Convergence of the autocorrelation error functions with respect to $\nu^{\text{chaos},*}$.

In particular, for the value $\nu^{\text{chaos},*} = \nu^{\text{chaos}} = 1000$, it can be noticed that the values of $err^{2,\text{class}}(\nu^{\text{chaos},*})$ and $err^{3,\text{class}}(\nu^{\text{chaos},*}, \nu)$ are significant when compared to $err^{1,\text{class}}(\nu)$, which introduces an additive bias in the identification.

Figure 3.5 shows a comparison between the marginal PDFs $\hat{p}_C^{\text{ref},k}$, $\hat{p}_C^{\text{exp},k}$, $\hat{p}_C^{\text{class},k}$ and $\hat{p}_C^{\text{new},k}$, for $1 \leq k \leq 3$. These PDFs are estimated using Kernel smoothing on the independent realizations gathered in the matrices $[C^{\text{ref}}(\nu^{\text{ref}})]$, $[C^{\text{exp}}(\nu)]$, $[C^{\text{class}}(6)]$ and $[C^{\text{new}}(6)]$, respectively, with $\nu^{\text{chaos},*} = 10^6 \gg \nu^{\text{chaos}} = 1000$. First, from only $\nu = 1000$ independent realizations of $\mathbf{C}$, it can be seen that the marginal PDFs are well described by the PCE random vectors $\mathbf{C}^{\text{new}}(6)$ and $\mathbf{C}^{\text{class}}(6)$. In particular, the PDFs tails are very well characterized. The PCE method is therefore an extremely efficient tool to build arbitrary multidimensional PDFs. Secondly, it can be noticed that, for a same computational cost (Z, Q) , the new PCE identification formulation leads us to better results than the classical one. Finally, to still improve these PCE, more trials in $\mathcal{O}_C$ would be necessary to better characterize $[y^{\text{class}}]$ and $[y^{\text{new}}]$. In order to obtain a PCE that corresponds still more precisely to the reference random vector $\mathbf{C}$, an increase of ν , that is to say, more information about $\mathbf{C}$, would have been required.

Relevance of the PCE compared to Kernel Mixture and PASM. From adequacy tests, likelihood estimations and graphical representations, the idea of this section is to show the assets of the new PCE formulation when dealing with the identification of multidimensional distributions from a limited knowledge on the random vector of interest $\mathbf{C}$ compared to Kernel Mixture (KM) and Prior Algebraic Stochastic Modeling (PASM). In this prospect, two PDFs $\hat{p}_C(\mathbf{x})$ and $\hat{p}_C^{\text{PASM}}(\mathbf{x}, \mathbf{w})$ are built using a KM approach and a PASM method. The input data of these models are still the matrix of independent realizations $[C^{\text{exp}}(\nu)] = [\mathbf{C}(\theta_1) \cdots \mathbf{C}(\theta_\nu)]$. Once the KM, the PASM and the two PCE projection matrices, $[y^{\text{class}}]$ and $[y^{\text{new}}]$, are constructed,

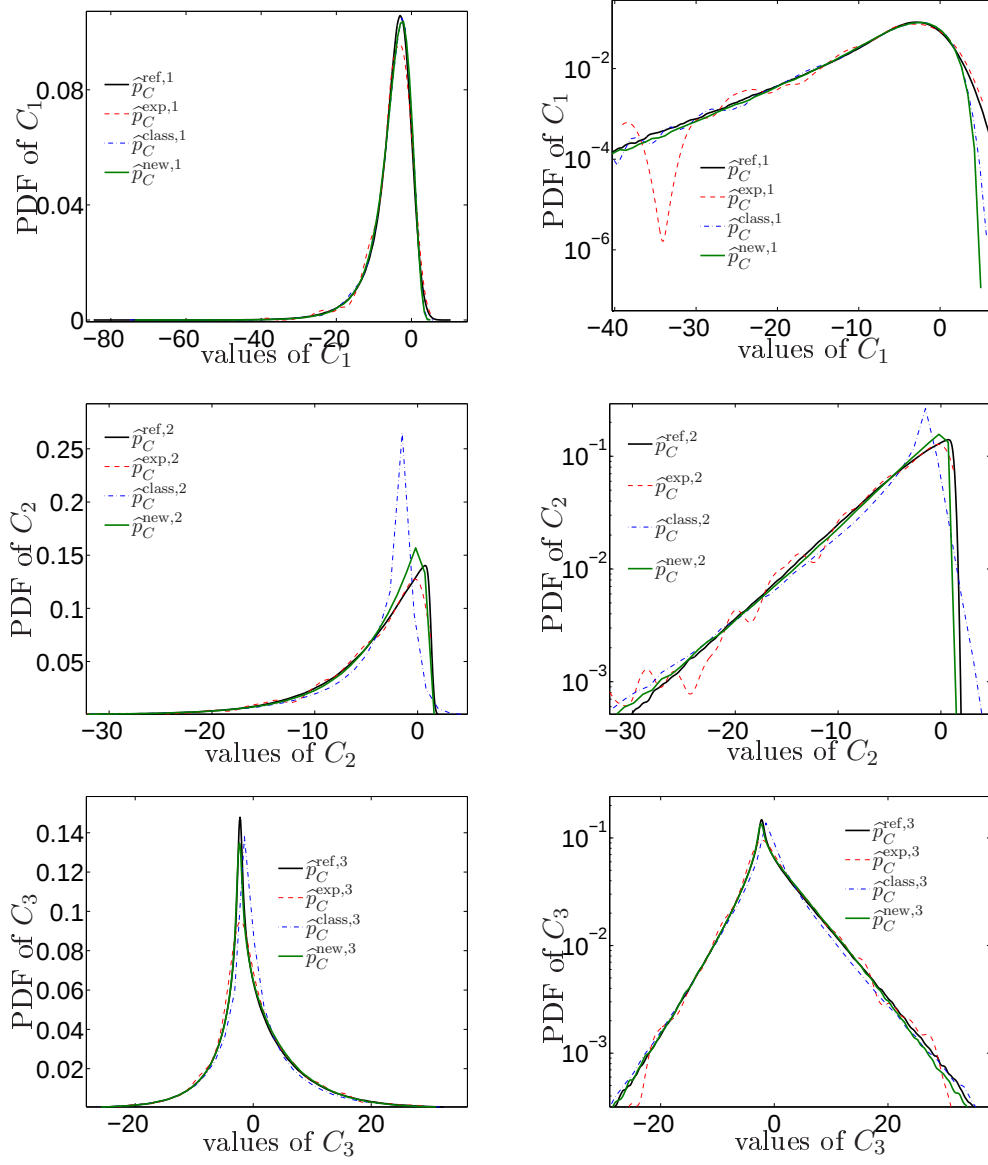


Figure 3.5: Comparison of the marginal PDFs of $\mathbf{C}$ and $\mathbf{C}^{\text{chaos}}(6)$.

P independent realizations are computed from the four distributions, from which comparisons to the reference solution are achieved. For this application, $P = 10^6$.

Construction of independent realizations

• Kernel Mixture.

Considering an independent Gaussian multidimensional Kernel, a nonparametric PDF $\hat{p}_C(\mathbf{x})$ is postulated as a sum of ν Gaussian PDFs $\{p_i, 1 \leq i \leq \nu\}$ to model $p_C(\mathbf{x})$:

$$\hat{p}_C(\mathbf{x}) = \sum_{i=1}^{\nu} \frac{1}{\nu} p_i(\mathbf{x}), \quad (3.108)$$

$$p_i(\mathbf{x}) = \prod_{m=1}^M \frac{1}{\sqrt{2\pi}h_m} \exp\left(-\frac{1}{2}\left(\frac{x_m - C_m^i}{h_m}\right)^2\right), \quad (3.109)$$

$$\mathbf{h} = \hat{\boldsymbol{\sigma}} \left(\frac{4}{(2+M)\nu} \right)^{1/(M+4)}, \quad (3.110)$$

where $\mathbf{x} \mapsto p_i(\mathbf{x})$ is the M -dimension multivariate Gaussian PDF, with mean value $\mathbf{C}^{(i)}$ and

covariance matrix $\begin{pmatrix} h_1 & 0 & \cdots & 0 \\ 0 & h_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & h_M \end{pmatrix}$, $\mathbf{h}$ is the multidimensional optimal Silverman band-

width, and $\hat{\sigma}_k$ is the empirical estimation of the standard deviation of each component C_k of $\mathbf{C}$. Let $\mathbf{C}^{\text{ker}}$ be the Kernel Mixture characterized by the PDF $\mathbf{x} \mapsto \hat{p}_C(\mathbf{x})$. The Q independent realizations $\{\mathbf{C}^{\text{ker},1}, \dots, \mathbf{C}^{\text{ker},Q}\}$ of $\mathbf{C}^{\text{ker}}$ are then computed and gathered in the matrix $[\mathbf{C}^{\text{ker}}]$.

• Prior Algebraic Stochastic Modeling.

From the ν independent realizations of $\mathbf{C}$, the M marginal cumulative distributions F_{C_m} of C_m , with $1 \leq m \leq M$, are estimated using a non parametric statistical method. In addition, a Gaussian copula $C_{\text{rank}}^{\text{gauss}}$ (see [2] for more details about the copula) based on the rank correlation is chosen (this type of copula has been chosen as it is the most commonly used in the PASM approaches):

$$C_{\text{rank}}^{\text{gauss}}(x_1, \dots, x_M) = \phi_{\text{rank}}^M(\phi^{-1}(x_1), \dots, \phi^{-1}(x_M)), \quad (3.111)$$

$$\phi_{\text{rank}}^M(\mathbf{u}) = \int_{-\infty}^{u_1} \cdots \int_{-\infty}^{u_M} \frac{1}{(2\pi)^{M/2} \sqrt{\det([R^{\text{rank}}])}} \exp\left(-\frac{1}{2}\mathbf{u}^T [R^{\text{rank}}] \mathbf{u}\right) du_1 \cdots du_M, \quad (3.112)$$

$$\phi(v) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^v \exp\left(-\frac{v^2}{2}\right) dv, \quad (3.113)$$

$$[R^{\text{rank}}]_{ij} = 2 \sin\left(\frac{\pi}{6} \rho_{ij}^S\right), \quad (3.114)$$

where ρ_{ij}^S is the Spearman correlation coefficient between C_i and C_j . Let $\mathbf{C}^{\text{cop}}$ be the random vector characterized by the copula $C_{\text{rank}}^{\text{gauss}}$ and the marginal cumulative distributions $\{F_{C_m}, 1 \leq m \leq M\}$. Q independent realizations of $\mathbf{C}^{\text{cop}}$ are thus gathered in the matrix $[\mathbf{C}^{\text{cop}}]$.

- **Polynomial chaos expansion.**

Finally, using the matrices $[y^{\text{class}}]$ and $[y^{\text{new}}]$ that have been previously defined, and a new $(6 \times P)$ real matrix $[\Psi(P)]$ of realizations, P independent realizations of $\mathbf{C}^{\text{class}}(6)$ and $\mathbf{C}^{\text{new}}(6)$ are gathered in the matrix $[\mathbf{C}^{\text{class}}] = [y^{\text{class}}][\Psi(P)]$ and $[\mathbf{C}^{\text{new}}] = [y^{\text{new}}][\Psi(P)]$.

Relevance of the PCE modeling when identifying multidimensional PDFs from a limited amount of independent realizations. Using the results of Parametric Statistics, this section assesses the relevance of the four methods to construct multidimensional PDFs. Three kinds of analysis are achieved: adequacy tests, 2D graphical representations, and multi-dimensional likelihood computations.

- **Adequacy tests.**

From the matrices of independent realizations $[\mathbf{C}^{\text{ker}}]$, $[\mathbf{C}^{\text{cop}}]$, $[\mathbf{C}^{\text{class}}]$ and $[\mathbf{C}^{\text{new}}]$, the estimations $\{\hat{F}_k^{\text{ker}}, 1 \leq k \leq M\}$, $\{\hat{F}_m^{\text{cop}}, 1 \leq m \leq M\}$, $\{\hat{F}_m^{\text{class}}, 1 \leq m \leq M\}$ and $\{\hat{F}_m^{\text{new}}, 1 \leq m \leq M\}$ of the cumulative distribution functions (CDF) of each components of $\mathbf{C}^{\text{ker}}$, $\mathbf{C}^{\text{cop}}$, $\mathbf{C}^{\text{class}}(6)$ and $\mathbf{C}^{\text{new}}(6)$ are respectively assessed. Let $\tilde{\mathbf{C}}^{(1)}, \dots, \tilde{\mathbf{C}}^{(M)}$ be the $(1 \times \nu)$ -dimension linear forms corresponding to the rows of $[\mathbf{C}^{\text{exp}}(\nu)]$. For $1 \leq m \leq M$, $\tilde{\mathbf{C}}^{(m)}$ gathers therefore the ν independent realizations of the component C_m of $\mathbf{C}$, which have been used to compute the statistical modelings. For $1 \leq m \leq M$, the Kolmogorov-Smirnov adequacy tests are then performed. For each component C_m of $\mathbf{C}$, the null distribution of the Kolmogorov-Smirnov statistics is computed under the null hypothesis that the ν independent realizations of $\tilde{\mathbf{C}}^{(m)}$ are drawn from the distribution of the chosen stochastic model. Table 3.1 gives the β -value for each stochastic model, which is defined as the probability of obtaining a test statistic at least as extreme as the one that was actually observed, assuming that the null hypothesis is true. Without surprise, this table allows us to verify that the modeling based on the Gaussian copula and the empirical PDFs of each components of $\mathbf{C}$ gives the best results. Moreover, with an error level of 5%, only the tests for the copula model and the PCE identification based on the new formulation are positive. The classical PCE and the Kernel mixture modelings are indeed less relevant to characterize the marginal PDFs of $\mathbf{C}$.

CDF	$\hat{F}_1^{\text{class}}$	$\hat{F}_1^{\text{new}}$	$\hat{F}_1^{\text{ker}}$	$\hat{F}_1^{\text{cop}}$
β -value	0.3779	0.6331	0.2142	0.9996
CDF	$\hat{F}_2^{\text{class}}$	$\hat{F}_2^{\text{new}}$	$\hat{F}_2^{\text{ker}}$	$\hat{F}_2^{\text{cop}}$
β -value	0.0000	0.0967	0.0000	0.4573
CDF	$\hat{F}_3^{\text{class}}$	$\hat{F}_3^{\text{new}}$	$\hat{F}_3^{\text{ker}}$	$\hat{F}_3^{\text{cop}}$
β -value	0.0000	0.8692	0.0411	0.9849

Table 3.1: Computation of the β -values corresponding to the different stochastic models.

- **Two-dimensions graphical analysis.**

From $[\mathbf{C}^{\text{ref}}(\nu^{\text{ref}})]$, $[\mathbf{C}^{\text{ker}}]$, $[\mathbf{C}^{\text{cop}}]$, $[\mathbf{C}^{\text{class}}]$ and $[\mathbf{C}^{\text{new}}]$, the estimations $\mathbf{x} \mapsto \hat{p}_{\mathbf{C}}^{\text{ref}}(\mathbf{x})$, $\mathbf{x} \mapsto \hat{p}_{\mathbf{C}}^{\text{ker}}(\mathbf{x})$, $\mathbf{x} \mapsto \hat{p}_{\mathbf{C}}^{\text{cop}}(\mathbf{x})$, $\mathbf{x} \mapsto \hat{p}_{\mathbf{C}}^{\text{class}}(\mathbf{x})$ and $\mathbf{x} \mapsto \hat{p}_{\mathbf{C}}^{\text{new}}(\mathbf{x})$ of the multidimensional PDF of $\mathbf{C}$, $\mathbf{C}^{\text{ker}}$, $\mathbf{C}^{\text{cop}}$, $\mathbf{C}^{\text{class}}(6)$, $\mathbf{C}^{\text{new}}(6)$ are respectively computed using the non parametric statistical

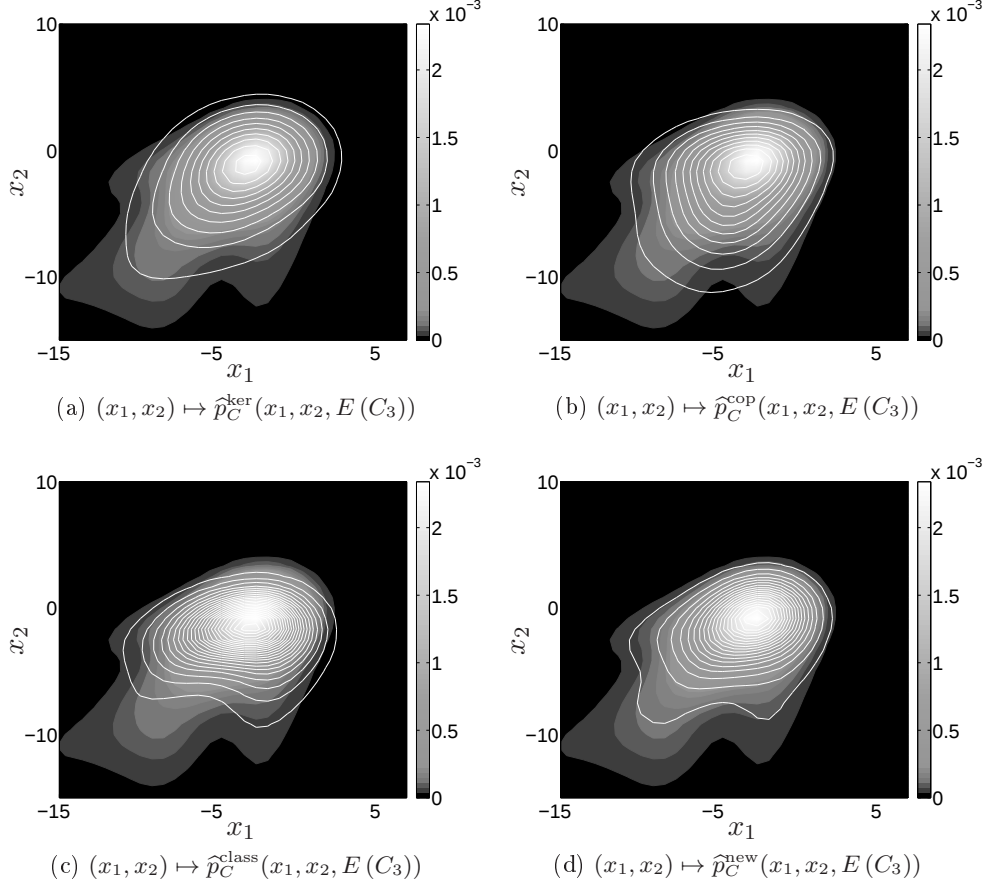


Figure 3.6: Comparison of 2D contours plots in the plane $[x_3 = E(C_3)]$.

estimation defined by Eq. (1.55). Projections of these functions are presented in Figures 3.6, 3.7 and 3.8. In each figure, the surface plot characterizes the reference PDF (based on the $\nu^{\text{ref}} = 2 \times 10^6$ independent realizations), and the contour plot refers to isovalues of the projected PDF of interest. It can therefore be seen that the new formulation of the PCE gives very good results in identifying multidimensional PDFs. In addition, in this example, the Kernel mixture model is more adapted than the copula based model to characterize the multidimensional PDFs.

• Likelihood estimations.

From Eq. (1.48), the multidimensional log-likelihood functions $\mathcal{L}_{C^{\text{ker}}}([C^{\text{exp}}(\nu)])$, $\mathcal{L}_{C^{\text{cop}}}([C^{\text{exp}}(\nu)])$, $\mathcal{L}_{C^{\text{class}}}([C^{\text{exp}}(\nu)])$ and $\mathcal{L}_{C^{\text{new}}}([C^{\text{exp}}(\nu)])$ are estimated from the realizations matrices $[C^{\text{exp}}(\nu)]$, $[C^{\text{ker}}]$, $[C^{\text{cop}}]$, $[C^{\text{class}}]$ and $[C^{\text{new}}]$, in order to evaluate the multidimensional relevance of the different stochastic models. In the same manner, $[C^{\text{ref}}(\nu^{\text{ref}})]_{1000}$ is defined as the 1000 first columns of $[C^{\text{ref}}(\nu^{\text{ref}})]$, and the log-likelihood functions $\mathcal{L}_{C^{\text{ker}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$, $\mathcal{L}_{C^{\text{cop}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$, $\mathcal{L}_{C^{\text{class}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$ and $\mathcal{L}_{C^{\text{new}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$ are computed. These values are gathered in Table 3.2. It can thus be verified that the new formulation of the PCE identification gives the best results when considering the maximization of the log-likelihood.

As a conclusion for this example, in low dimension, it can be seen that the new formulation of the PCE identification is very relevant when trying to identify multidimensional distributions

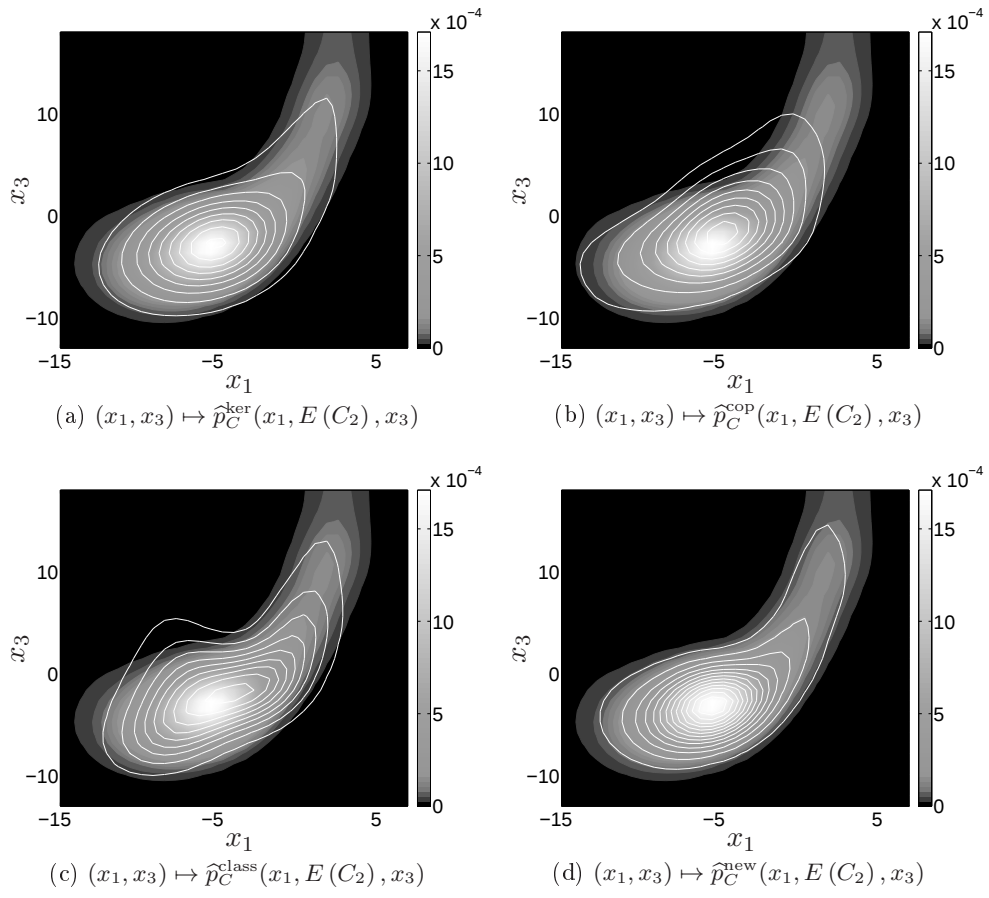


Figure 3.7: Comparison of 2D contours plots in the plane $[x_2 = E(C_2)]$.

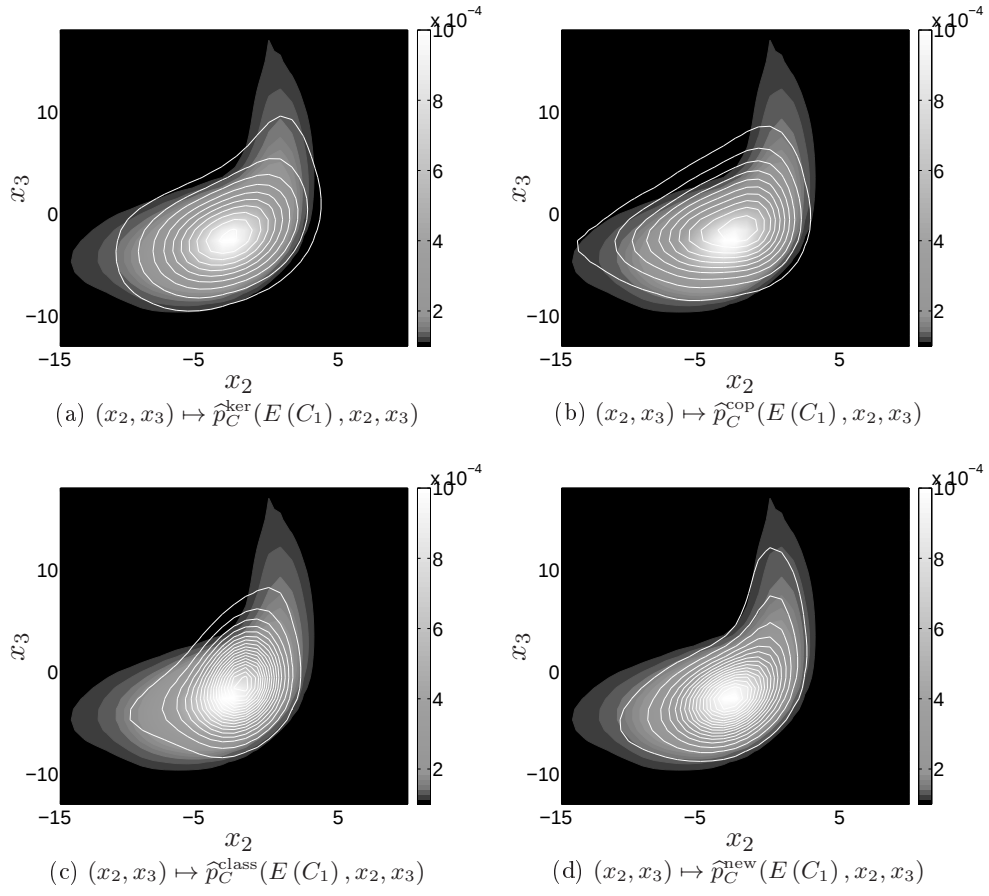


Figure 3.8: Comparison of 2D contours plots in the plane $[x_1 = E(C_1)]$.

$\mathcal{L}_{\mathcal{C}^{\text{ker}}}([C^{\text{exp}}(\nu)])$	$\mathcal{L}_{\mathcal{C}^{\text{cop}}}([C^{\text{exp}}(\nu)])$	$\mathcal{L}_{\mathcal{C}^{\text{class}}}([C^{\text{exp}}(\nu)])$	$\mathcal{L}_{\mathcal{C}^{\text{new}}}([C^{\text{exp}}(\nu)])$
$-8.0712.10^3$	$-8.7530.10^3$	$-8.1844.10^3$	$-7.8624.10^3$
$\mathcal{L}_{\mathcal{C}^{\text{ker}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$	$\mathcal{L}_{\mathcal{C}^{\text{cop}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$	$\mathcal{L}_{\mathcal{C}^{\text{class}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$	$\mathcal{L}_{\mathcal{C}^{\text{new}}}([C^{\text{ref}}(\nu^{\text{ref}})]_{1000})$
$-8.1933.10^3$	$-8.5535.10^3$	$-8.1797.10^3$	$-7.8457.10^3$

Table 3.2: Computation of the multidimensional log-likelihood corresponding to the different stochastic models.

from a limited number of measurements. Indeed, it allows us to build multidimensional distributions that are still relevant for experimental data that have not been used in the identification process.

3.4.2 Application in high dimension

The idea of this second application is to underline the capability of the new PCE formulation to carry out convergence analysis in high dimension. Indeed, as it has been shown in Section 3.3.1, for a given value of ν^{chaos} , when the size N of the polynomial basis increases, and more specially when the maximum degree p of the polynomial basis becomes high, the difference $\frac{1}{\nu^{\text{chaos}}}[\Psi][\Psi]^T - [I_N]$ introduces a significant numerical bias which perturbs the classical PCE identification. In opposite, the new PCE formulation, which avoids computational autocorrelation errors, allows the numerical algorithms to be much more stable and to give more relevant results.

Generation of a high dimension random vector Using the same notations than in Section 3.4.1, let $[X^{\text{HD}}]$ be a $(M \times N)$ real matrix whose entries are randomly generated, such that random vector $\mathbf{C}$ is given by:

$$\mathbf{C} = [X^{\text{HD}}]\Psi(\xi^{\text{exp}}), \quad (3.115)$$

$$\xi^{\text{exp}} = (\xi_1^{\text{exp}}, \xi_2^{\text{exp}}, \dots, \xi_{N_g}^{\text{exp}}), \quad (3.116)$$

where $\{\xi_\ell^{\text{exp}}, 1 \leq \ell \leq N_g\}$ is a set of N_g independent normalized Gaussian random variables. As in Section 3.4.1, we define a $(M \times \nu)$ real matrix $[C^{\text{exp}}(\nu)]$, which gathers ν independent realizations of $\mathbf{C}$:

$$[C^{\text{exp}}(\nu)] = [X^{\text{HD}}][\Psi^{\text{exp}}], \quad (3.117)$$

$$[\Psi^{\text{exp}}] = [\Psi(\xi^{\text{exp}}(\theta_1)) \quad \dots \quad \Psi(\xi^{\text{exp}}(\theta_\nu))]. \quad (3.118)$$

The components of the random vector $\mathbf{C}$ are again strongly dependent. As a numerical illustration, it is supposed that $\nu = 1000$, $p^{\text{exp}} = 9$, $N_g^{\text{exp}} = 3$, $N = (9 + 3)!/(9! 3!) = 220$, $M = 50$. A high value of p^{exp} is deliberately chosen, in order to emphasize the difficulties of the classical PCE formulation to carry out convergence analysis in high dimension. Nevertheless, this high value for the maximal polynomial order implies an ill-conditioning of $[\Psi^{\text{exp}}]$, such that $\mathbf{C}$ can have very high values.

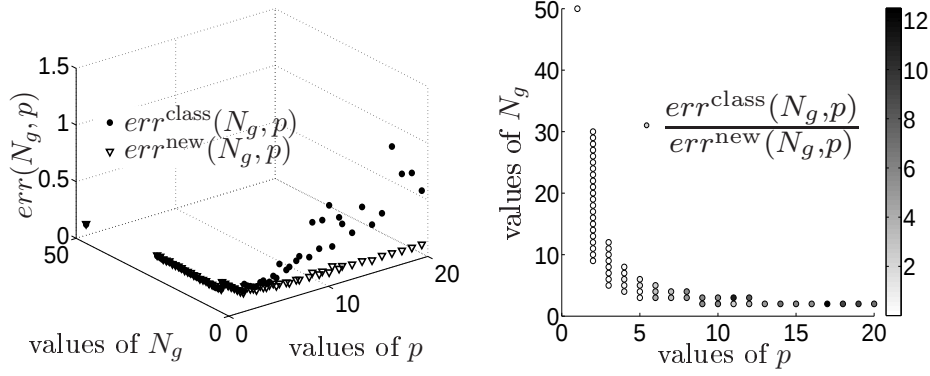


Figure 3.9: Comparison of the results for the convergence analysis of the two PCE identification formulations.

Identification of the PCE truncation parameters According to Eq. (1.34), the truncated PCE, $\mathbf{C}^{\text{chaos}}(N)$, of $\mathbf{C}$ is given by:

$$\mathbf{C}^{\text{chaos}}(N) = [y]\Psi(\xi). \quad (3.119)$$

Eq. (1.69) implies that the number N_y of elements of $[y]$ has to be higher than $M(M+1)/2$. When M is large, this leads us to the identification of thousands of coefficients. However, as it has been said in Section 1.5.4, the higher N_y is, the less precise is the PCE identification, for a given computational cost (Z, Q) . This motivates the definition of a new set $\tilde{\mathcal{Q}}(p^{\max}, N^{\max})$, such that the optimal values p^{opt} and N_g^{opt} are given by:

$$\tilde{\mathcal{Q}}(p^{\max}, N^{\max}) = \{(p, N_g), N_g \leq M, p \leq p^{\max}, (N_g + p)! / (N_g! p!) \leq N^{\max}\}, \quad (3.120)$$

$$(p^{\text{opt}}, N_g^{\text{opt}}) = \arg \min_{(p, N_g) \in \tilde{\mathcal{Q}}(p^{\max}, N^{\max})} \text{err}(N_g, p), \quad (3.121)$$

where error $\text{err}(N_g, p)$ is defined by Eq. (1.66), and is computed with respect to a fixed choice for the computational cost (Z, Q) . For a fixed value $\nu^{\text{chaos}} = 1000$, the detrimental influence of the autocorrelation errors err^2 and err^3 of Eqs. (3.55) and (3.56) can then be noticed in Figure 3.9, when high values of N (and more specially high values of p) are considered. The error functions $\text{err}^{\text{class}}(N_g, p)$ and $\text{err}^{\text{new}}(N_g, p)$ correspond, respectively, to the classical formulation and the new formulation of the PCE identification. It can be seen that for $p \geq 8$, the ratio $\text{err}^{\text{class}}(N_g, p) / \text{err}^{\text{new}}(N_g, p)$ becomes greater than five, whereas the two methodologies are globally similar for low values of p . Hence, the accuracy of the classical method seems to be limited to low values of p and is therefore less relevant for convergence analysis which handle high polynomial orders. At last, the five lowest values of the numerical assessments of $\text{err}^{\text{new}}(N_g, p)$ are gathered in Table 3.3. It can be seen that the new formulation allows finding back the couple $(p^{\text{exp}}, N_g^{\text{exp}})$ as the minimum of the error function. Nevertheless, keeping in mind that the lowest N is, the easiest the identification is, this result also shows that using the couple $(p, N_g) = (11, 2)$ could be interesting.

PCE Identification From the ν independent realizations of $\mathbf{C}$, a PCE identification using the new formulation can be computed for the truncation parameters $p = 9$ and $N_g = 3$, which

couples (p, N_g)	(11,2)	(9,3)	(7,4)	(6,5)	(2,27)
values of N	78	220	330	462	406
$err^{new}(N_g, p)$	0.06104	0.06005	0.06228	0.06301	0.06521

Table 3.3: Lowest values of $err^{new}(N_g, p)$ with respect to (p, N_g) .

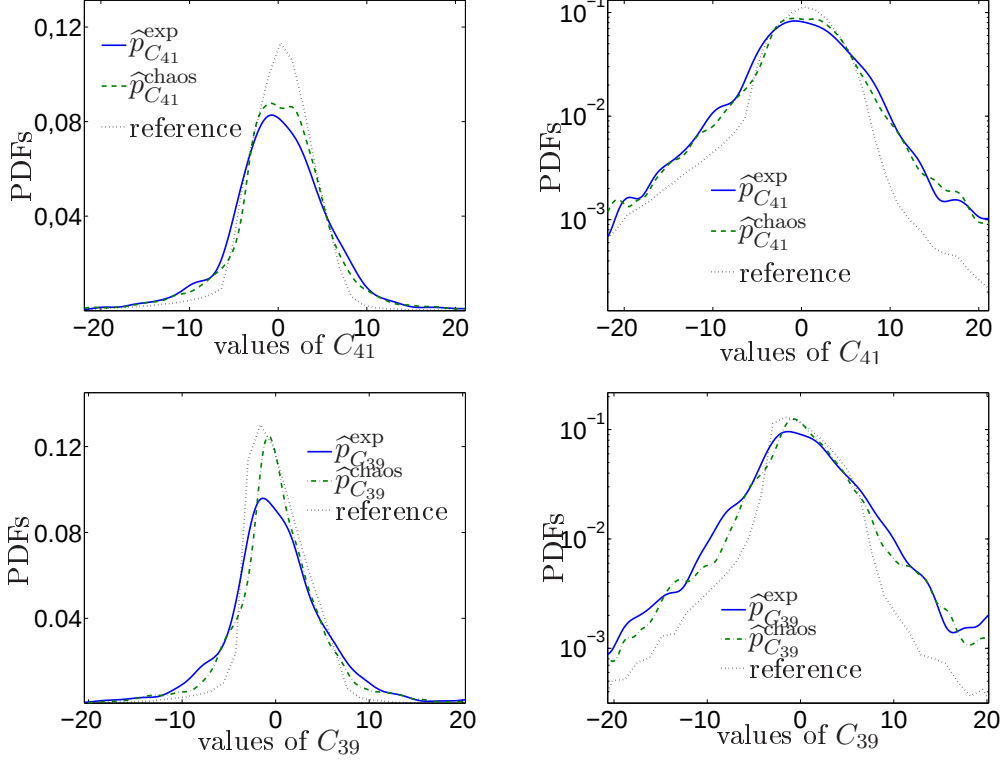


Figure 3.10: Graphs of the estimated marginal PDFs for two particular components of $\mathbf{C}$.

correspond to $N = 220$. The results of the numerical identification with a computational cost of $(Z = 10, Q = 1000)$ are given in Figure 3.10. The values of $(Z = 10, Q = 1000)$ have once again been chosen for the PCE error function $err(N_g, p)$ to be independent on them. In this figure, the marginal PDFs $\hat{p}_{C_{41}}^{chaos}$ and $\hat{p}_{C_{39}}^{chaos}$ of $C_{41}^{new}(220)$ and $C_{39}^{new}(220)$ are compared to the experimental estimations $\hat{p}_{C_{41}}^{exp}$ and $\hat{p}_{C_{39}}^{exp}$ of the components C_{41} and C_{39} , respectively. The values $C_{41}^{new}(220)$ and $C_{39}^{new}(220)$ correspond to the minimum and to the maximum values of the unidimensional error function $err_k(3, 11)$, for $1 \leq k \leq 50$, which is defined by Eq. (1.65). In order to evaluate the distance between these estimations and the true marginal PDFs of $\mathbf{C}$, the marginal PDFs estimated by the non parametric statistical Kernel method, with $\nu^{ref} = 2 \times 10^5$ independent realizations of C_{41} and C_{39} , are added to the figures. These PDFs are considered as the reference. These figures therefore emphasize that the new PCE identification method allows building a stochastic model of the distribution of $\mathbf{C}$ that suits the experimental marginal PDFs.

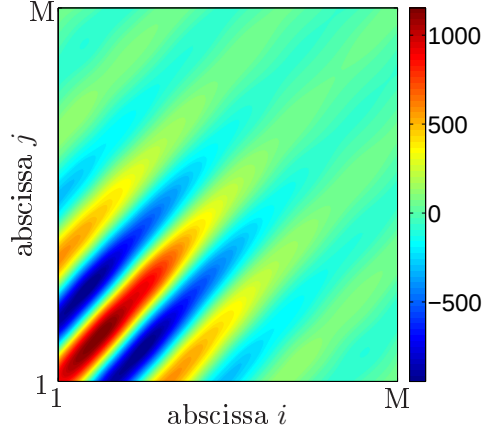


Figure 3.11: Representation of $E[\mathbf{C} \otimes \mathbf{C}]_{i,j}$, $1 \leq i, j \leq M$.

3.4.3 Application in very high dimension when the available information is very limited

Generation of the random vector to identify As in the previous applications, we define a centered M -dimension random vector $\mathbf{C}$ from its PCE formulation, such that:

$$\mathbf{C} = [D][X^{HD}]\Psi(\boldsymbol{\xi}^{\text{exp}}), \quad (3.122)$$

where $[X^{HD}]$ is a $(M \times N)$ real matrix whose entries are randomly generated under the constraint $[X^{HD}][X^{HD}]^T = [I_M]$, $\{\xi_1^{\text{exp}}, \dots, \xi_{N_g}^{\text{exp}}\}$ is a set of N_g independent normalized Gaussian random variables, $\{\Psi_1(\boldsymbol{\xi}^{\text{exp}}), \dots, \Psi_N(\boldsymbol{\xi}^{\text{exp}})\}$ gathers N polynomial function of $\boldsymbol{\xi}^{\text{exp}}$ that are statistically orthonormal and for which maximal order is p , and $[D]$ is a $(M \times M)$ real orthogonal matrix, such that:

$$E[\mathbf{C} \otimes \mathbf{C}] = [D][D]^T. \quad (3.123)$$

A representation of the chosen matrix $E[\mathbf{C} \otimes \mathbf{C}]$ can be found in Figure 3.11. The components of $\mathbf{C}$ are thus chosen on purpose very correlated.

Two sets of independent realizations of $\mathbf{C}$ are then introduced:

- $\mathcal{S}^{\text{exp}} = \{\mathbf{C}(\theta_1), \dots, \mathbf{C}(\theta_\nu)\}$ corresponds to the available information about the identification in inverse of the multidimensional distribution of $\mathbf{C}$,
- $\mathcal{S}^{\text{ref}} = \{\mathbf{C}(\Theta_1), \dots, \mathbf{C}(\Theta_{\nu^{\text{ref}}})\}$ is a reference set, which will only be used to evaluate the relevance of the identification process.

In this section, we choose $M = 150 > \nu = 100$, such that the information about $\mathbf{C}$ is very limited compared to its dimension, and the rank of the empirical estimator of the covariance of $\mathbf{C}$, $[\hat{R}_{CC}(\nu)] = 1/\nu[C^{\text{exp}}(\nu)][C^{\text{exp}}(\nu)]^T$ is inferior to ν . In addition, in the following, $N_g = 5$, $p = 5$, such that $N = 252$, and $\nu^{\text{ref}} = 4,000$.

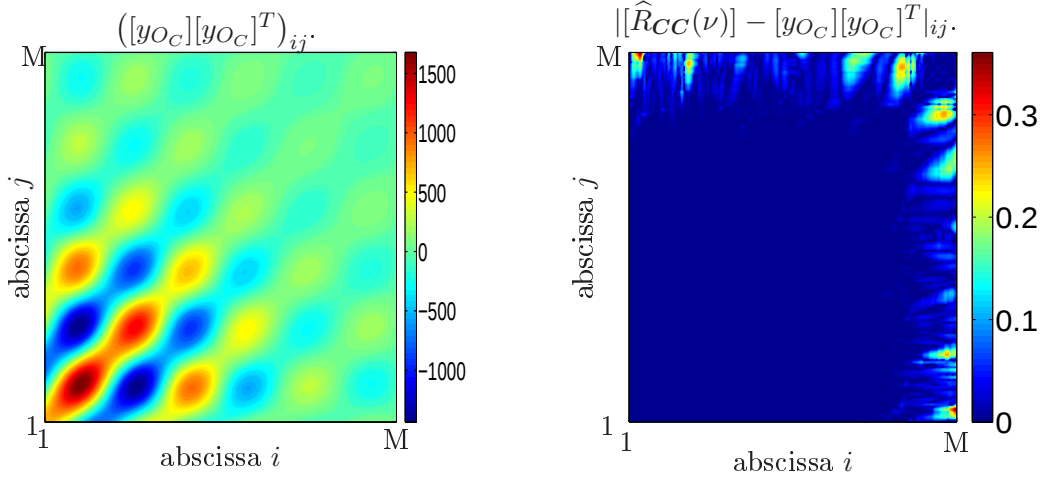


Figure 3.12: Comparison between the covariance of $\mathbf{C}^{\text{chaos}}(N)$ and the empirical estimator of the covariance of $\mathbf{C}$.

PCE identification For ξ a N_g -dimension random vector whose component are independent and normally distributed, it has been shown in Section 1.5 that the PCE projection matrix, $[y]$, of $\mathbf{C}$, such that $\mathbf{C} \approx \mathbf{C}^{\text{chaos}}(N) = [y]\Psi(\xi)$, has to be searched to maximize the likelihood of $\mathbf{C}^{\text{chaos}}(N)$ at the experimental points gathered in $\mathcal{S}^{\text{exp}}$, under the orthogonality constraints:

$$[y][y]^T = E[\mathbf{C} \otimes \mathbf{C}]. \quad (3.124)$$

As only a limited set of $\nu < M$ realizations of $\mathbf{C}$ are available, this constraint cannot be exactly verified. From the iterative algorithm based on the development achieved in Section 3.2.5, matrices $[y^*]$ that verify $[y^*][y^*]^T \approx [\hat{R}_{CC}(\nu)]$ can however be generated, such that the PCE matrix $[y]$ can be identified from the limited available information about $\mathbf{C}$. Let $[y_{O_C}]$ be the approximation of $[y]$ corresponding to a computational cost given by $(Z = 10, Q = 10^3)$, such that $\mathbf{C}^{\text{chaos}}(N) = [y_{O_C}]\Psi(\xi)$. A comparison between $[\hat{R}_{CC}(\nu)]$ and the covariance matrix of $\mathbf{C}^{\text{chaos}}(N)$, $[y_{O_C}][y_{O_C}]^T$, can thus be seen in Figure 3.12. It can be verified that the algorithm proposed in Section 3.2.5 allows us to make $[y_{O_C}][y_{O_C}]^T$ be equal to $[\hat{R}_{CC}(\nu)]$ almost everywhere but in restricted zones where the components of $[\hat{R}_{CC}(\nu)]$ are very low.

From $\nu^{\text{ref}} = 4,000$ new realizations of ξ , ν^{ref} independent realizations of $\mathbf{C}^{\text{chaos}}(N)$ can be deduced, and are gathered in the set $\mathcal{S}^{\text{chaos}} = \{\mathbf{C}^{\text{chaos}}(N, \theta_1), \dots, \mathbf{C}^{\text{chaos}}(N, \theta_{\nu^{\text{ref}}})\}$. Hence, let $\hat{p}_{C_m}^{\text{exp}}$, $\hat{p}_{C_m}^{\text{ref}}$ and $\hat{p}_{C_m}^{\text{chaos}}$ be the empirical estimations of the PDFs of C_m and $C_m^{\text{chaos}}(N)$ that have been computed from the sets $\mathcal{S}^{\text{exp}}$, $\mathcal{S}^{\text{ref}}$ and $\mathcal{S}^{\text{chaos}}$ respectively. Three particular components of C_m and $C_m^{\text{chaos}}(N)$ are then compared in figure 3.13. It can be seen that the PDFs $\hat{p}_{C_m}^{\text{chaos}}$ are very close to the reference PDFs $\hat{p}_{C_m}^{\text{ref}}$, and even closer than $\hat{p}_{C_m}^{\text{exp}}$, such that even when dealing with very high dimensional problems with very limited available information ($M = 150 > \nu = 100$), the proposed PCE method appears to give very promising results.

In order to emphasize that not only the marginal PDFs of $\mathbf{C}$ are well characterized but its whole distribution, let $\{b_m, 1 \leq m \leq M\}$ be a set of M orthonormal functions that are defined on $\Omega = [0, 1]$, and let $\mathcal{J}_1$, $\mathcal{J}_2$ and $\mathcal{J}_3$ be three random indices permutations such that:

$$\mathcal{J}_p = \{j_1^{(p)}, \dots, j_M^{(p)}\} \subset \{1, \dots, M\}, \quad j_1^{(p)} \neq \dots \neq j_M^{(p)}, \quad 1 \leq p \leq 3. \quad (3.125)$$

This allows us to define three triplets of random fields,

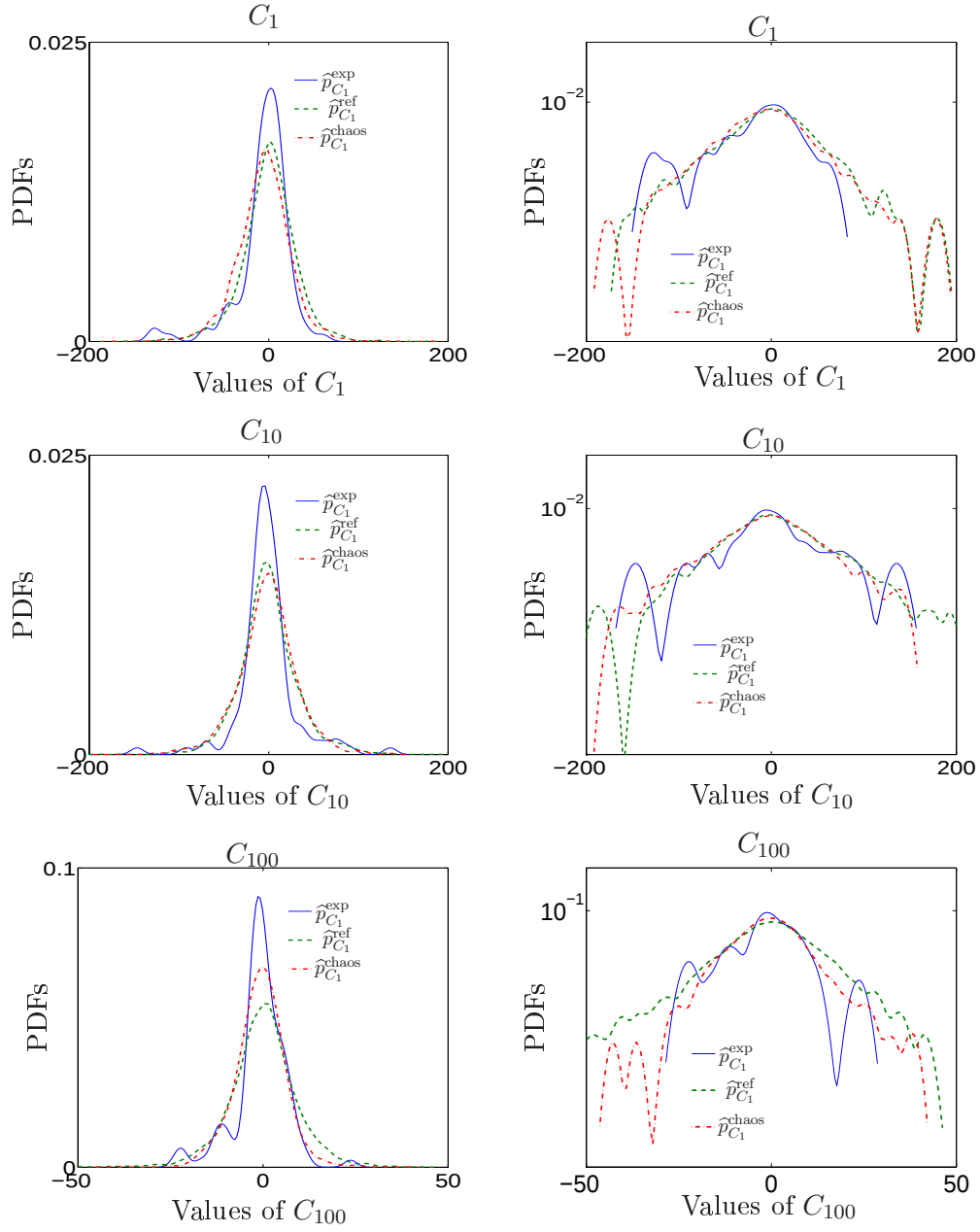


Figure 3.13: Comparison of the PDFs of three particular components of $\mathbf{C}$ and $\mathbf{C}^{\text{chaos}}(N)$.

$$(X_1, X_1^{\text{chaos}}, X_1^{\text{Gauss}}), (X_2, X_2^{\text{chaos}}, X_2^{\text{Gauss}}), (X_3, X_3^{\text{chaos}}, X_3^{\text{Gauss}}),$$

such that for $1 \leq p \leq 3$:

$$\begin{aligned} X_p &= \sum_{m=1}^M C_m b_{j_m^{(p)}} \\ X_p^{\text{chaos}} &= \sum_{m=1}^M C_m^{\text{chaos}}(N) b_{j_m^{(p)}} \\ X_p^{\text{Gauss}} &= \sum_{m=1}^M C_m^{\text{Gauss}} b_{j_m^{(p)}}, \end{aligned} \tag{3.126}$$

where $\mathbf{C}^{\text{Gauss}} = (C_1^{\text{Gauss}}, \dots, C_M^{\text{Gauss}})$ is a Gaussian vector for which mean and correlation are equal to the ones of $\mathbf{C}$. Therefore, the statistical properties of these random fields strongly depend on the dependencies between the components of their projection coefficients. Comparing the statistical properties of these random fields is thus a method to compare the global relevance of the characterization of the multidimensional distribution of $\mathbf{C}$.

Thanks to $\nu^{\text{ref}} = 4,000$ realizations of $\mathbf{C}$, $\mathbf{C}^{\text{chaos}}(N)$ and $\mathbf{C}^{\text{Gauss}}$, we then have access to ν^{ref} independent realizations of X_p , X_p^{chaos} and X_p^{Gauss} . In order to compare the statistical information that is included in these realizations, we denote by $N_{\text{up}}(X_p(\theta_q), u)$, $N_{\text{up}}(X_p^{\text{chaos}}(\theta_q), u)$, $N_{\text{up}}(X_p^{\text{Gauss}}(\theta_q), u)$ the numbers of upcrossings (see [16] for more details about the upcrossings) of the level u by the q^{th} realization, $X_i(\theta_q)$, $X_i^{\text{chaos}}(\theta_q)$, $X_i^{\text{Gauss}}(\theta_q)$, of X_p , X_p^{chaos} and X_p^{Gauss} respectively over the length $[0, 1]$. At last, we define $\mathcal{D}_i$, $1 \leq i \leq 10$ the domains such that for each level u , $\mathcal{D}_i$ gathers $i/10$ of the values of $\{N_{\text{up}}(X_p(\theta_1), u), \dots, N_{\text{up}}(X_p(\theta_{\nu^{\text{ref}}}), u)\}$. These domains for the three considered permutations are thus compared to contour plots that corresponds to the equivalent domains for random fields X_p^{chaos} and X_p^{Gauss} in Figure 3.14. The very good agreement between the domains of X_p and X_p^{chaos} , whereas the domains of X_p and X_p^{Gauss} do not match correctly, is an other illustration of the relevance of the PCE method to identify in inverse from a finite set of independent realizations the multidimensional distributions of an unknown random vector, even when the components of this vector are strongly correlated and very dependent.

3.5 Conclusions

In this chapter, it has been shown in what extent the PCE method gives very promising results when trying to identify in inverse the multidimensional distribution of high dimensional random vectors. For this method to numerically give relevant results, two adaptations of the classical formulation presented in Chapter 1 have been emphasized. First, iterative algorithms have been described to optimize the trials of random matrices under orthogonality constraints. Secondly, a method to numerically stabilize the matrix of realizations of the statistical polynomial basis has been introduced. The interest of these two adaptations has then been underlined on three applications based on simulated data. Finally, the method proposed allows making the PCE range reachable for many engineering applications with many degrees of freedom.

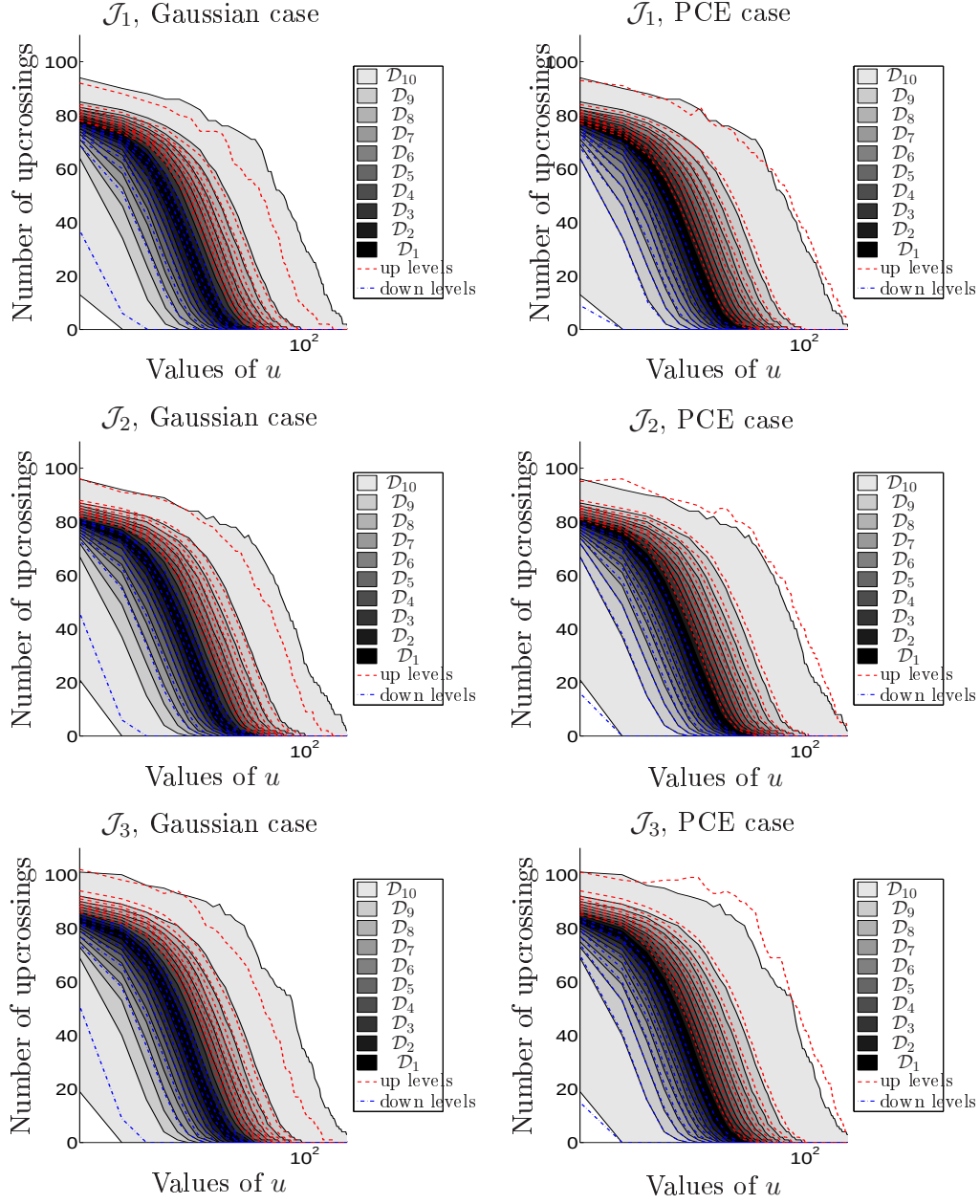


Figure 3.14: Comparison of the distributions of the upcrossings corresponding to the realizations of X_p , X_p^{chaos} and X_p^{Gauss} .

Chapter 4

Karhunen-Loève expansion revisited for vector-valued random fields

4.1 Introduction

As presented in Chapters 1 and 2, the Karhunen-Loève (KL) expansion has been used in many scientific fields to efficiently reduce the statistical dimension of random fields. This expansion can then be coupled to a polynomial chaos expansion (PCE), for which formulation has been presented in detail in Chapters 1 and 3, to completely characterize the distribution of random fields [50, 51, 22, 62, 78, 79, 34, 66, 80, 69, 68, 81]. In particular, it has been shown in Chapter 1 that the importance of the KL expansion stems from its optimality in the sense that, due to the orthogonal projection theorem in Hilbert spaces, it minimizes the total mean-squared error. In other words, for any integer M and any multivariate random field $\mathbf{X}$, it can be extracted from the KL basis associated with $\mathbf{X}$ the M -dimension family that minimizes the total mean-squared error among all the M -dimension families that have been extracted from a countable Hilbertian basis.

When considering vector-valued random fields, as it will be the case when considering the track-geometry random field, this error can be written as a sum of weighted local errors, where the local errors and the weights are respectively the normalized mean-squared errors and the signal energies associated with each component of $\mathbf{X}$. Therefore, when minimizing the total mean-squared error, we minimize in priority the local errors associated with the components of $\mathbf{X}$ that have the highest signal energies. If the KL projection of the random field $\mathbf{X}$ is then used to propagate variability in mechanical systems, it has therefore to be kept in mind that the particular components of lowest signal energy will not necessary be realistic nor well characterized. If the quantities of interest of the studied system are however very dependent on a precise description of these components, such an optimal KL family may not be relevant and give biased results.

In this prospect, in addition to the classical mean-squared error, two local-global projection errors are introduced in this work:

- ε_{β}^2 corresponds to another weighted sum of local errors, for which weights are *a priori* or *a posteriori* chosen from sensitivity analysis;
- ε_{∞}^2 refers to the maximal value of the local errors associated with each component of random field $\mathbf{X}$.

Indeed, these errors illustrate two classical expectations. On the first hand, error ε_{β}^2 leads to projection families that are particularly adapted to the components of $\mathbf{X}$ of highest chosen weight. If, for a given quantity of interest, the importance of each component of $\mathbf{X}$ can be evaluated from a sensibility analysis, these weights can thus be chosen in order to maximize the relevance of the projection basis to analyze this chosen quantity of interest. On the other hand, if no information is available about the importance of each component of $\mathbf{X}$, making these weights be equal corresponds to the case where no component of $\mathbf{X}$ is favored in the error to be minimized. In such a case, there is however no reason for the minimization of this equally weighted error to lead to a projection family for which each local error would be the same. This thus motivates the introduction of error ε_{∞}^2 , which forces us to search projection families, for which the description precision would be close for each component.

Based on an original scaled expansion of $\mathbf{X}$, the idea of this work is therefore to propose a method to identify the optimal families that respectively minimize errors ε_{β}^2 and ε_{∞}^2 .

In Section 4.2, the scaled expansion is described. In particular, it will be shown how such a formalism allows the identification of the two former optimal basis to be constructed. Section 4.3 illustrates the possibilities of such an expansion on an application based on simulated data.

4.2 Scaled expansion and optimal basis for vector-valued random fields

In this section, the definition of the two local-global errors ε_{β}^2 and ε_{∞}^2 is first presented. The proposed scaled expansion is then introduced for vector-valued random fields. It is finally shown in what extent such a decomposition can lead to the minimization of these two errors.

4.2.1 Local-global errors and optimal basis

Theoretical framework

Adapting the notations of Chapter 2 to the vectorial case, for $Q \geq 1$, let $\mathcal{P}^{(Q)}(\Omega)$ be the space of all the second-order $\mathbb{R}^Q$ -valued random fields, indexed by the compact interval $\Omega = [0, S]$, where it is reminded that $S < +\infty$. Let $\mathbb{H}^{(Q)} = L^2(\Omega, \mathbb{R}^Q)$ be the space of all the square integrable functions on Ω , with values in $\mathbb{R}^Q$, equipped with the inner product $(\cdot, \cdot)$, such that for all $\mathbf{u}$ and $\mathbf{v}$ in $\mathbb{H}^{(Q)}$,

$$(\mathbf{u}, \mathbf{v}) = \int_{\Omega} \mathbf{u}(s)^T \mathbf{v}(s) ds. \quad (4.1)$$

Let $\mathbf{X} = \{(X_1(s), \dots, X_Q(s)), s \in \Omega\}$ be an element of $\mathcal{P}^{(Q)}(\Omega)$. Without loss of generality, it is once again supposed that the mean value of $\mathbf{X}$ is equal to zero:

$$E[\mathbf{X}(s)] = \mathbf{0}, \forall s \in \Omega. \quad (4.2)$$

It is recalled that the signal energy of $\mathbf{X}$, $\|\mathbf{X}\|_{\mathcal{P}^{(Q)}(\Omega)}$, is written:

$$\|\mathbf{X}\|_{\mathcal{P}^{(Q)}(\Omega)} \stackrel{\text{def}}{=} \sqrt{E[(\mathbf{X}, \mathbf{X})]}. \quad (4.3)$$

In the following, $\mathcal{F}^{(M)} = \{\mathbf{f}^m, 1 \leq m \leq M\}$ refers to a set of M deterministic functions that has been extracted from any countable Hilbertian basis of $\mathbb{H}^{(Q)}$. The projection of random field $\mathbf{X}$ on $\mathcal{F}^{(M)}$ is then written $\widehat{\mathbf{X}}^{\mathcal{F}^{(M)}}$. The total normalized mean-squared error associated with

$\mathcal{F}^{(M)}$ is denoted as $\varepsilon^2(\mathcal{F}^{(M)})$ and can thus be written as a sum of weighted **local** normalized mean-squared errors,

$$\varepsilon_q^2(\mathcal{F}^{(M)}) = \frac{\|X_q - \widehat{X}_q^{(M)}\|_{\mathcal{P}(\Omega)}^2}{\|X_q\|_{\mathcal{P}(\Omega)}^2}, \quad 1 \leq q \leq Q, \quad (4.4)$$

associated with each component X_q of random field $\mathbf{X}$:

$$\varepsilon^2(\mathcal{F}^{(M)}) = \sum_{q=1}^Q \left\{ \frac{\|X_q\|_{\mathcal{P}(\Omega)}^2}{\|\mathbf{X}\|_{\mathcal{P}^{(Q)}(\Omega)}^2} \right\} \varepsilon_q^2(\mathcal{F}^{(M)}). \quad (4.5)$$

Optimality of the KL expansion

The matrix-valued covariance function, $[R_{XX}]$, of centered random field $\mathbf{X}$ is introduced as:

$$[R_{XX}(s, s')] = E[\mathbf{X}(s) \otimes \mathbf{X}(s')], \quad \forall (s, s') \in \Omega^2. \quad (4.6)$$

It is assumed that $[R_{XX}]$ is square integrable on $\Omega \times \Omega$, that is to say

$$\|[R_{XX}]\|_{\mathbb{M}}^2 \stackrel{\text{def}}{=} \int_{\Omega} \int_{\Omega} \|[R_{XX}(s, s')]\|_F^2 ds ds' < +\infty, \quad (4.7)$$

with $\|\cdot\|_F$ the Frobenius norm of matrices. It is reminded that the KL basis, $\mathcal{K} = \{\mathbf{k}^m, 1 \leq m\}$, associated with $\mathbf{X}$, can be constructed as a countable Hilbertian basis of $\mathbb{H}^{(Q)}$, which is constituted of the eigenfunctions of covariance matrix-valued function $[R_{XX}]$, such that:

$$\int_{\Omega} [R_{XX}(s, s')] \mathbf{k}^m(s') ds' = \lambda_m \mathbf{k}^m(s), \quad s \in \Omega, \quad 1 \leq m, \quad (4.8)$$

$$(\mathbf{k}^m, \mathbf{k}^j) = \delta_{mj}, \quad \lambda_1 \geq \lambda_2 \geq \dots \rightarrow 0, \quad \sum_{m \geq 1} \lambda_m^2 < +\infty, \quad (4.9)$$

Issues concerning the solving of the integral eigenvalue problem, defined by Eq. (4.8), which is usually called Fredholm problem, can be found in [21, 44, 45]. Due to the orthogonal projection theorem in Hilbert space, for all $M \geq 1$, projection family $\mathcal{K}^{(M)} = \{\mathbf{k}^m, 1 \leq m \leq M\}$ is thus optimal in the sense that, for all family $\mathcal{F}^{(M)}$:

$$\varepsilon^2(\mathcal{K}^{(M)}) \leq \varepsilon^2(\mathcal{F}^{(M)}). \quad (4.10)$$

Let $\widetilde{\mathbf{X}}^{\mathcal{K}^{(M)}}$ be the projection of $\mathbf{X}$ on $\mathcal{K}^{(M)}$. Family $\mathcal{K}^{(M)}$ being orthonormal, it comes:

$$\widetilde{\mathbf{X}}^{(M)} = \sum_{m=1}^M \sqrt{\lambda_m} \mathbf{k}^m \xi_m, \quad (4.11)$$

where $\boldsymbol{\xi} = (\xi_1, \dots, \xi_M)$ is a centered random vector, for which components are uncorrelated and with variance equal to 1. In particular, if $\mathbf{X}$ is a Gaussian random field, the components of $\boldsymbol{\xi}$ are normally distributed and statistically independent.

Local-global errors

From Eq. (4.5), minimizing ε^2 amounts therefore to minimizing in priority the local errors corresponding to the components of $\mathbf{X}$ that have the highest weights $\frac{\|X_q\|_{\mathcal{P}(\Omega)}^2}{\|\mathbf{X}\|_{\mathcal{P}^{(Q)}(\Omega)}^2}$. In other words, for given values of p , q and M , if $\|X_p\|_{\mathcal{P}(\Omega)} \gg \|X_q\|_{\mathcal{P}(\Omega)}$, the minimization of ε^2 can lead to the identification of a M -dimension truncated Karhunen-Loève family associated with $\mathbf{X}$, $\mathcal{K}^{(M)}$, such that $\varepsilon_p^2(\mathcal{K}^{(M)}) \ll \varepsilon_q^2(\mathcal{K}^{(M)})$. Consequently, if X_p and X_q are independent, a two steps approach, based on the definition of two different families (one for X_p and the components of $\mathbf{X}$ that depend on X_p , one for the other components of $\mathbf{X}$ that do not depend on X_p) would be more relevant. On the contrary, if X_p and X_q are indeed dependent, more elements have to be added in $\mathcal{K}^{(M)}$ to make ε_q^2 decrease, or another choice for the error function to be minimized has to be considered.

In this prospect, two local-global projection errors are introduced in this work, ε_β^2 and ε_∞^2 , such that for any β in $]0, +\infty[^Q$:

$$\varepsilon_\beta^2 = \sum_{q=1}^Q \beta_q^2 \varepsilon_q^2, \quad (4.12)$$

$$\varepsilon_\infty^2 = \max_{1 \leq q \leq Q} \{\varepsilon_q^2\}. \quad (4.13)$$

As presented in Section 4.1, if the reduction of the statistical complexity of random field $\mathbf{X}$ is carried out as a first step in a propagation of variability in mechanical systems, minimizing these two errors instead of error ε^2 should allow us to improve the relevance of the projection basis, whether the importance of each component of $\mathbf{X}$ for a given quantity of interest can be evaluated from a sensibility analysis or not.

4.2.2 Scaled expansion

Let $\mathbf{O}$ be an element of $\mathcal{S}^{(Q)}(1) = \{\mathbf{O} \in]0, 1[^Q, \sum_{q=1}^Q O_q^2 = 1\}$. This allows us to define the scaled random field, $\mathbf{Y}(\mathbf{O})$, such that:

$$\mathbf{Y}(\mathbf{O}) = [\text{Diag}(\mathbf{O})] \mathbf{X}, \quad (4.14)$$

$$[\text{Diag}(\mathbf{O})] = \begin{bmatrix} O_1 & 0 & \cdots & 0 \\ 0 & O_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & O_Q \end{bmatrix}. \quad (4.15)$$

The autocorrelation function, $[R_{\mathbf{Y}\mathbf{Y}}(\mathbf{O})]$, of $\mathbf{Y}(\mathbf{O})$ is thus equal to:

$$[R_{\mathbf{Y}\mathbf{Y}}(\mathbf{O})] = [\text{Diag}(\mathbf{O})] [R_{\mathbf{X}\mathbf{X}}] [\text{Diag}(\mathbf{O})]. \quad (4.16)$$

The family $\mathcal{K}^{(M)}(\mathbf{O}) = \{\mathbf{k}^m(\mathbf{O}), 1 \leq m\}$ is thus denoted as the Karhunen-Loève family associated with random field $\mathbf{Y}(\mathbf{O})$, such that:

$$\mathbf{Y}(\mathbf{O}) = \sum_{m=1}^{+\infty} \mathbf{k}^m(\mathbf{O}) \sqrt{\lambda_m(\mathbf{O})} \xi_m(\mathbf{O}), \quad (4.17)$$

$$\lambda_m(\mathbf{O}) = \langle (\mathbf{Y}(\mathbf{O}), \mathbf{k}^m(\mathbf{O})) , (\mathbf{Y}(\mathbf{O}), \mathbf{k}^m(\mathbf{O})) \rangle, \quad \xi_m(\mathbf{O}) = \frac{(\mathbf{Y}(\mathbf{O}), \mathbf{k}^m(\mathbf{O}))}{\sqrt{\lambda_m(\mathbf{O})}}. \quad (4.18)$$

where it is reminded that, by construction, family $\mathcal{K}^{(M)}(\mathbf{O})$ is orthonormal in $\mathbb{H}^{(Q)}$, and projection coefficients $\{\xi_m(\mathbf{O}), m \geq 1\}$ are uncorrelated:

$$(\mathbf{k}^m(\mathbf{O}), \mathbf{k}^j(\mathbf{O})) = E[\xi_m(\mathbf{O})\xi_j(\mathbf{O})] = \delta_{mj}, \quad 1 \leq m, j. \quad (4.19)$$

Since $O_q \neq 0$ for all $1 \leq q \leq Q$, matrix $[\text{Diag}(\mathbf{O})]$ is invertible. Therefore, the projection of random field $\mathbf{X}$ on family $\mathcal{K}^{(M)}(\mathbf{O})$, that is denoted as $\widehat{\mathbf{X}}^{(M)}(\mathbf{O})$, is given by:

$$\widehat{\mathbf{X}}^{(M)}(\mathbf{O}) = \sum_{m=1}^M [\text{Diag}(\mathbf{O})]^{-1} \mathbf{k}^m(\mathbf{O}) \sqrt{\lambda_m(\mathbf{O})} \xi_m(\mathbf{O}), \quad 1 \leq M \quad (4.20)$$

The elements of $\mathcal{K}^{(M)}(\mathbf{O})$ are once again ordered such that the variance of the projection random variables are sorted in a decreasing order:

$$\lambda_1(\mathbf{O}) \geq \lambda_2(\mathbf{O}) \geq \dots \rightarrow 0. \quad (4.21)$$

According to Eqs. (4.4) and (4.5), for all $1 \leq M$, we finally have:

$$\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) = 1 - \frac{O_q^{-2}}{\|\mathbf{X}_q\|_{\mathcal{P}(\Omega)}^2} \sum_{m=1}^M \lambda_m(\mathbf{O}) \int_{\Omega} \{k_q^m(\mathbf{O}, s)\}^2 ds, \quad 1 \leq q \leq Q, \quad (4.22)$$

$$\varepsilon^2(\mathcal{K}^{(M)}(\mathbf{O})) = 1 - \frac{1}{\|\mathbf{X}\|_{\mathcal{P}^{(Q)}(\Omega)}^2} \sum_{q=1}^Q O_q^{-2} \sum_{m=1}^M \lambda_m(\mathbf{O}) \int_{\Omega} \{k_q^m(\mathbf{O}, s)\}^2 ds. \quad (4.23)$$

It can be verified that if $\mathbf{O} = \frac{1}{\sqrt{Q}}(1, \dots, 1)$, the scaled expansion coincides with the classical and direct KL expansion associated with $\mathbf{X}$, defined in Section 4.2.1.

4.2.3 Properties of the scaled expansion

This section aims at emphasizing the main properties of the scaled expansion, on which the minimization of local-global errors ε_{β}^2 and ε_{∞}^2 will be based. First, the continuity of the applications $\mathbf{O} \mapsto \varepsilon_{\beta}^2(\mathcal{K}^{(M)}(\mathbf{O}))$ and $\mathbf{O} \mapsto \varepsilon_{\infty}^2(\mathcal{K}^{(M)}(\mathbf{O}))$ on $S^{(Q)}(1)$ will be shown. Then, the mechanisms induced by the scaled expansion and its optimality are presented.

Lemma 1 *Random field $\mathbf{Y}(\mathbf{O})$ and its realizations are continuous in $\mathbf{O}$ with respect to the L_2 norm on $S^{(Q)}(1)$.*

□ **Proof:** Let $\mathbf{O}$ and $\mathbf{O}^*$ be two elements of $S^{(Q)}(1)$.

1. We have:

$$\begin{aligned} \|\mathbf{Y}(\mathbf{O}) - \mathbf{Y}(\mathbf{O}^*)\|_{\mathcal{P}^{(Q)}(\Omega)}^2 &= \sum_{q=1}^Q (O_q - O_q^*)^2 \|\mathbf{X}_q\|_{\mathcal{P}(\Omega)}^2, \\ &\leq C_Y \|\mathbf{O} - \mathbf{O}^*\|_{\mathbb{R}^Q}^2, \end{aligned} \quad (4.24)$$

where $\|\cdot\|_{\mathbb{R}^Q}$ is the Euclidian norm on $\mathbb{R}^Q$ and $C_Y = \max_{1 \leq q \leq Q} \|X_q\|_{\mathcal{P}(\Omega)}^2$ is a positive constant that is independent of $\mathbf{O}$ and $\mathbf{O}^*$. The application $\mathbf{O} \mapsto \mathbf{Y}(\mathbf{O})$ is therefore continuous on $\mathcal{S}^{(Q)}(1)$ with respect to the L_2 norm.

2. In the same manner, let $\mathbf{X}(\theta)$ be a realization of $\mathbf{X}$, such that, by construction, $\mathbf{Y}(\mathbf{O}, \theta) = [\text{Diag}(\mathbf{O})]\mathbf{X}(\theta)$ and $\mathbf{Y}(\mathbf{O}^*, \theta) = [\text{Diag}(\mathbf{O}^*)]\mathbf{X}(\theta)$ are the corresponding realizations of $\mathbf{Y}(\mathbf{O})$ and $\mathbf{Y}(\mathbf{O}^*)$ respectively. Therefore:

$$\begin{aligned} \|\mathbf{Y}(\mathbf{O}, \theta) - \mathbf{Y}(\mathbf{O}^*, \theta)\|_{L_2}^2 &\stackrel{\text{def}}{=} (\mathbf{Y}(\mathbf{O}, \theta) - \mathbf{Y}(\mathbf{O}^*, \theta), \mathbf{Y}(\mathbf{O}, \theta) - \mathbf{Y}(\mathbf{O}^*, \theta)) \\ &\leq \|\mathbf{O} - \mathbf{O}^*\|_{\mathbb{R}^Q}^2 \left[\max_{1 \leq q \leq Q} \{(X_q(\theta), X_q(\theta))\} \right]. \end{aligned} \quad (4.25)$$

As $\max_{1 \leq q \leq Q} \{(X_q(\theta), X_q(\theta))\}$ is a positive constant that is independent of $\mathbf{O}$ and $\mathbf{O}^*$, the application $\mathbf{O} \mapsto \mathbf{Y}(\mathbf{O}, \theta)$ is continuous on $\mathcal{S}^{(Q)}(1)$ with respect to the norm $\|\cdot\|_{L_2}$.

□

Equation (4.17) and Lemma 1 yield that for any values of the set of random variables $\{\xi_m(\mathbf{O}), 1 \leq m\}$, whose mean values are equal to zero and variances are equal to one, the application $\mathbf{O} \mapsto \sum_{1 \leq m} \sqrt{\lambda_m(\mathbf{O})} \mathbf{k}^m(\mathbf{O}) \xi_m(\mathbf{O})$ is continuous on $\mathcal{S}^{(Q)}(1)$ with respect to the L_2 norm. This motivates the introduction of the following hypothesis, that will be required for the next propositions to be valid.

Hypothesis 1 *For all $1 \leq m$, the applications $\mathbf{O} \mapsto \sqrt{\lambda_m(\mathbf{O})} \mathbf{k}^m(\mathbf{O})$ are supposed to be continuous on $\mathcal{S}^{(Q)}(1)$ with respect to the norm $\|\cdot\|_{L_2}$.*

Proposition 2 *Under Hypothesis 1, the applications $\mathbf{O} \mapsto \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$ are continuous with respect to the Euclidian norm on $\mathcal{S}^{(Q)}(1)$, for all $1 \leq q \leq Q$.*

□ **Proof:** If Hypothesis 1 is verified, due to the continuity properties of the product, of the sum, and of the integral over a closed interval, it can be deduced that for all $1 \leq M$,

$$\mathbf{O} \mapsto \sum_{m=1}^M \lambda_m(\mathbf{O}) \int_{\Omega} \{k_q^m(\mathbf{O}, s)\}^2 ds, \quad 1 \leq q \leq Q, \quad (4.26)$$

are continuous with respect to the Euclidian norm on $\mathcal{S}^{(Q)}(1)$. According to Eq. (4.22), this leads us to the continuity on $\mathcal{S}^{(Q)}(1)$ of the applications $\mathbf{O} \mapsto \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$, for all $1 \leq q \leq Q$.

□

Corrolary 2 *Under Hypothesis 1, the applications $\mathbf{O} \mapsto \varepsilon_{\beta}^2(\mathcal{K}^{(M)}(\mathbf{O}))$ and $\mathbf{O} \mapsto \varepsilon_{\infty}^2(\mathcal{K}^{(M)}(\mathbf{O}))$ are continuous with respect to the Euclidian norm on $\mathcal{S}^{(Q)}(1)$.*

□ **Proof:** By construction of errors ε_{β}^2 and ε_{∞}^2 , defined by Eqs. (4.12) and (4.13), this corrolary is a direct consequence of Proposition 2. □

Proposition 3 *Under Hypothesis 1, for all $1 \leq M$, application $\mathbf{O} \mapsto \varepsilon_{\infty}^2(\mathcal{K}^{(M)}(\mathbf{O}))$ admits a minimal value, $\mathbf{O}_{\infty}^{(M)}$, in $\mathcal{S}^{(Q)}(1)$.*

□ **Proof:** Under Hypothesis 1, Corollary 2 yields that application $\mathbf{O} \mapsto \varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O}))$ is continuous with respect to the Euclidian norm on $\mathcal{S}^{(Q)}(1)$ for all $1 \leq M$. It admits therefore a minimal value in any closed subset $\widehat{\mathcal{S}}(\epsilon) = \{\mathbf{O} \in [\epsilon, 1 - \epsilon]^Q, \sum_{q=1}^Q O_q^2 = 1\}$, for all $0 < \epsilon < 1$.

Then, for $1 \leq q \leq Q$, if O_q tends to zero, it can be noticed that $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$ tends to its maximal value as the weight of X_q in the global minimization is almost zero. This leads us to the fact that it exists $0 < \epsilon^* < 1$ sufficiently small, such that for all $\mathbf{O}$ and $\mathbf{O}^*$ in $\widehat{\mathcal{S}}(\epsilon^*)$ and $\mathcal{S}^{(Q)}(1) \setminus \widehat{\mathcal{S}}(\epsilon^*)$ respectively, $\varepsilon_\infty^2(\mathbf{O}) \leq \varepsilon_\infty^2(\mathbf{O}^*)$. In other words, it exists ϵ^* in $]0, 1[$ and $\mathbf{O}_\infty^{(M)}$ in $\mathcal{S}^{(Q)}(1)$ such that:

$$\mathbf{O}_\infty^{(M)} = \arg \min_{\mathbf{O} \in \widehat{\mathcal{S}}(\epsilon^*)} \{\varepsilon_\infty^2(\mathbf{O})\} = \arg \min_{\mathbf{O} \in \mathcal{S}^{(Q)}(1)} \{\varepsilon_\infty^2(\mathbf{O})\}. \quad (4.27)$$

□

The importance of such a vector $\mathbf{O}_\infty^{(M)}$ for the minimization of error ε_∞^2 will be discussed in Section 4.2.5. Although the perturbation of $[R_{\mathbf{X}\mathbf{X}}]$, defined by Eq. (4.16), is quadratic with respect to vector $\mathbf{O}$, there is no theoretical result in the perturbation theory field that could guarantee the validity of Hypothesis 1 in the general case. From a discrete point of view, applications $\mathbf{O} \mapsto \sqrt{\lambda_m(\mathbf{O})} \mathbf{k}^m(\mathbf{O})$ can however always be considered as continuous, as for any discontinuous application $\mathcal{A}$, it exists a continuous application $\mathcal{A}^*$, such that the projections of $\mathcal{A}$ and $\mathcal{A}^*$ on the same discretized space are the same. Hence, in the following, it is supposed that we are within the framework of Hypothesis 1.

The next Lemma and Proposition aim now at emphasizing how the scaled expansion could be used to favor or put at a disadvantage on purpose the characterization of a particular component of $\mathbf{X}$.

Lemma 2 *For all $\mathbf{O}$ in $\mathcal{S}^{(Q)}(1)$ and for all $\mathcal{F}^{(M)}$ in $(\mathbb{H}^{(Q)})^M$, we have:*

$$\sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \leq \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{F}^{(M)}). \quad (4.28)$$

□ **Proof:** The proof of this lemma is detailed in Appendix A. □

In other words, Lemma 2 underlines that for all $\mathbf{O}$ in $\mathcal{S}^{(Q)}(1)$, family $\mathcal{K}^{(M)}(\mathbf{O})$ is M -optimal for $\mathbf{X}$ regarding error $\sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2$. For $1 \leq p \neq q \leq Q$, imposing $O_p^2 \|X_p\|_{\mathcal{P}(\Omega)}^2 > O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2$ tends therefore to favor the characterization of X_p rather than the one of X_q . This can be seen from the following proposition:

Proposition 4 *For any $\mathbf{O} = (O_1, \dots, O_Q)$ in $\mathcal{S}^{(Q)}(1)$ and for all κ such that $0 < \kappa < \left\{\sum_{q=1}^{Q-1} O_q^2\right\}^{-1/2}$, the vector $\mathbf{O}^* = \left(\kappa O_1, \dots, \kappa O_{Q-1}, \sqrt{1 - \kappa^2 \sum_{q=1}^{Q-1} O_q^2}\right)$ is in $\mathcal{S}^{(Q)}(1)$. For $\kappa = 1$, we have $\mathbf{O} = \mathbf{O}^*$ and κ can be smaller or larger than 1. We then have:*

$$\left\{\varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) - \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}))\right\} \{\kappa^2 - 1\} \geq 0. \quad (4.29)$$

□ **Proof:**

1. If $\mathbf{O} = (O_1, \dots, O_Q)$ is in $\mathcal{S}^{(Q)}(1)$, then $\sum_{q=1}^Q O_q^2 = 1$. Hence, if $0 < \kappa < \left\{\sum_{q=1}^{Q-1} O_q^2\right\}^{-1/2}$, $\sum_{q=1}^Q (O_q^*)^2 = 1$, which shows that $\mathbf{O}^*$ is in $\mathcal{S}^{(Q)}(1)$.

2. Moreover, Lemma 2 yields:

$$\begin{cases} \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \leq \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)), \\ \sum_{q=1}^Q (O_q^*)^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) \leq \sum_{q=1}^Q (O_q^*)^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})), \end{cases} \quad (4.30)$$

which can, for all c_a and c_b in $\mathbb{R}^+$, be written in a more compact form as:

$$\sum_{q=1}^Q \|X_q\|_{\mathcal{P}(\Omega)}^2 \left\{ \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) - \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \right\} \left\{ c_a O_q^2 - c_b (O_q^*)^2 \right\} \geq 0. \quad (4.31)$$

Choosing $c_b = 1$ and $c_a = \kappa^2$ yields:

$$\left\{ \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) - \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O})) \right\} \{ \kappa^2 - 1 \} \geq 0. \quad (4.32)$$

□

Hence, if $\kappa \geq 1$, that is to say if the weights of all components of $\mathbf{X}$, but the one of X_Q , have been increased in the choice of $\mathbf{O}^*$, the projection of X_Q on $\mathcal{K}^{(M)}(\mathbf{O}^*)$ will be less precise than its projection on $\mathcal{K}^{(M)}(\mathbf{O})$ because $\varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) \geq \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}))$. On the contrary, if $\kappa \leq 1$, the weight of X_Q in the scaled expansion defined in Section 4.2.2 is increased by comparison to the other components of $\mathbf{X}$, such that the projection of X_Q on $\mathcal{K}^{(M)}(\mathbf{O}^*)$ will be better than its projection on $\mathcal{K}^{(M)}(\mathbf{O})$ because $\varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}^*)) \leq \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}))$.

By playing on the values of the components of $\mathbf{O}$, the scaled expansion thus appears to be able to favor or put at a disadvantage **on purpose** the characterization of a particular component of $\mathbf{X}$. The goal of the next sections is therefore to define a method to minimize errors ε_{β}^2 and ε_{∞}^2 , based on this scaled expansion.

4.2.4 Minimization of a weighted sum of local errors

The minimization of error ε_{β}^2 , defined by Eq. (4.12), is a direct consequence of Lemma 2. Indeed, for all β in $\mathcal{S}^{(Q)}(1)$, it can directly be seen that the choice

$$O_q^{\beta} = \frac{\beta_q}{\|X_q\|_{\mathcal{P}(\Omega)}}, \quad 1 \leq q \leq Q, \quad (4.33)$$

leads us to the minimization of error ε_{β}^2 , such that:

$$\mathcal{K}^{(M)}(\mathbf{O}^{\beta}) = \arg \min_{\mathcal{F}^{(M)} \in (\mathbb{H}^{(Q)})^M} \left\{ \varepsilon_{\beta}^2(\mathcal{F}^{(M)}) \right\}, \quad 1 \leq M. \quad (4.34)$$

Hence, just by considering the KL expansion of $\mathbf{Y}(\mathbf{O}) = [\text{Diag}(\mathbf{O})] \mathbf{X}$ rather than $\mathbf{X}$, it is possible to construct projection families that could favor particular components of $\mathbf{X}$, from *a priori* or *a posteriori* choices for β .

4.2.5 Minimization of the maximal value of the local errors

The optimal families $\mathcal{F}_\infty^{(M)}$, which minimize error ε_∞^2 , that is to say such that:

$$\mathcal{F}_\infty^{(M)} = \arg \min_{\mathcal{F}^{(M)} \in (\mathbb{H}^{(Q)})^M} \left\{ \varepsilon_\infty^2(\mathcal{F}^{(M)}) \right\} = \arg \min_{\mathcal{F}^{(M)} \in (\mathbb{H}^{(Q)})^M} \left\{ \max_{1 \leq q \leq Q} \varepsilon_q^2(\mathcal{F}^{(M)}) \right\}, \quad (4.35)$$

have been introduced to minimize the local errors associated with each component of $\mathbf{X}$. These projection families stem however from a Min-Max optimization on the very large space $(\mathbb{H}^{(Q)})^M$, such that their direct numerical identification can be very difficult. As the dimension of $\mathcal{S}^{(Q)}(1)$ is comparatively very small, the idea presented in this section is thus to use the former scaled expansion, defined in Section 4.2.2, to approximate $\mathcal{F}_\infty^{(M)}$ as the solution of an optimization problem with respect to $\mathbf{O}$ in $\mathcal{S}^{(Q)}(1)$, rather than an optimization problem with respect to $\mathcal{F}^{(M)}$ in $(\mathbb{H}^{(Q)})^M$. For all $1 \leq M$, we thus define $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})$ as the scaled basis associated with the vector $\mathbf{O}_\infty^{(M)}$, such that:

$$\mathbf{O}_\infty^{(M)} = \arg \min_{\mathbf{O} \in \mathcal{S}^{(Q)}(1)} \left\{ \varepsilon_\infty^2(\mathbf{O}) \right\}, \quad (4.36)$$

for which existence stems from Proposition 3.

Whereas vector $\mathbf{O}^\beta$, defined by Eq. (4.33), is independent of M , it has to be reminded that vector $\mathbf{O}_\infty^{(M)}$ depends on M in the general case.

This section aims first at quantifying the distance between $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})$ and $\mathcal{F}_\infty^{(M)}$. In the two dimensional case ($Q = 2$), it will be shown in particular that $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)}) = \mathcal{F}_\infty^{(M)}$. At last, based on Proposition 4, an algorithm to numerically solve Eq. (4.36) is presented.

Quantification of the error introduced by the approximated identification problem

Lemma 3 *For all $M \geq 1$ and for all $\mathbf{O}$ in $\mathcal{S}^{(Q)}(1)$, the relevance of $\mathcal{K}^{(M)}(\mathbf{O})$ to minimize error ε_∞^2 can be assessed as:*

$$0 \leq \varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O})) - \varepsilon_\infty^2(\mathcal{F}_\infty^{(M)}) \leq \mathcal{UB}(\mathbf{O}), \quad (4.37)$$

where:

$$\mathcal{UB}(\mathbf{O}) \stackrel{\text{def}}{=} \frac{\sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \delta_q^2(\mathcal{K}^{(M)}(\mathbf{O}))}{\sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2}, \quad (4.38)$$

$$0 \leq \delta_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \stackrel{\text{def}}{=} \varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O})) - \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})), \quad 1 \leq q \leq Q. \quad (4.39)$$

□ **Proof:** The first inequality $\varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O})) \geq \varepsilon_\infty^2(\mathcal{F}_\infty^{(M)})$ is a direct consequence of the optimality of $\mathcal{F}_\infty^{(M)}$. Let $\mathbf{O}$ be an element in $\mathcal{S}^{(Q)}(1)$. From Lemma 2, it can therefore be deduced that:

$$\begin{aligned} \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) &\leq \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2(\mathcal{F}_\infty^{(M)}) \\ &\leq \varepsilon_\infty^2(\mathcal{F}_\infty^{(M)}) \left\{ \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \right\}, \end{aligned} \quad (4.40)$$

such that, by definition of $\{\delta_1^2(\mathcal{K}^{(M)}(\mathbf{O})), \dots, \delta_Q^2(\mathcal{K}^{(M)}(\mathbf{O}))\}$:

$$\left\{ \varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O})) - \varepsilon_\infty^2(\mathcal{F}_\infty^{(M)}) \right\} \left\{ \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \right\} \leq \sum_{q=1}^Q O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \delta_q^2(\mathcal{K}^{(M)}(\mathbf{O})), \quad (4.41)$$

which proves the second part of the inequality. $\square$

This Lemma emphasizes that the closer the local errors are, the more relevant projection family $\mathcal{K}^{(M)}(\mathbf{O})$ is. In particular, quantity $\mathcal{UB}(\mathbf{O}_\infty^{(M)})$ defines an upper bound for the error introduced by the consideration of the approximated problem defined by Eq. (4.36). Lemma 3 leads us moreover to the following proposition:

Proposition 5 *If the following equalities are verified:*

$$\varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = \dots = \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})), \quad (4.42)$$

then the family $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})$ minimizes ε_∞^2 .

$\square$ **Proof:** By construction, if $\varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = \dots = \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)}))$, then $\delta_1^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = \dots = \delta_Q^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = 0$. Hence, from Lemma 3, we get $\varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = \varepsilon_\infty^2(\mathcal{F}_\infty^{(M)})$, such that $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)}) = \mathcal{F}_\infty^{(M)}$. $\square$

Identification of the optimal scaling vector

By construction, it can directly be seen that, for all $\alpha > 0$, $\mathcal{K}^{(M)}(\mathbf{O}) = \mathcal{K}^{(M)}(\alpha \mathbf{O})$. Hence, if the conditions of Proposition 5 are fulfilled, that is to say if $\varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) = \dots = \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)}))$, the scaling vector $\mathbf{O}_\infty^{(M)}$ is solution of the following problem:

$$\mathcal{K}^{(M)}(\mathbf{O}) = \mathcal{K}^{(M)} \left([\text{Diag}(\mathbf{O})] \boldsymbol{\epsilon}^2(\mathcal{K}^{(M)}(\mathbf{O})) \right), \quad (4.43)$$

where the matrix $[\text{Diag}(\mathbf{O})]$ is defined by Eq. (4.15), and where:

$$\boldsymbol{\epsilon}^2(\mathcal{K}^{(M)}(\mathbf{O})) = \left(\varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O})), \dots, \varepsilon_Q^2(\mathcal{K}^{(M)}(\mathbf{O})) \right). \quad (4.44)$$

This motivates the following iterative algorithm for the identification of scaling vector $\mathbf{O}_\infty^{(M)}$. For given parameters τ and γ :

$$\left[\begin{array}{l} \text{Initialize } \mathbf{O} = \left(\frac{1}{\|X_1\|_{\mathcal{P}(\Omega)}}, \dots, \frac{1}{\|X_Q\|_{\mathcal{P}(\Omega)}} \right) \\ \text{Normalize } \mathbf{O} \\ \text{for } i = 1 : N_{\max} \\ \quad \text{Compute } \mathcal{K}^{(M)}(\mathbf{O}) \\ \quad \text{if } \mathcal{UB}(\mathbf{O}) > \tau : \\ \quad \quad O_q = O_q \cdot (\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})))^\gamma, \quad 1 \leq q \leq Q \\ \quad \quad \text{Normalize } \mathbf{O} \\ \quad \text{else} \\ \quad \quad \text{Break loop for} \\ \quad \text{end if} \\ \text{end for} \\ \mathbf{O}_\infty^{(M)} = \mathbf{O}. \end{array} \right. \quad (4.45)$$

Parameter τ corresponds to the chosen precision of the numerical convergence, whereas γ controls the speed of the convergence and has to be adapted to avoid numerical instabilities. For our applications, γ will be chosen equal to $1/2$. In such an algorithm, at each iteration $(n+1)$, the weight of X_q in the KL expansion, $\left(O_q^{(n+1)}\right)^2 \|X_q\|_{\mathcal{P}(\Omega)}^2$, is updated with respect to the local error $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}^{(n)}))$ of the former step. Hence, the weights of the less well characterized components of $\mathbf{X}$, for which local errors $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}^{(n)}))$ are the highest at iteration n , will be increased the most at the new iteration $(n+1)$. In the general case, no convergence property for this algorithm has been proved yet, but under the following conditions:

$$\lim_{O_q^2 \rightarrow 1} \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \leq \min_{1 \leq p \neq q \leq Q} \left\{ \lim_{O_q^2 \rightarrow 1} \varepsilon_p^2(\mathcal{K}^{(M)}(\mathbf{O})) \right\}, \quad 1 \leq q \leq Q, \quad (4.46)$$

it is assumed that the algorithm defined by Eq. (4.45) gives very promising results for the minimization of function $\mathbf{O} \mapsto \varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O}))$ in a very few number of iterations. In other words, in cases where the weight of X_q in the scaled expansion is much higher than the weights of the other components $\{X_p, 1 \leq p \neq q \leq Q\}$, if X_q still remains badly characterized, then there is no reason for such an algorithm to converge to a satisfying result. In practice, these conditions are not very restrictive, and are most of the time verified for correlated vector-valued random fields.

In particular, under Hypothesis 1, when dealing with a two dimensional case ($Q = 2$, $\mathbf{O} = (O_1, \sqrt{1 - O_1^2})$), Propositions 2 and 4 yield that errors functions $O_1 \mapsto \varepsilon_1^2(O_1)$ and $O_1 \mapsto \varepsilon_2^2(O_1)$ are continuous and respectively decreases and increases with respect to O_1 in $]0, 1[$. Therefore, if the conditions defined by Eq. (4.46) are fulfilled, it exists $\mathbf{O}_\infty^{(M)}$ in $\mathcal{S}^{(Q)}(1)$ such that $\varepsilon_1^2(\mathbf{O}_\infty^{(M)}) = \varepsilon_2^2(\mathbf{O}_\infty^{(M)})$. Therefore, according to Proposition 5, optimal basis $\mathcal{F}_\infty^{(M)}$ could be in these cases exactly identified from the solving of the optimization problem that is defined by Eq. (4.36).

4.3 Application

Most of the results emphasized in Section 4.2 are illustrated in this section on a practical example. This section is divided in three parts: first, a particular $\mathbb{R}^4$ -valued random field is generated from its Karhunen-Loève expansion; then the influence of scaling vector $\mathbf{O}$ on the local errors is emphasized; at last, it is shown in what extent the scaled expansion allows us to identify optimal families $\mathcal{F}_\infty^{(M)}$ and $\mathcal{F}_\beta^{(M)}$ for several values of β in $\mathcal{S}^{(Q)}(1)$ and any values of $M \geq 1$.

4.3.1 Generation of a vector-valued random field

In this application, the dimension of random field $\mathbf{X}$, Q , is chosen equal to 4, and $\Omega = [0, 1]$. A particular matrix-valued covariance function, $[R_{\mathbf{X}\mathbf{X}}]$, is then postulated, for which some projections are represented in Figures 4.1 and 4.2. Random field $\mathbf{X}$, which is still supposed to be centered, can thus be written as:

$$\mathbf{X} = \sum_{m=1}^{+\infty} \sqrt{\lambda_m} \mathbf{k}^m \xi_m, \quad (4.47)$$

where, for all $m \geq 1$, couples $(\lambda_m, \mathbf{k}^m)$ are solution of the Fredholm problem associated with $[R_{\mathbf{X}\mathbf{X}}]$:

$$\int_{\Omega} [R_{\mathbf{X}\mathbf{X}}(s, s')] \mathbf{k}^m(s') ds' = \lambda_m \mathbf{k}^m(s), \quad \forall s \in \Omega, \quad (4.48)$$

and coefficients $\{\xi_m, m \geq 1\}$ are uncorrelated random variables. For the sake of simplicity, these coefficients are moreover considered independent and normally distributed, which amounts to supposing that $\mathbf{X}$ is Gaussian. In particular, $[R_{\mathbf{X}\mathbf{X}}]$ has been chosen such that $\|X_1\|_{\mathcal{P}(\Omega)} > \|X_2\|_{\mathcal{P}(\Omega)} > \|X_3\|_{\mathcal{P}(\Omega)} > \|X_4\|_{\mathcal{P}(\Omega)}$. Further details about the generation of $[R_{\mathbf{X}\mathbf{X}}]$ can be seen in Appendix B. As an illustration, a particular realization, $\mathbf{X}(\theta)$, of $\mathbf{X}$ is represented in Figure 4.3. From Eq. (4.22), it is reminded that for any value of $\mathbf{O}$ in $\mathcal{S}^{(Q)}(1)$, for all $1 \leq q \leq Q$, and for all $M \geq 1$, errors $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$ can directly be computed by the scaled expansion.

4.3.2 Influence of the scaling vector on the local errors

According to Section 4.2, by introducing vector $\mathbf{O} = (O_1, O_2, O_3, O_4)$, we should be able to balance the values of local errors ε_q^2 , for $1 \leq q \leq 4$. In particular, it has been shown in Section 4.2.3 that for $\mathbf{O} = \frac{1}{\sqrt{3+\kappa^2}}(1, 1, 1, \kappa)$ and for all $1 \leq M$, $\varepsilon_4^2(\mathcal{K}^{(M)}(\mathbf{O}))$ decreases with respect to κ on $]0, +\infty[$. Hence, if κ tends to zero, $\varepsilon_4^2(\mathcal{K}^{(M)}(\mathbf{O}))$ is bound to converge to its maximal value, as the weight of X_4 in the minimization of $\sum_{q=1}^4 O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2$ becomes negligible. On the contrary, if κ tends to infinity, $\varepsilon_4^2(\mathcal{K}^{(M)}(\mathbf{O}))$ will tend to its minimal value, as the minimization of $\sum_{q=1}^4 O_q^2 \|X_q\|_{\mathcal{P}(\Omega)}^2 \varepsilon_q^2$ will completely be driven by ε_4^2 . This phenomenon can be seen in Figure 4.4, where the evolution of local errors $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$ with respect to κ is represented.

In the same manner, the results concerning the two dimensions case can be illustrated from this four dimensions case, by imposing:

$$\mathbf{O} = \frac{1}{\sqrt{O_1^2 + 2 \times 10^{-10} + O_4^2}} (O_1, 10^{-5}, 10^{-5}, O_4). \quad (4.49)$$

Indeed, in such a case, the weights of X_2 and X_3 will always be negligible. In Figure 4.5, it can therefore be seen that when ratio O_4/O_1 increases, $\varepsilon_4^2(\mathcal{K}^{(M)}(\mathbf{O}))$ decreases from its maximal value to its minimal value, whereas $\varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O}))$ increases from its minimal value to its maximal value. As

$$\left\{ \begin{array}{l} \lim_{O_4^2/O_1^2 \rightarrow 0} \varepsilon_1^2(\mathcal{K}^{(M)}(\mathbf{O})) < \min_{2 \leq q \leq 4} \left\{ \lim_{O_4^2/O_1^2 \rightarrow 0} \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \right\}, \\ \lim_{O_4^2/O_1^2 \rightarrow +\infty} \varepsilon_4^2(\mathcal{K}^{(M)}(\mathbf{O})) < \min_{1 \leq q \leq 3} \left\{ \lim_{O_4^2/O_1^2 \rightarrow 0} \varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O})) \right\}, \end{array} \right. \quad (4.50)$$

it exists a value for O_4/O_1 in $]0, +\infty[$ such that ε_1^2 and ε_4^2 are equal. This value allows us therefore to identify a projection family which is M -optimal for $\mathbf{X}$ with respect to the error $\max_{p \in \{1,4\}} \{\varepsilon_p^2\}$.

4.3.3 Identification of the optimal basis

In Section 4.2, for all β in $\mathcal{S}^{(Q)}(1)$, optimal projection families $\mathcal{F}_\beta^{(M)}$ and $\mathcal{F}_\infty^{(M)}$ have been introduced as the solutions of the two following optimization problems:

$$\mathcal{F}_\beta^{(M)} = \arg \min_{\mathcal{F}^{(M)} \in (\mathbb{H}^{(Q)})^M} \left\{ \varepsilon_\beta^2(\mathcal{F}^{(M)}) \right\}, \quad (4.51)$$

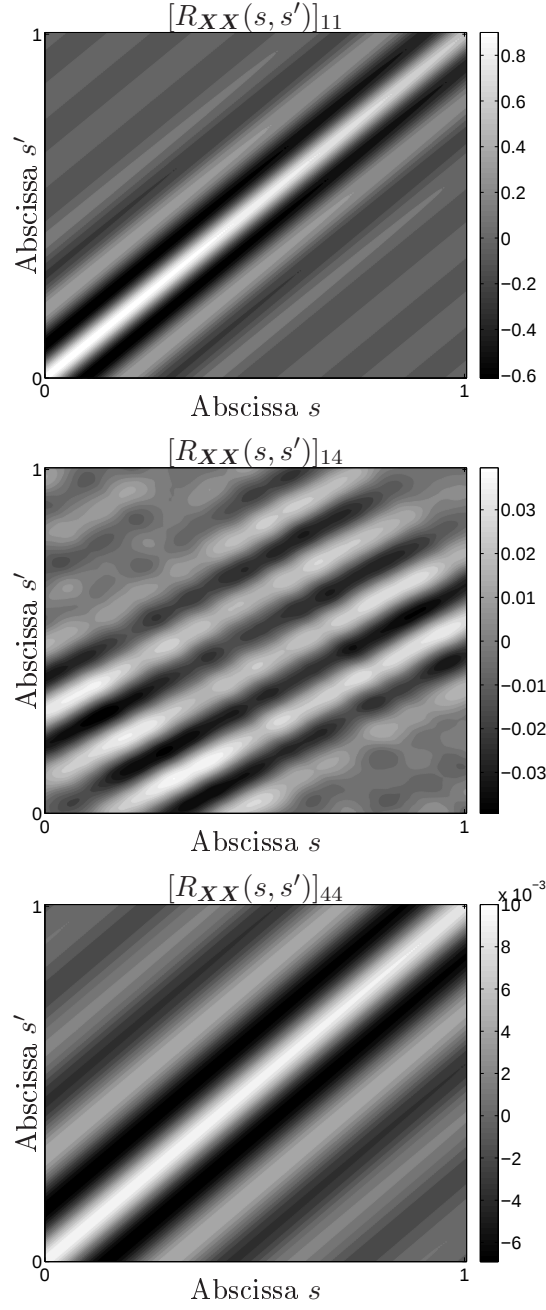


Figure 4.1: Representations of $[R_{XX}]_{11}$, $[R_{XX}]_{14}$ and $[R_{XX}]_{44}$ on $[0, 1] \times [0, 1]$.

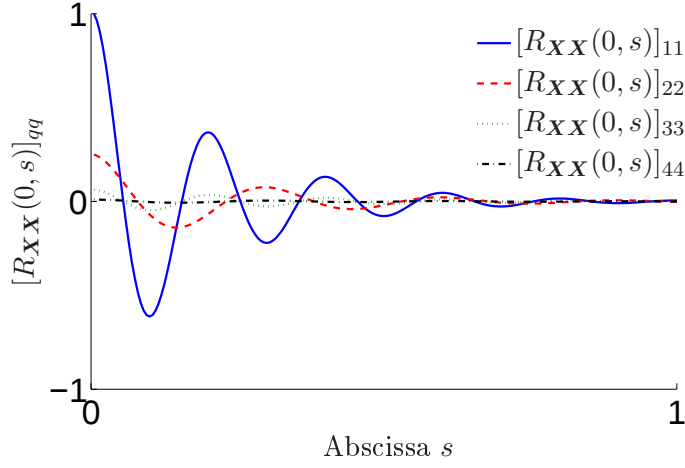


Figure 4.2: Representations of $s \mapsto [R_{\mathbf{X}\mathbf{X}}(0, s)]_{qq}$, for $1 \leq q \leq 4$.

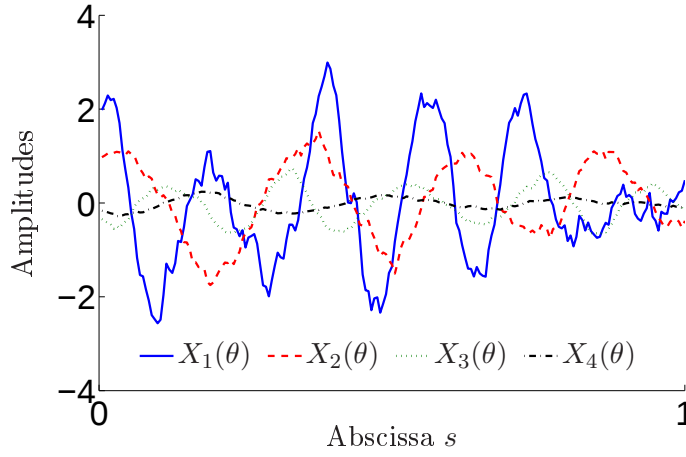


Figure 4.3: Representation of $\mathbf{X}(\theta)$.

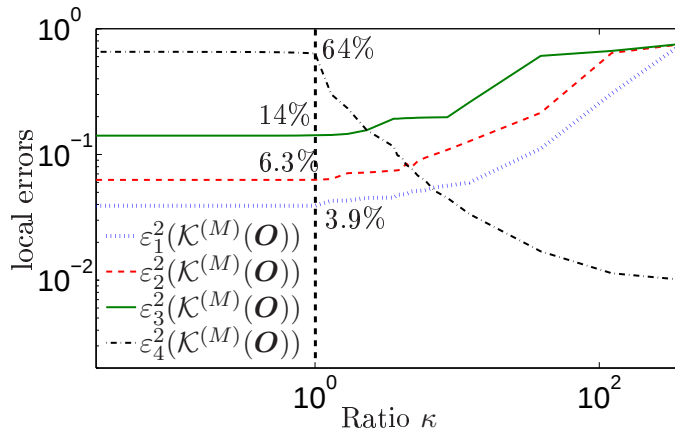


Figure 4.4: Comparison of local errors $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$, when $\mathbf{O} = \frac{1}{\sqrt{3+\kappa^2}}(1, 1, 1, \kappa)$, with respect to κ , for $M = 50$.

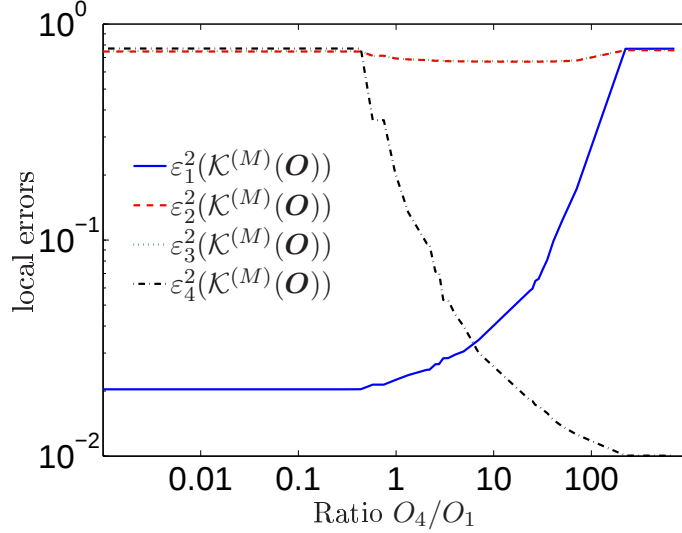


Figure 4.5: Comparison of local errors $\varepsilon_q^2(\mathcal{K}^{(M)}(\mathbf{O}))$, when $\mathbf{O} = \frac{1}{\sqrt{O_1^2 + 2 \cdot 10^{-10} + O_4^2}} (O_1, 10^{-5}, 10^{-5}, O_4)$ with respect to ratio O_4/O_1 , for $M = 50$.

$$\mathcal{F}_\infty^{(M)} = \arg \min_{\mathcal{F}^{(M)} \in (\mathbb{H}^{(Q)})^M} \left\{ \varepsilon_\infty^2(\mathcal{F}^{(M)}) \right\}. \quad (4.52)$$

In particular, for all $M \geq 1$, the choice

$$\beta = \left(\|X_1\|_{\mathcal{P}(\Omega)}, \|X_2\|_{\mathcal{P}(\Omega)}, \|X_3\|_{\mathcal{P}(\Omega)}, \|X_4\|_{\mathcal{P}(\Omega)} \right) \quad (4.53)$$

leads to the identification of the classical Karhunen-Loève family, which is called $\mathcal{F}_{L_2}^{(M)}$, for $\mathbf{X}$. The corresponding local errors, $\varepsilon_q^2(\mathcal{F}_{L_2}^{(M)})$ can then be compared. In Figure 4.4, for $M = 50$, it can be seen that $\varepsilon_1^2(\mathcal{F}_{L_2}^{(50)}) = 3.9\%$, $\varepsilon_2^2(\mathcal{F}_{L_2}^{(50)}) = 6.3\%$, $\varepsilon_3^2(\mathcal{F}_{L_2}^{(50)}) = 14\%$ and $\varepsilon_4^2(\mathcal{F}_{L_2}^{(50)}) = 64\%$. Due to the fact that $\|X_1\|_{\mathcal{P}(\Omega)} > \|X_2\|_{\mathcal{P}(\Omega)} > \|X_3\|_{\mathcal{P}(\Omega)} > \|X_4\|_{\mathcal{P}(\Omega)}$, it can thus be verified that the direct Karhunen-Loève expansion favors the description of component X_1 , whereas component X_4 is not precisely characterized.

As explained in Section 4.2, other values for β have to be considered in order to improve the characterization of X_4 . For instance, the choice $\beta = (0.5, 0.5, 0.5, 0.5)$ corresponds to the minimization of the mean value of the local errors, $\varepsilon_\mu^2 = \frac{1}{4} \sum_{q=1}^4 \varepsilon_q^2$. Let $\mathcal{F}_\mu^{(M)}$ be the corresponding optimal family. Any other value for β can nevertheless be chosen. For instance, let $\mathcal{F}_\beta^{(M)}$ be the M -dimension optimal family corresponding to the case $\beta = (0.1, 2, 1, 0.5) / 2.2935$. At last, family $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})$ is introduced as the numerical solution of the algorithm defined by Eq. (4.45), with $\tau = 10^{-3}$ and $\gamma = 1/2$.

In this prospect, Figures 4.6 and 4.7 allow us to numerically illustrate that projection families $\mathcal{F}_\beta^{(M)}$, $\mathcal{F}_\mu^{(M)}$, $\mathcal{F}_{L_2}^{(M)}$ and $\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})$ can be identified from the scaled expansion, such that for any $M \geq 1$:

- $\varepsilon_\beta^2(\mathcal{F}_\beta^{(M)}) \leq \min \left\{ \varepsilon_\beta^2(\mathcal{F}_{L_2}^{(M)}), \varepsilon_\beta^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})), \varepsilon_\beta^2(\mathcal{F}_\mu^{(M)}) \right\},$
- $\varepsilon_\mu^2(\mathcal{F}_\mu^{(M)}) \leq \min \left\{ \varepsilon_\mu^2(\mathcal{F}_{L_2}^{(M)}), \varepsilon_\mu^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})), \varepsilon_\mu^2(\mathcal{F}_\beta^{(M)}) \right\},$

- $\varepsilon_\infty^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})) \leq \min \left\{ \varepsilon_\infty^2(\mathcal{F}_{L_2}^{(M)}), \varepsilon_\infty^2(\mathcal{F}_\beta^{(M)}), \varepsilon_\infty^2(\mathcal{F}_\mu^{(M)}) \right\},$
- $\varepsilon^2(\mathcal{F}_{L_2}^{(M)}) \leq \min \left\{ \varepsilon^2(\mathcal{F}_\beta^{(M)}), \varepsilon^2(\mathcal{K}^{(M)}(\mathbf{O}_\infty^{(M)})), \varepsilon^2(\mathcal{F}_\mu^{(M)}) \right\}.$

In particular, for $M = 100$:

$$\left\{ \begin{array}{l} \varepsilon_1^2(\mathcal{F}_{L_2}^{(100)}) = 1.7\% \\ \varepsilon_2^2(\mathcal{F}_{L_2}^{(100)}) = 3.0\% \\ \varepsilon_3^2(\mathcal{F}_{L_2}^{(100)}) = 5.8\% \\ \varepsilon_4^2(\mathcal{F}_{L_2}^{(100)}) = 17\% \end{array} \right\}, \quad \left\{ \begin{array}{l} \varepsilon_1^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \\ \varepsilon_2^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \\ \varepsilon_3^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \\ \varepsilon_4^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \end{array} \right\}, \quad (4.54)$$

$$\left\{ \begin{array}{l} \varepsilon^2(\mathcal{F}_{L_2}^{(100)}) = 2.3\% \\ \varepsilon^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \\ \varepsilon_\infty^2(\mathcal{F}_{L_2}^{(100)}) = 17\% \\ \varepsilon_\infty^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) = 3.0\% \end{array} \right\}. \quad (4.55)$$

Whereas family $\mathcal{F}_{L_2}^{(100)}$ can put at a disadvantage the description of a particular component of $\mathbf{X}$ to minimize ε^2 , family $\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})$ tries to equilibrate the precision of the description of each component. To do so, the local error of some components can increase to make the other decrease. Indeed, in this example, $\varepsilon_1^2(\mathcal{F}_{L_2}^{(100)}) < \varepsilon_1^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)}))$ whereas $\varepsilon_4^2(\mathcal{F}_{L_2}^{(100)}) > \varepsilon_4^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)}))$. From Eq. (4.37), it can moreover be seen that in this case:

$$\left| \varepsilon_\infty^2(\mathcal{F}_\infty^{(100)}) - \varepsilon_\infty^2(\mathcal{K}^{(100)}(\mathbf{O}_\infty^{(M)})) \right| \leq \tau = 0.1\%. \quad (4.56)$$

4.4 Conclusions

In spite of the increasing computational power that has encouraged the development of computational models with always more degrees of freedom, statistical reduction methods, such as the Karhunen-Loève expansion, still have a big role to play to make the solving of these problems faster and more robust. When dealing with $\mathbb{R}^Q$ -valued random fields $\mathbf{X} = (X_1, \dots, X_Q)$, it has however been shown in this chapter that the direct truncated KL expansion, which minimizes the total mean-squared error, tends to better characterize the components of $\mathbf{X}$ that have the highest signal energy. In this context, a particular adaptation of the KL expansion has been proposed. Based on a scaling transformation of $\mathbf{X}$, this original decomposition allows defining projection basis that can favor or put at a disadvantage **on purpose** the characterization of a particular component of $\mathbf{X}$. This expansion appears to be also very relevant to identify projection basis that minimize the maximal value of the local errors of $\mathbf{X}$. Finally, when interested in studying complex systems that are excited by vector-valued random fields (one can think about the interactions between trains and track irregularities, buildings and earthquakes, harbors and swell, etc.), the method proposed opens new opportunities to adapt the projection basis with respect to the quantities of interest of the systems.

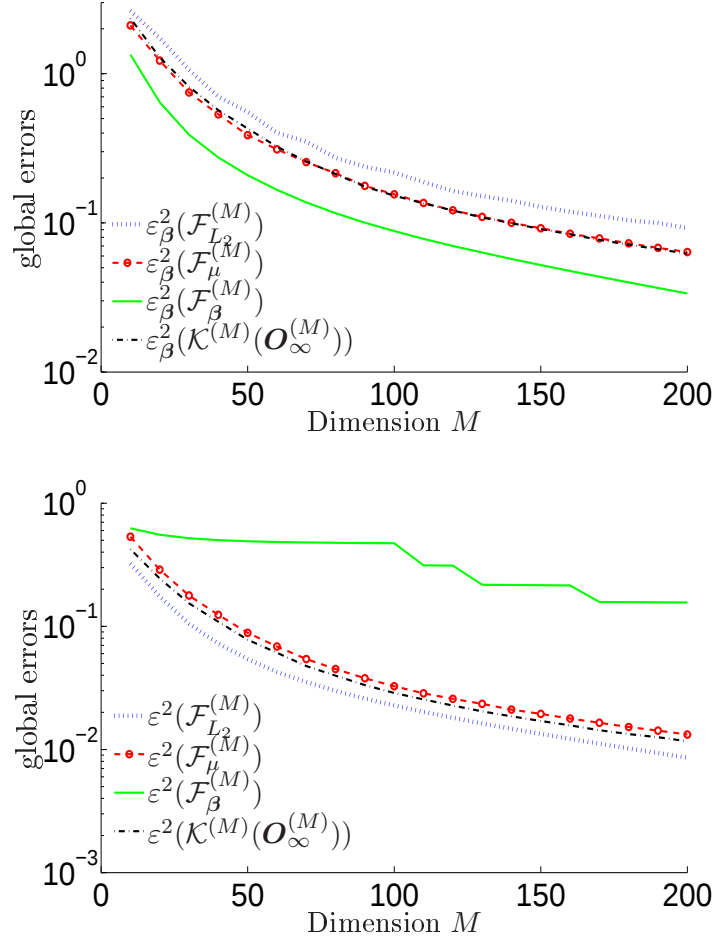


Figure 4.6: Evolution of errors ε_{β}^2 and ε^2 with respect to the dimension of the projection family, M .

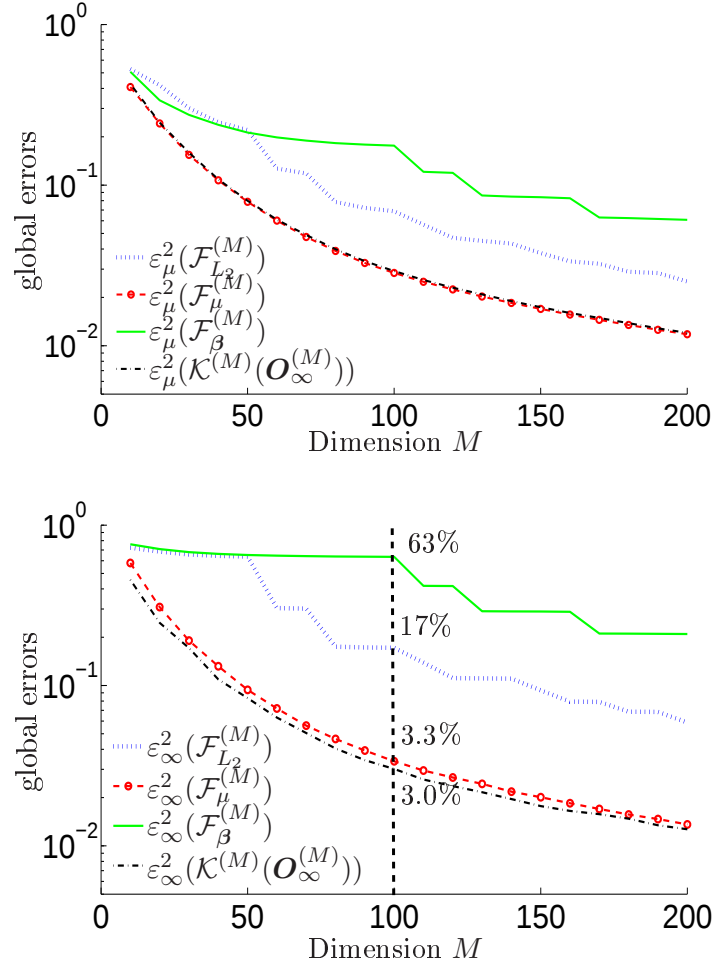


Figure 4.7: Evolution of errors ε_μ^2 and ε_∞^2 with respect to the dimension of the projection family, M .

Chapter 5

Experimental identification of the railway track stochastic modeling

5.1 Introduction

The expected benefits of simulation in the railway field are multiple: robust and optimized conception, shorter and cheaper certification procedure, better knowledge of the critical situations of the track-vehicle system, optimization of the maintenance. However, if simulation is introduced in certification and conception processes, it has to be very representative of the physical behavior of the system. The model has thus to be fully validated and the simulations have to be raised on a realistic and representative set of excitations.

A particular attention has therefore to be paid to the track geometry, which is the main source of excitation of the train. Two description scales can then be considered for this geometry. On the first hand, the track design, which corresponds to the mean line position of a perfect track is decided once for all at the building of a new track. This description is characterized by three curvilinear quantities: the vertical curvature c_V , the horizontal curvature c_H , and the track superelevation c_L . On the other hand, for a fixed track design, the actual positions of the rails are in constant evolution, which is mostly due to the interactions between the train, the track and the substructure. The irregularities appearing during the track lifecycle are of four types (see Figure 5.1): lateral and vertical alignment irregularities x_1 and x_2 on the one hand, cant deficiencies x_3 and gauge irregularities x_4 on the other hand. Therefore, each rail position $\mathbf{R}_{\ell/r}$ (ℓ refers to the left rail whereas r refers to the right rail) can be written as the sum of a mean position $\mathbf{M}_{\ell/r}$, which only depends on the curvilinear abscissa of the track, s , the track gauge E , and the three parameters of the track design, c_H , c_V and c_L , and a deviation toward this mean position $\mathbf{I}_{\ell/r}$, which only depends on the track irregularities:

$$\mathbf{R}_{\ell/r}(s) = \mathbf{M}_{\ell/r}(s) + \mathbf{I}_{\ell/r}(s), \quad (5.1)$$

$$\mathbf{M}_{\ell/r}(s) = \mathbf{O}_{\text{NT}}(s) \pm \frac{E}{2} \mathbf{N}(s), \quad (5.2)$$

$$\mathbf{I}_{\ell/r}(s) = \{x_2(s) \pm x_3(s)\} \mathbf{B}(s) + \{x_1(s) \pm x_4(s)\} \mathbf{N}(s), \quad (5.3)$$

where $-$ goes with the subscript ℓ and $+$ goes with r in the symbol $\pm$, $\mathbf{O}_{\text{NT}}(s) = (\mathbf{M}_{\ell}(s) + \mathbf{M}_{/r}(s))/2$ is the mean position of the two rails, and $(\mathbf{O}_{\text{NT}}(s), \mathbf{T}(s), \mathbf{N}(s), \mathbf{B}(s))$ is the Frenet frame. Hence, a track geometry $\mathcal{T}$ of total length S^{tot} is completely characterized by the knowledge of seven curvilinear functions:

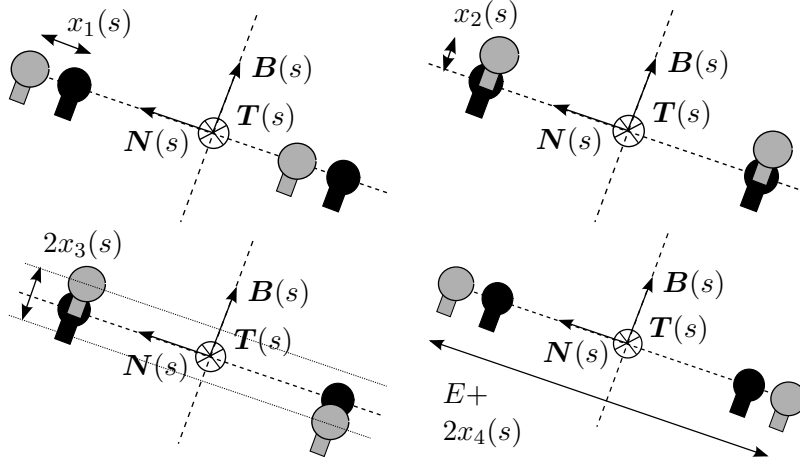


Figure 5.1: Parametrization of the track irregularities (for each rail, the mean position is represented in black, whereas the real position is in grey).

$$\mathcal{T} = \{(x_1(s), x_2(s), x_3(s), x_4(s), c_L(s), c_H(s), c_V(s)), s \in [0, S^{\text{tot}}]\}. \quad (5.4)$$

However, as the mean line of the track geometry is chosen at the building of a new line, this work is only devoted to the modeling of the track irregularity vector

$$\mathbf{x} = (x_1, x_2, x_3, x_4), \quad (5.5)$$

where x_1 , x_2 , x_3 and x_4 are the four types of track irregularities previously introduced.

Made up of straight lines and curves at its construction, the new track evolves gradually due to the train dynamics and is regularly subjected to maintenance operations. During their lifecycles, trains are therefore confronted with very different running conditions. The track-vehicle system being strongly nonlinear, the dynamic behavior of trains has thus to be analyzed not only for a few track portions but for these whole realms of possibility.

In reply to these expectations, the measurements of the train IRIS 320 are of great interest. Indeed, this one has been running continuously since 2007 over the French railway network, measuring and recording the track geometry of the main national lines. Based on these experimental measurements, a complete parametrization of the track geometry and of its variability would be of great concern in specification, security and certification prospects, to be able to generate track geometries that are realistic and representative of a whole railway network.

In this context, this chapter develops a stochastic modeling of the track geometry, which is based on an inverse identification of the statistical properties of a vector-valued random field from measured data. These data being complete, this modeling allows us to generate numerically track geometries that are physically realistic and statistically representative of the set of available track measurements. Moreover, these tracks can be used in any deterministic railway dynamic software to characterize the dynamic behavior of the train.

Hence, this modeling could bring an innovative technical answer to introduce numerical methods and treatments in the maintenance and certification processes.

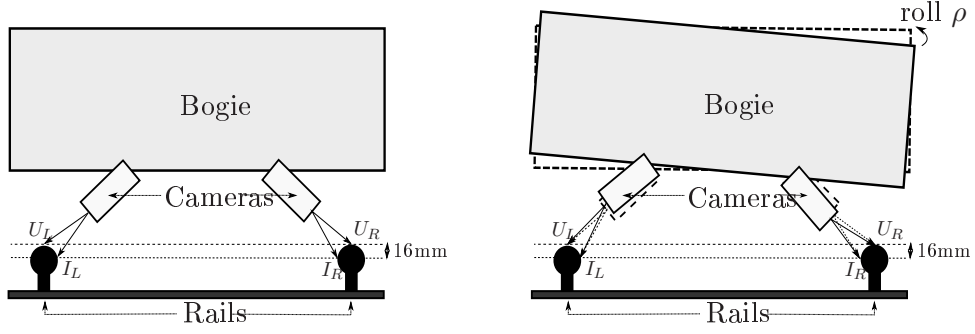


Figure 5.2: Experimental protocol

5.2 Experimental measurements and signal processing

5.2.1 Collection of the experimental inputs for the modeling

Experimental protocol The measurement train IRIS 320 is running continuously since 2007, and monitors the track geometry thanks to two laser cameras, three accelerometers and three rate gyros. Pictures of the track geometry are taken, whereas the accelerometers and rate gyros register the movements of the bogie at the sampling frequency of 10kHz. More precisely, the laser cameras measure the distance toward four particular points of the rails (see Figure 5.2):

- the left and right upper points of the rails U_L and U_R ;
- the left and right interior points I_L and I_R that are placed 16mm under the upper points of the rails.

From these positions, four deviation fields for the rails positions are deduced: d_1 , d_2 , d_3 and d_4 .

Measurements post-processing As the cameras are fixed to one of the bogie of the train, the bogie own movements, which can be characterized by three translations and three rotations, introduce a bias in the measurements, which has to be removed. As an illustration, the three rotation angles of the bogie in a particular curve are represented in Figure 5.3, whereas Figure 5.2 shows the bias induced in the measurements by the roll angle ρ of the bogie. Hence, from a filtering and integration process, the irregularity vector $\mathbf{x}(s) = (x_1(s), x_2(s), x_3(s), x_4(s))$ is deduced from the four data $(d_1(s), d_2(s), d_3(s), d_4(s))$ at each abscissa s of the track (see Figure 5.4). These measurements are also post-processed in order to remove the measurement anomalies, which are mostly due to the absence of signal or to the presence of points and crossings. After these post-treatments, N_{por} track portions of different lengths are available for the modeling.

5.2.2 Local-global approach

In this work, the track irregularity vector of a complete railway track of total length S^{tot} is assumed to be a centered second-order random field,

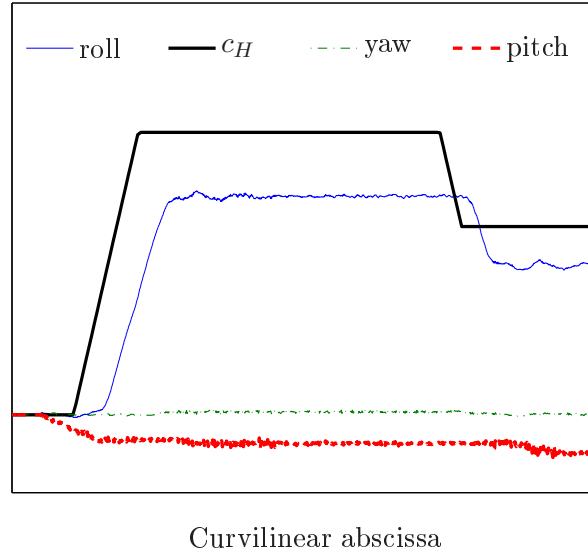


Figure 5.3: Rotation angles of the bogie in a horizontal curve

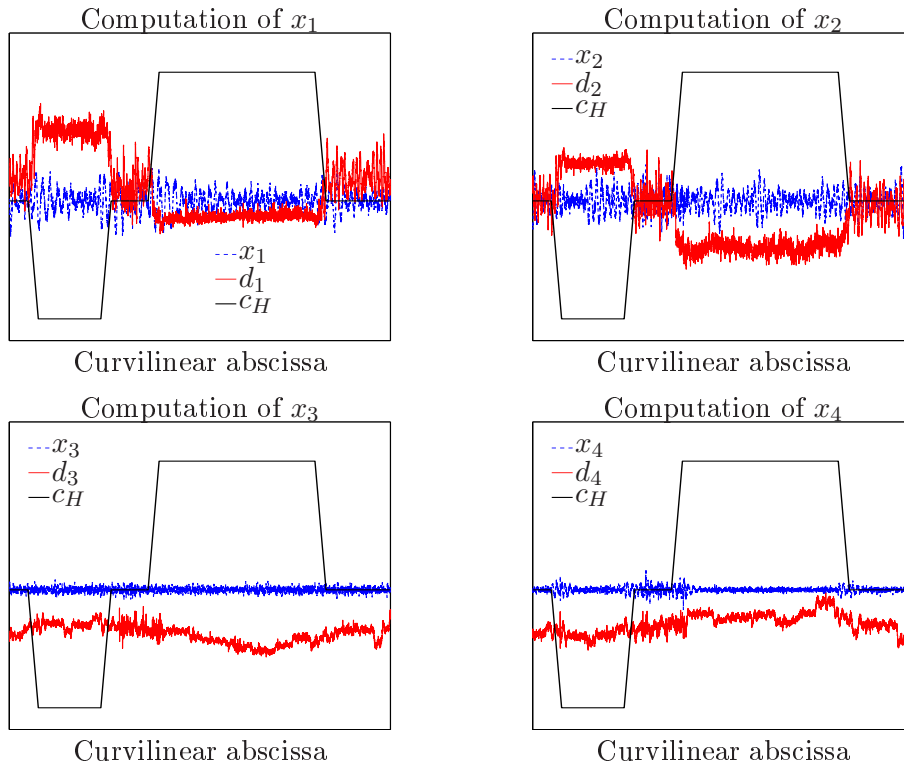


Figure 5.4: Filtering of the experimental data

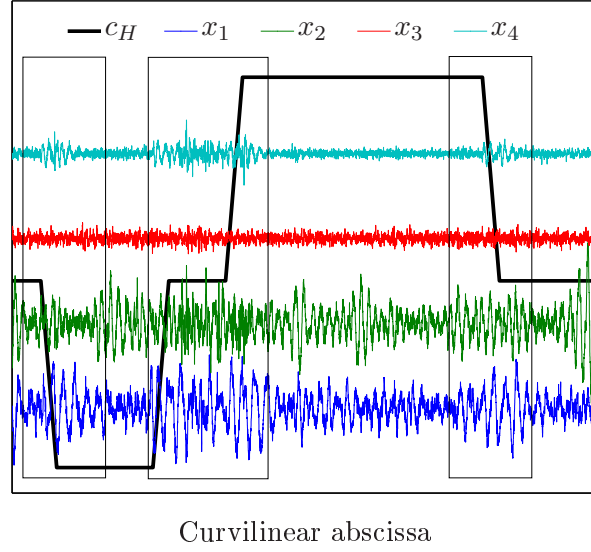


Figure 5.5: Influence of the horizontal curvature on the track irregularities

$$\mathbf{X} = (X_1, X_2, X_3, X_4), \quad (5.6)$$

such that:

$$E[\mathbf{X}(s)] = \mathbf{0}, \quad s \in [0, S^{\text{tot}}]. \quad (5.7)$$

Due to the specific interaction between the train and the track, this random field strongly depends on the horizontal curvature (the influence of the vertical curvature is negligible in terms of track geometry and will not be discussed in the following) and thus on the direction of circulation, as it can be seen in Figure 5.5 (for a better visualization, the four centered irregularity fields have been translated on purpose in this figure). This random field is therefore non-stationary.

Moreover, 200,000 particular values of X_1, X_2, X_3, X_4 are randomly chosen among the available measurements of these four irregularities. Four empirical estimations of the PDFs of X_1, X_2, X_3, X_4 , which are denoted by $\hat{p}_{X_1}, \hat{p}_{X_2}, \hat{p}_{X_3}, \hat{p}_{X_4}$, are then compared over the same closed domain $[\text{LB}, \text{UB}]$ to the corresponding Gaussian PDFs $\mathcal{N}(0, \hat{\sigma}_{X_1}), \mathcal{N}(0, \hat{\sigma}_{X_2}), \mathcal{N}(0, \hat{\sigma}_{X_3}), \mathcal{N}(0, \hat{\sigma}_{X_4})$, in Figure 5.6. From these experimental observations, the track irregularity random field is not Gaussian.

This motivates the introduction of a local-global approach for the characterization of the distribution of the track irregularity random field. This local-global approach is based on the hypothesis that a whole railway track can be considered as the concatenation of a series of independent track portions of same length S , for which physical and statistical properties are the same. Length S plays therefore a key role in the modeling procedure, and its value has to be carefully evaluated. In order to choose length S such that its sensitivity on the stochastic modeling is minimized, ν track portions of same length L , $\{z^{(1)}, \dots, z^{(\nu)}\}$, have been collected from the available measurements of the railway network of interest. For any value of S in $[0, L]$, we denote by $\{y^{(1)}(S), \dots, y^{(\nu)}(S)\}$ the ν new track geometries of total length L , which are then built from the concatenation of track subsections of length S that have been randomly chosen in $\{z^{(1)}, \dots, z^{(\nu)}\}$.

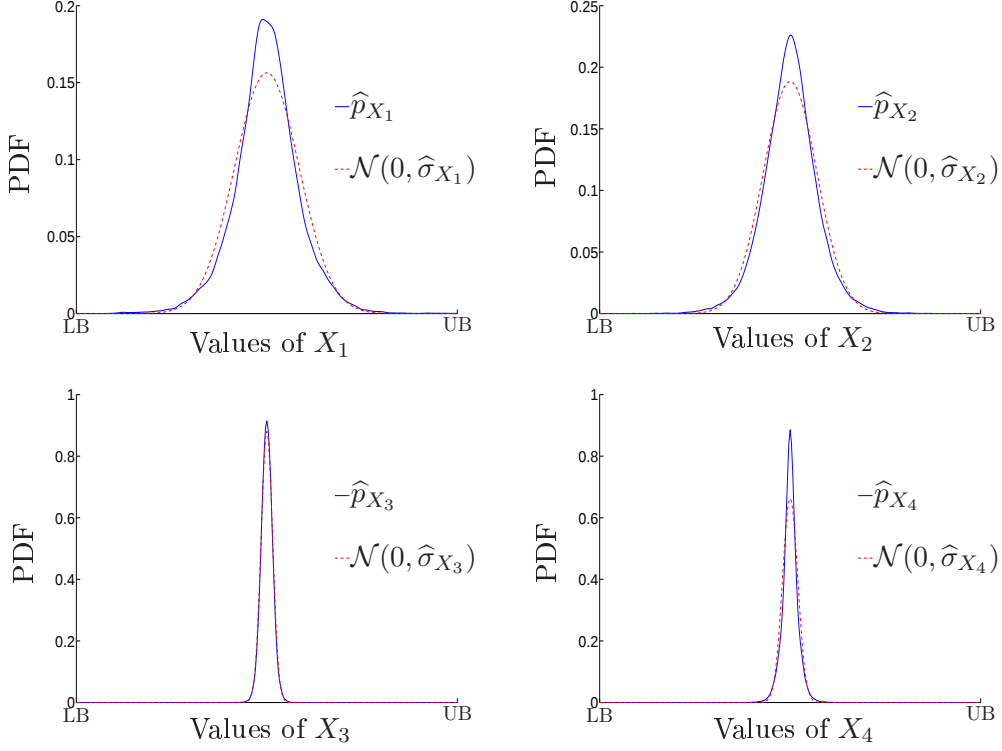


Figure 5.6: Analysis of the marginal distributions of the four track irregularities

Three errors functions, for which definitions can be found in Appendix C, are therefore introduced in this work to quantify the influence of S :

- a covariance error, $err_{\text{cov}}^2(S)$: fixing S to a particular value amounts to supposing that for $|s - s'| \geq S$, $E[\mathbf{X}(s)\mathbf{X}(s')^T]$ is negligible;
- a spectral error, $err_{\text{spect}}^2(S)$: generating complete track geometries from the concatenation of several track portions of length S introduces an artificial periodicity and is likely to degrade the low-frequency characterization of $\mathbf{X}$;
- an estimation error, $err_{\text{est}}^2(S)$: the higher S is, the smaller the number of independent realizations for $\mathbf{X}$, $\nu^{\text{exp}}(S)$, can be extracted from the complete measurements of the railway network of interest. This error is therefore directly related to the estimation accuracy of the covariance function of $\mathbf{X}$, and to the identification precision of the PCE coefficients, on which the modeling will then be based. With reference to the Central Limit Theorem (see [15] for further details), we simply choose $err_{\text{est}}^2(S) = 1/\sqrt{\nu^{\text{exp}}(S)}$ to illustrate this phenomenon.

For the chosen railway network, based on these sets of track geometries of same lengths L , errors $err_{\text{cov}}^2(S)$, $err_{\text{spect}}^2(S)$ and $err_{\text{est}}^2(S)$ are represented in Figure 5.7. When S increases, it can be verified that $err_{\text{cov}}^2(S)$ and $err_{\text{spect}}^2(S)$ decrease whereas $err_{\text{est}}^2(S)$ increases. Length S has thus to be chosen as the right balance between these three error functions.

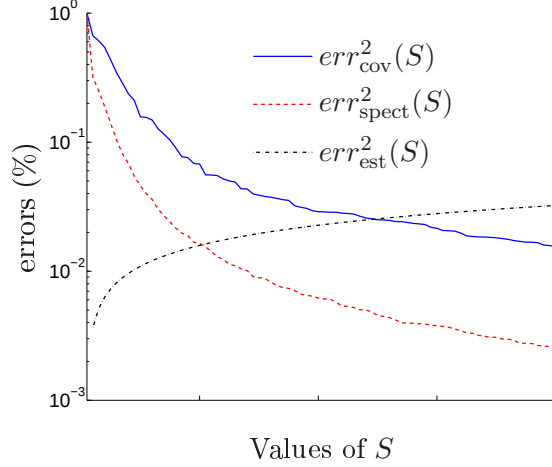


Figure 5.7: Graphs of errors $err_{\text{cov}}^2(S)$, $err_{\text{spect}}^2(S)$ and $err_{\text{est}}^2(S)$ for the computation of the local-global length S .

For confidentiality reasons, the exact value of S is however not given in this work, and the spatial quantities will be normalized by length S in the following. From the available experimental data, $\nu = 1,889$ track portions of same length S , which are denoted by $\{\mathbf{X}(\theta_n), 1 \leq n \leq \nu\}$, are extracted to represent the maximal available information about $\mathbf{X}$. Based on these experimental measurements, which can be seen as a finite set of independent realizations, the next sections aim at completely parameterizing the track irregularity random field, based on the theoretical development that have been presented in Chapters 2, 3 and 4.

5.3 Optimal reduced basis

The first step of the identification of $\mathbf{X}$ corresponds to a revisited Karhunen-Loève (KL) decomposition. This original decomposition, which is presented in detail in Chapters 2 and 4, makes a point of maximizing the representativeness of the projection basis with respect to the limited available information.

5.3.1 Direct KL expansion and projection biases

Let $\Omega = [0, S]$. Using the same notations as in Chapters 2 and 4, we define $[\hat{R}_{\mathbf{X}\mathbf{X}}(\nu)]$ as the empirical estimator of the covariance of $\mathbf{X}$, which has been computed from the ν available realizations of $\mathbf{X}$, and for all $1 \leq M$, let $\hat{\mathcal{K}}^{(M)} = \{\hat{\mathbf{k}}^m, 1 \leq m \leq M\}$ be the set gathering the M eigenfunctions of highest eigenvalues in the Fredholm problem associated with $[\hat{R}_{\mathbf{X}\mathbf{X}}(\nu)]$. The approximation $(s, s') \mapsto [\hat{R}_{\mathbf{X}\mathbf{X}}(\nu, s, s')]_{11}$ of $(s, s') \mapsto E[X_1(s) \otimes X_1(s')]$ is shown in Figure 5.8. This figure emphasizes a quasi symmetry along the first bisector. The functions $s \mapsto [\hat{R}_{\mathbf{X}\mathbf{X}}(\nu, s, 0)]_{qp}, 1 \leq q, p \leq 4$, can thus be used to condense and compare the covariance information of different track irregularities. In Figure 5.8, it can thus be noticed that the covariance matrices are very different from one track irregularity to another.

In addition, we denote by $\varepsilon_q^2(\hat{\mathcal{K}}^{(M)})$ the normalized local projection error such that for $1 \leq q \leq 4$:

$$\varepsilon_q^2(\hat{\mathcal{K}}^{(M)}) \stackrel{\text{def}}{=} \left\| X_q - \hat{X}_q^{\mathcal{K}^{(M)}} \right\|_{\mathcal{P}(\Omega)}^2 / \|X_q\|_{\mathcal{P}(\Omega)}^2, \quad (5.8)$$

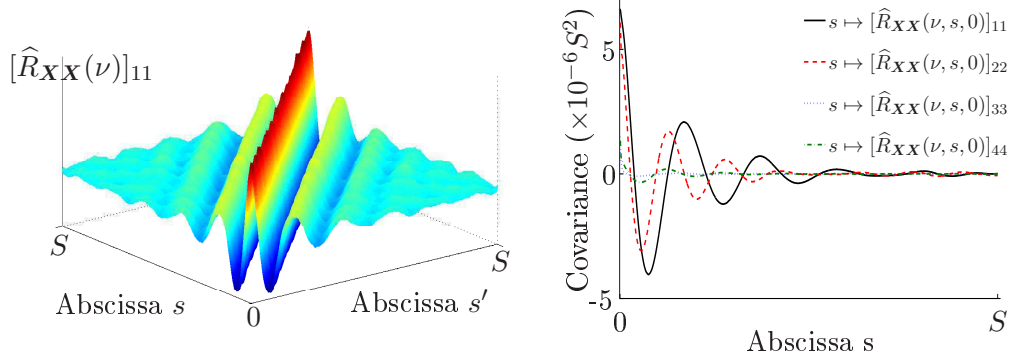


Figure 5.8: Graphs of projections of $[\hat{R}_{XX}(\nu)]$.

where $\hat{X}_q^{\mathcal{K}^{(M)}}$ is the projection of X_q on $\{\hat{\mathcal{K}}_q^m, 1 \leq m \leq M\}$ and where $\|\cdot\|_{\mathcal{P}(\Omega)}^2$ is the norm defined by Eq. (4.3). From Eq. (4.5), the total normalized mean-squared error associated with the projection of $\mathbf{X}$ on $\hat{\mathcal{K}}^{(M)}$, $\varepsilon^2(\hat{\mathcal{K}}^{(M)})$, verifies therefore:

$$\varepsilon^2(\hat{\mathcal{K}}^{(M)}) = \sum_{q=1}^4 \beta_q^2 \varepsilon_q^2(\hat{\mathcal{K}}^{(M)}), \quad (5.9)$$

$$\beta_q^2 = \frac{\|X_q\|_{\mathcal{P}(\Omega)}^2}{\|\mathbf{X}\|_{\mathcal{P}^{(4)}(\Omega)}^2}.$$

From the available realizations of $\mathbf{X}$, we have:

$$\beta_1^2 > \beta_2^2 \gg \beta_4^2 > \beta_3^2. \quad (5.10)$$

The signal energy associated with each track irregularity being different, as shown in Chapter 4, projection family $\hat{\mathcal{K}}^{(M)}$ is bound to describe in priority irregularities X_1 and X_2 rather than X_3 and X_4 . The phenomenon is shown in Figure 5.9, where the evolutions of the LOO estimations, $\varepsilon_{LOO}^2(\hat{\mathcal{K}}^{(M)})$, $\varepsilon_{q,LOO}^2(\hat{\mathcal{K}}^{(M)})$, of the total and local mean-square errors are represented with respect to the size M . In particular, for $M = 500$ and $M = 2000$, although $\varepsilon_{LOO}^2(\hat{\mathcal{K}}^{(500)}) = 10.0\%$ and $\varepsilon_{LOO}^2(\hat{\mathcal{K}}^{(2000)}) = 1.26\%$, we have:

$$\begin{cases} \varepsilon_{1,LOO}^2(\hat{\mathcal{K}}^{(500)}) = 3.52\%, \\ \varepsilon_{2,LOO}^2(\hat{\mathcal{K}}^{(500)}) = 7.77\%, \\ \varepsilon_{3,LOO}^2(\hat{\mathcal{K}}^{(500)}) = 40.3\%, \\ \varepsilon_{4,LOO}^2(\hat{\mathcal{K}}^{(500)}) = 33.0\%, \end{cases} \quad \begin{cases} \varepsilon_{1,LOO}^2(\hat{\mathcal{K}}^{(2000)}) = 0.816\%, \\ \varepsilon_{2,LOO}^2(\hat{\mathcal{K}}^{(2000)}) = 0.615\%, \\ \varepsilon_{3,LOO}^2(\hat{\mathcal{K}}^{(2000)}) = 6.10\%, \\ \varepsilon_{4,LOO}^2(\hat{\mathcal{K}}^{(2000)}) = 3.09\%. \end{cases} \quad (5.11)$$

5.3.2 Optimization of the projection basis

As presented in Chapters 2 and 4, two kinds of improvement can be brought to enrich the direct KL projection family associated with $[\hat{R}_{XX}(\nu)]$. Indeed, in the railway community, the cant deficiency X_3 and the gauge irregularity X_4 are generally considered as the most dangerous irregularities, and are therefore more carefully monitored and maintained. Their signal energy

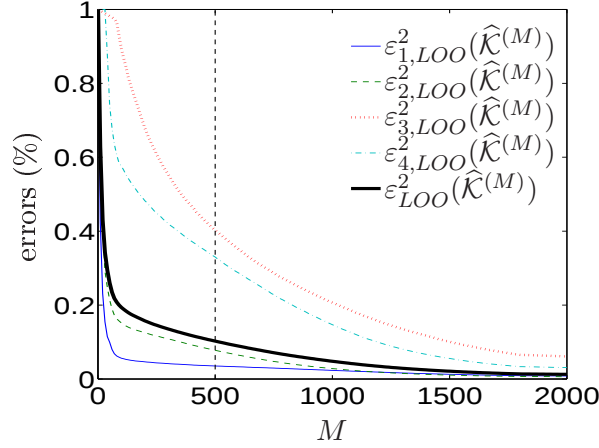


Figure 5.9: Evolution of the projection errors with respect to the dimension M of the projection family.

is lower than the signal energy of the two other track irregularities X_1 and X_2 , although their importance on the train dynamics is likely to be higher. Hence, a scaled expansion can be considered to avoid the numerical bias that is introduced by the differences in the signal energies of the components of $\mathbf{X}$. Then, as the available information about $\mathbf{X}$ is limited to a set of independent realizations, the true covariance function of $\mathbf{X}$ is unknown, and the extension to solve the classical Fredholm equation that is presented in Chapter 2 could allow us to improve the relevance of $\hat{\mathcal{K}}^{(M)}$ to characterize $\mathbf{X}$ for a given value of M .

In this prospect, generalizing the notations of Chapters 2 and 4, for $\mathbf{O}$ in $\mathcal{S}^{(4)}(1)$ and α in $[0, 1]$, $\mathcal{K}(\alpha, \mathbf{O}) = \{\mathbf{k}^m(\alpha, \mathbf{O}), 1 \leq m\}$ is introduced as the orthonormal projection basis that gathers the solutions of the Fredholm problem associated with the $(Q \times Q)$ matrix-valued function $[A(\alpha, \mathbf{O})]$, such that:

$$\int_{\Omega} [A(\alpha, \mathbf{O}, s, s')] \mathbf{k}^m(\alpha, \mathbf{O}, s') ds' = \lambda_m(\alpha, \mathbf{O}) \mathbf{k}^m(\alpha, \mathbf{O}, s), \quad s \in \Omega, \quad (5.12)$$

$$(\mathbf{k}^m(\alpha, \mathbf{O}, s), \mathbf{k}^p(\alpha, \mathbf{O}, s)) = \delta_{mp} \quad \lambda_1(\alpha, \mathbf{O}) \geq \lambda_2(\alpha, \mathbf{O}) \geq \dots \rightarrow 0, \quad (5.13)$$

$$[A(\alpha, \mathbf{O})] = \alpha [\hat{R}_{\mathbf{Y}\mathbf{Y}}(\mathbf{O}, \nu)] + (1 - \alpha) [\tilde{R}_{\mathbf{Y}\mathbf{Y}}(\mathbf{O}, \nu)], \quad (5.14)$$

$$[\hat{R}_{\mathbf{Y}\mathbf{Y}}(\mathbf{O}, \nu)] = [\text{Diag}(\mathbf{O})][\hat{R}_{\mathbf{X}\mathbf{X}}(\nu)][\text{Diag}(\mathbf{O})], \quad (5.15)$$

$$[\tilde{R}_{\mathbf{Y}\mathbf{Y}}(\mathbf{O}, \nu)] = [\text{Diag}(\mathbf{O})][\tilde{R}_{\mathbf{X}\mathbf{X}}(\nu)][\text{Diag}(\mathbf{O})], \quad (5.16)$$

where matrices $[\text{Diag}(\mathbf{O})]$ and $[\tilde{R}_{\mathbf{X}\mathbf{X}}(\nu)]$ are defined by Eqs. (4.15), (2.43) and (2.45). For $M \geq 1$, if $\mathcal{K}^{(M)}(\alpha, \mathbf{O})$ gathers the M first elements of $\mathcal{K}(\alpha, \mathbf{O})$, this leads us to search the optimal projection family for $\mathbf{X}$, $\mathcal{F}_{\text{opt}}^{(M)}$, as the solution of the following optimization problem:

$$\mathcal{F}_{\text{opt}}^{(M)} = \mathcal{K}^{(M)}(\alpha^{\text{opt}}(M), \mathbf{O}^{\text{opt}}(M)), \quad (5.17)$$

$$(\alpha^{\text{opt}}(M), \mathbf{O}^{\text{opt}}(M)) = \arg \min_{(\alpha, \mathbf{O}) \in [0, 1] \times \mathcal{S}^{(4)}(1)} \varepsilon_{\infty, \text{LOO}}^2(\mathcal{K}^{(M)}(\alpha, \mathbf{O})), \quad (5.18)$$

α	$\mathbf{O}$	$\varepsilon_{LOO,1}^2(\hat{\mathcal{K}}^{(500)})$	$\varepsilon_{LOO,2}^2(\hat{\mathcal{K}}^{(500)})$	$\varepsilon_{LOO,3}^2(\hat{\mathcal{K}}^{(500)})$	$\varepsilon_{LOO,4}^2(\hat{\mathcal{K}}^{(500)})$
1	(0.5, 0.5, 0.5, 0.5)	3.52%	7.77%	40.3%	33.3%
1	$\mathbf{O}^\beta$	5.88%	16.3%	17.8%	23.9%
1	$\mathbf{O}^{\text{opt}}(500, \alpha)$	18.3%	18.3%	18.3%	18.3%
$\alpha^{\text{opt}}(500, \mathbf{O})$	(0.5, 0.5, 0.5, 0.5)	2.95%	5.62%	30.2%	25.9%
$\alpha^{\text{opt}}(500, \mathbf{O})$	$\mathbf{O}^\beta$	4.77%	13.7%	12.6%	17.8%
$\alpha^{\text{opt}}(500)$	$\mathbf{O}^{\text{opt}}(500)$	13.9%	13.9%	13.9%	13.9%

Figure 5.10: Influence of the choices for α and $\mathbf{O}$ on the local mean-square errors.

$$\varepsilon_{\infty, LOO}^2(\mathcal{K}^{(M)}(\alpha, \mathbf{O})) = \max_{1 \leq q \leq 4} \varepsilon_{q, LOO}^2(\mathcal{K}^{(M)}(\alpha, \mathbf{O})). \quad (5.19)$$

This problem is solved coupling the iterative algorithm defined by Eq. (2.51) with $\tau = 10^{-4}$ and $\gamma = 1/2$ for $\mathbf{O}$, and an algorithm based on a dichotomy for α . In order to illustrate the advantage of such an approach, Table 5.10 compares the LOO errors associated with particular values of α and $\mathbf{O}$, that stem from optimizations on α and/or $\mathbf{O}$, where $\mathbf{O}^\beta = \left(\frac{1}{\|X_1\|_{\mathcal{P}(\Omega)}}, \frac{1}{\|X_2\|_{\mathcal{P}(\Omega)}}, \frac{1}{\|X_3\|_{\mathcal{P}(\Omega)}}, \frac{1}{\|X_4\|_{\mathcal{P}(\Omega)}} \right)$, $\alpha^{\text{opt}}(500, \mathbf{O})$ is the optimal value of α in $[0, 1]$ for a given value of $\mathbf{O}$, and $\mathbf{O}^{\text{opt}}(500, \alpha)$ is the optimal value of $\mathbf{O}$ in $\mathcal{S}^{(4)}(1)$ for a given value α .

For a same dimension $M = 500$, this double adaptation of the classical KL expansion for the track irregularity random field allows us to divide the maximal value of the local mean square errors by three. Richer definitions for $[A(\alpha, \mathbf{O})]$ should lead us to even better results, but to do so in very high dimension with very limited information, as it is the case here, a method to optimize the solving of Eq. (5.18) would be required, which has not been made in this thesis.

5.3.3 Choice of the dimension of the spatial projection parameter

For any value of M , the optimization problem defined by Eq. (5.18) allows us to identify projection basis that are particularly well adapted to each component of $\mathbf{X}$. The optimal value of M can therefore be searched with respect to a chosen threshold for $\varepsilon_{\infty, LOO}^2(\mathcal{K}^{(M)}(\alpha^{\text{opt}}(M), \mathbf{O}^{\text{opt}}(M)))$. In the following, for $M = 2000$, we denote by $\mathcal{F}_{\text{opt}}^{(2000)} = \{\mathbf{f}^m, 1 \leq m \leq 2000\}$ this basis, which allows the maximal value of the local errors, $\varepsilon_{\infty, LOO}^2(\mathcal{K}^{(2000)}(\alpha^{\text{opt}}(2000), \mathbf{O}^{\text{opt}}(2000)))$, to be lower than 0.5%. From Eq. (5.11), it can be noticed that thanks to the two proposed adaptations of the KL expansion, the four local errors associated with $\mathcal{F}_{\text{opt}}^{(2000)}$ are lower than the minimal value of the local errors associated with $\hat{\mathcal{K}}^{(2000)}$. Such a high value for M will be justified more in detail in Chapter 6 from the train dynamics analysis. We compare in Figure 5.11 several graphs of eigenfunctions $\mathbf{f}^m = (f_1^m, f_2^m, f_3^m, f_4^m)$. For a better visualization, the mean values of the different subvectors, which are zero, are deliberately translated.

5.4 PCE identification in very high dimension

5.4.1 Sorting with respect to the horizontal curvature

As presented in the Section 5.2.2, the track geometry strongly depends on the horizontal curvature. In this prospect, four classes of track portions can be introduced:

- the **alignment**, for which the horizontal curvature c_H is zero.

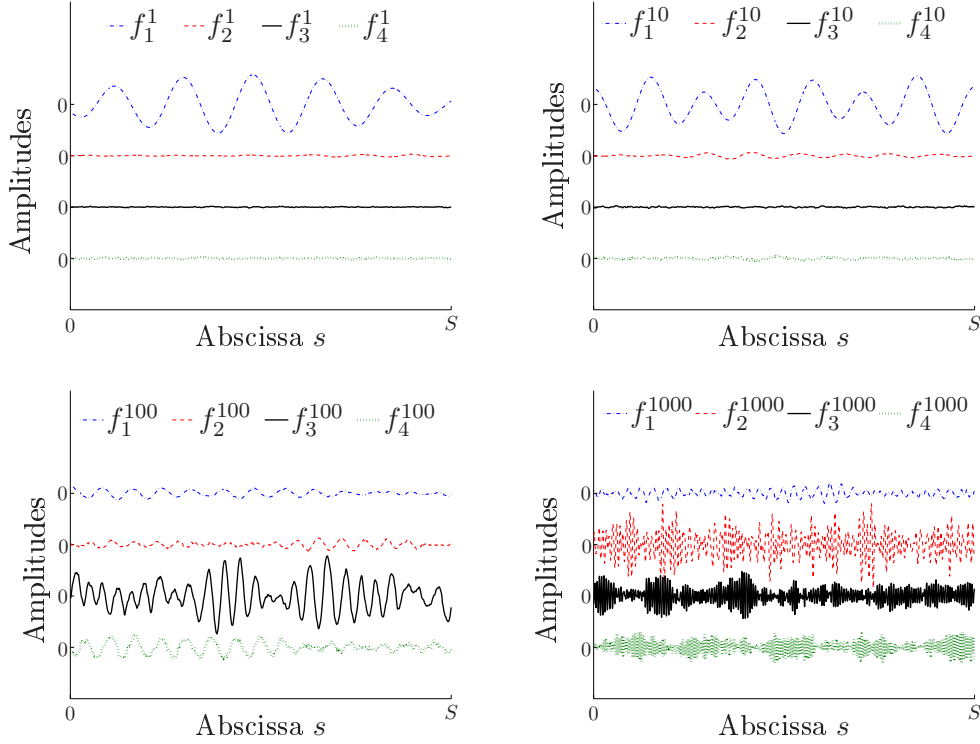


Figure 5.11: Graphs of four particular eigenfunctions $\mathbf{f}^1$, $\mathbf{f}^{10}$, $\mathbf{f}^{100}$ and $\mathbf{f}^{1000}$ ($\times 10^{-3}S$).

- the **established curve**, for which the horizontal curvature c_H is constant and not zero.
- the **curve entrance**, for which the absolute value of the curvature is linearly increasing.
- the **curve exit**, for which the absolute value of the curvature is linearly decreasing.

The statistical properties of $\mathbf{X}$ are therefore different in each of these four classes, such that instead of one stochastic modeling, four stochastic modelings of the track geometry over a length S are needed to accurately characterize the track geometry variabilities. We thus denote by $\mathbf{X}^{(A)}$ (in alignment), $\mathbf{X}^{(EC)}$ (in curve entrance), $\mathbf{X}^{(C)}$ (in curve) and $\mathbf{X}^{(SC)}$ (in curve exit) the four projections of random field $\mathbf{X}$ in the different curvature classes. These four random fields can then be projected on the 2,000-dimension deterministic family $\mathcal{F}_{\text{opt}}^{(2000)}$, which was introduced in the previous section, such that:

$$\left\{ \begin{array}{l} \mathbf{X}^{(A)} \approx \widehat{\mathbf{X}}^{(A)} = \sum_{m=1}^{2000} C_m^{(A)} \mathbf{f}^m, \\ \mathbf{X}^{(EC)} \approx \widehat{\mathbf{X}}^{(EC)} = \sum_{m=1}^{2000} C_m^{(EC)} \mathbf{f}^m, \\ \mathbf{X}^{(C)} \approx \widehat{\mathbf{X}}^{(C)} = \sum_{m=1}^{2000} C_m^{(C)} \mathbf{f}^m, \\ \mathbf{X}^{(SC)} \approx \widehat{\mathbf{X}}^{(SC)} = \sum_{m=1}^{2000} C_m^{(SC)} \mathbf{f}^m. \end{array} \right. \quad (5.20)$$

Finally, characterizing the variability of the track geometry amounts to identifying the multi-dimensional distributions of the four 2000-dimension random vectors, $\mathbf{C}^{(A)} = (C_1^{(A)}, \dots, C_{2000}^{(A)})$, $\mathbf{C}^{(EC)} = (C_1^{(EC)}, \dots, C_{2000}^{(EC)})$, $\mathbf{C}^{(C)} = (C_1^{(C)}, \dots, C_{2000}^{(C)})$ and $\mathbf{C}^{(SC)} = (C_1^{(SC)}, \dots, C_{2000}^{(SC)})$, for which components are dependent.

Independent realizations of these four random vectors have to be extracted from the sorting of the N_{por} available measurements with respect to the horizontal curvature. This sorting is based on a four-step method, which is illustrated in Figures 5.12 and 5.13:

- First, the true horizontal curvature, c_H , which is piecewise linear, is deduced from the on-track measured horizontal curvature, $c_H^{\text{on track}}$.
- Secondly, the positions of the beginnings and the ends of the curvature classes are localized.
- Then, for each curvature class, a series of measurements of same length S is extracted, and is denoted by $\mathbf{x}_A^i$, $\mathbf{x}_C^j$, $\mathbf{x}_{EC}^k$, $\mathbf{x}_{SC}^\ell$ for the alignment, the curve, the curve entrance and the curve exit cases. The length of the curve entrances and exits being generally lower than S , an overlapping is tolerated, such that some small track portions can be used in two modelings. Under the local-global hypothesis, $\nu_A = 414$, $\nu_{EC} = 482$, $\nu_C = 522$ and $\nu_{SC} = 471$ track portions of same length S are extracted from the complete railway network of total length S^{tot} . These measurements are supposed to be independent realizations of the random fields $\mathbf{X}^{(A)}$, $\mathbf{X}^{(EC)}$, $\mathbf{X}^{(C)}$ and $\mathbf{X}^{(SC)}$ respectively.
- Finally, these realizations are projected on $\mathcal{F}_{\text{opt}}^{(2000)}$ to compute the corresponding realizations of $\mathbf{C}^{(A)}$, $\mathbf{C}^{(EC)}$, $\mathbf{C}^{(C)}$ and $\mathbf{C}^{(SC)}$.

The same approach will be used to identify the distributions of these four random vectors, but only the identification of $\mathbf{C}^{(A)}$ will be presented in the following.

5.4.2 PCE identification

Let $\{\mathbf{C}^{(A)}(\theta_1), \dots, \mathbf{C}^{(A)}(\theta_{\nu_A})\}$ be the ν_A available realizations of random vector $\mathbf{C}^{(A)}$. The mean value of random field $\mathbf{X}^{(A)}$ being zero for all s in Ω , this vector is also centered, such that:

$$E[\mathbf{C}^{(A)}] = \mathbf{0}. \quad (5.21)$$

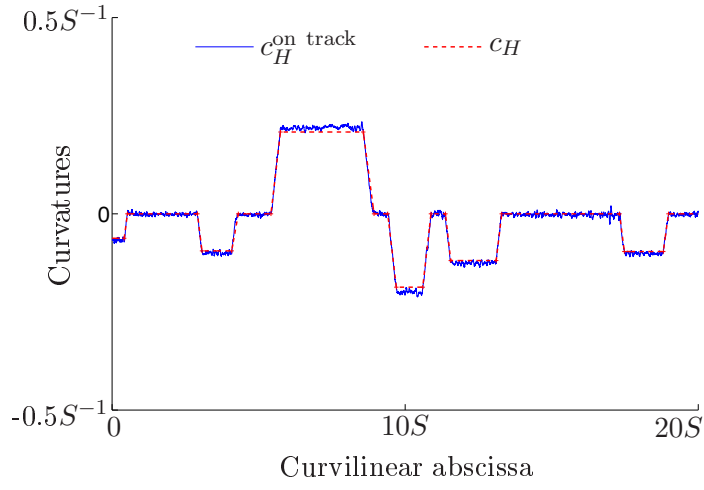


Figure 5.12: Extraction of the true horizontal curvature from the in line measurements.

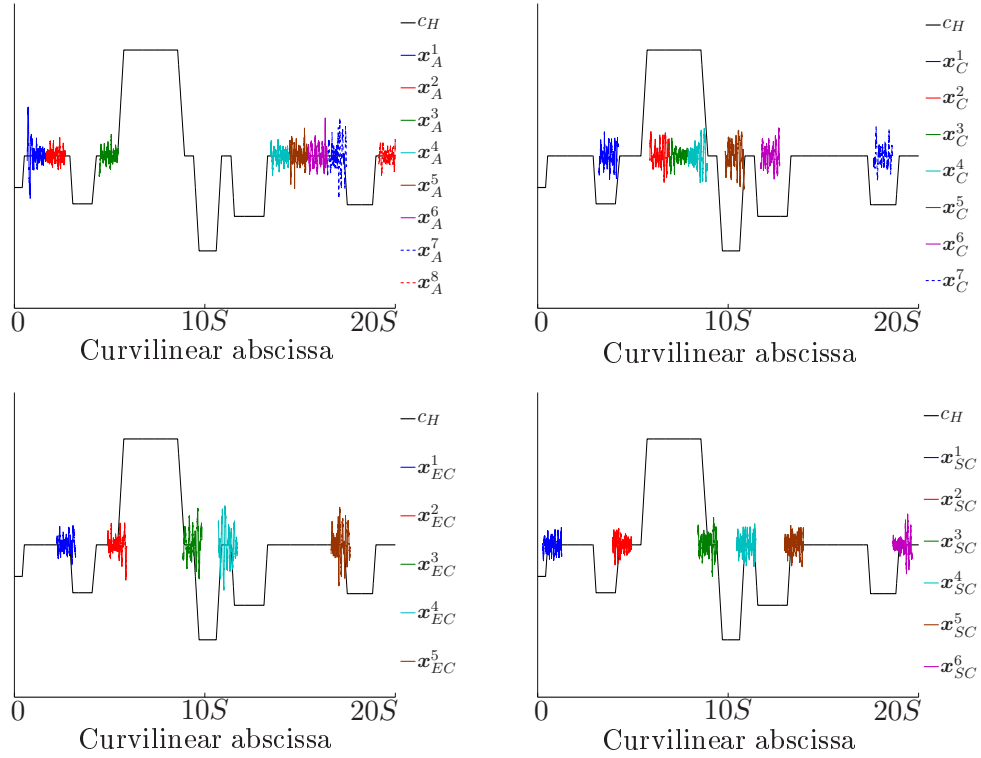


Figure 5.13: Extraction of the track irregularities for each curvature class

Moreover, from these ν_A realizations, the covariance matrix of $\mathbf{C}^{(A)}$, which is written $[\widehat{R}_{CC}^{(A)}(\nu_A)]$, can be estimated as:

$$[\widehat{R}_{CC}^{(A)}(\nu_A)] = \frac{1}{\nu_A} \sum_{n=1}^{\nu_A} \mathbf{C}^{(A)}(\theta_n) \otimes \mathbf{C}^{(A)}(\theta_n). \quad (5.22)$$

Given this information, the multidimensional distribution of $\mathbf{C}^{(A)}$ is identified from a PCE approach. In this prospect, let $\boldsymbol{\xi} = (\xi_1, \dots, \xi_{N_g})$ be a random vector for which components are independent and uniformly distributed between -1 and 1. A uniform germ for the PCE is chosen, as it appears to be more stable in very high polynomial dimensions as shown in Section 3.3.2. The corresponding Hilbertian basis of all N_g -dimension random vectors is the set of the multidimensional Legendre polynomials, which are denoted by $\{\psi_j(\boldsymbol{\xi}), 1 \leq j\}$. In agreement with the theoretical developments of Section 1.5, this basis is truncated to its N elements of total polynomial order lower than p . At last, we define $\mathbf{C}^{\text{chaos},(A)}(N)$ as the projection of $\mathbf{C}^{(A)}$ on this truncated basis, such that:

$$\mathbf{C}^{(A)} \approx \mathbf{C}^{\text{chaos},(A)}(N) = \sum_{j=1}^N \mathbf{y}^{j,(A)} \psi_j(\boldsymbol{\xi}) = [\mathbf{y}^{(A)}] \boldsymbol{\Psi}(\boldsymbol{\xi}). \quad (5.23)$$

For given values of N_g and N , the projection matrix $[\mathbf{y}^{(A)}]$ is computed to maximize the likelihood of $[\mathbf{y}^{(A)}] \boldsymbol{\Psi}(\boldsymbol{\xi})$ at the independent realizations of $\mathbf{C}^{(A)}$ under the approximated constraint $[\mathbf{y}^{(A)}][\mathbf{y}^{(A)}]^T \approx [\widehat{R}_{CC}^{(A)}(\nu_A)]$ with the iterative algorithm described in Section 3.2.5, as the dimension of $\mathbf{C}^{(A)}$ is much higher than ν_A . According to Figure 5.14, where $err(N, N_g)$ is plotted as a function of N for different values of N_g , truncation parameters N_g and N are chosen equal to 3 and 2,925 respectively, which corresponds to the maximal polynomial order $p = 24$ for the reduced polynomial basis. Moreover, Figure 5.15 compares the empirical estimations, $\widehat{p}_{C_m^{(A)}}$ and $\widehat{p}_{C_m^{\text{chaos},(A)}(N)}$, of the PDFs of three particular components of $\mathbf{C}^{(A)}$ and $\mathbf{C}^{\text{chaos},(A)}(N)$ respectively. A normal PDF associated with the variance of $C_m^{(A)}$, $\widehat{p}_{C_m^{(A)}}^{\text{Gauss}}$, has also been added to this figures.

Following exactly the same approaches, the three projection matrices $[\mathbf{y}^{(EC)}]$, $[\mathbf{y}^{(C)}]$ and $[\mathbf{y}^{(SC)}]$ are also identified. The convergence analysis for these expansions has moreover given the same results as for the alignment case, such that we get:

$$\mathbf{C}^{(EC)} = [\mathbf{y}^{(EC)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \quad \mathbf{C}^{(C)} = [\mathbf{y}^{(C)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \quad \mathbf{C}^{(SC)} = [\mathbf{y}^{(SC)}] \boldsymbol{\Psi}(\boldsymbol{\xi}). \quad (5.24)$$

5.5 Generation of a whole track geometry

Once truncation parameters M , N , N_g have been identified according to convergence analysis, spatial projection family $\mathcal{F}^{(2000)} = \{\mathbf{f}^m, 1 \leq m \leq 2000\}$ has been computed, and PCE projection matrices $[\mathbf{y}^{(A)}]$, $[\mathbf{y}^{(EC)}]$, $[\mathbf{y}^{(C)}]$ and $[\mathbf{y}^{(SC)}]$ have been calculated, the track irregularity random field is completely characterized:

$$\begin{cases} \mathbf{X}^{(A)} \approx \widetilde{\mathbf{X}}^{(A)} = [F^{(2000)}][\mathbf{y}^{(A)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \\ \mathbf{X}^{(EC)} \approx \widetilde{\mathbf{X}}^{(EC)} = [F^{(2000)}][\mathbf{y}^{(EC)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \\ \mathbf{X}^{(C)} \approx \widetilde{\mathbf{X}}^{(C)} = [F^{(2000)}][\mathbf{y}^{(C)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \\ \mathbf{X}^{(SC)} \approx \widetilde{\mathbf{X}}^{(SC)} = [F^{(2000)}][\mathbf{y}^{(SC)}] \boldsymbol{\Psi}(\boldsymbol{\xi}), \end{cases} \quad (5.25)$$

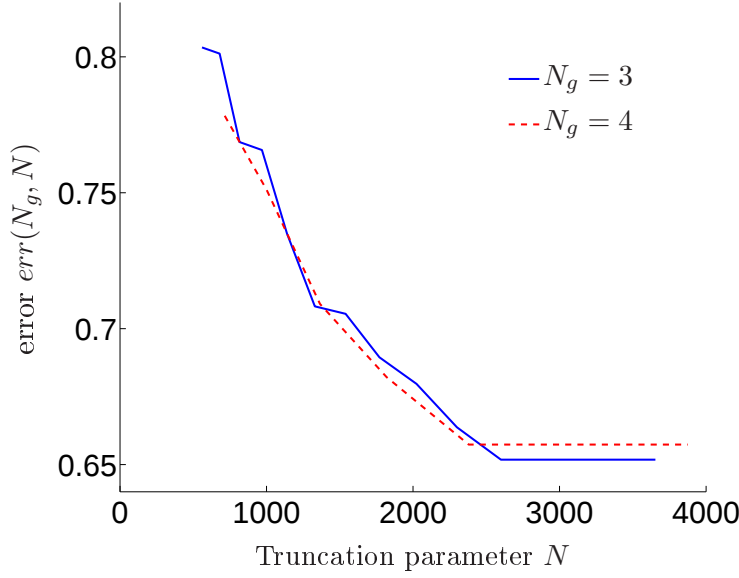


Figure 5.14: Identification of the PCE truncation parameters.

$$[F^{(2000)}] = [\mathbf{f}^1 \ \mathbf{f}^2 \ \dots \ \mathbf{f}^{2000}]. \quad (5.26)$$

For each realization of random vector $(\xi_1, \dots, \xi_{N_g})$, a representative and realistic track geometry of length S can be generated for the alignment, the established curve, the curve entrance and the curve exit cases. Thanks to the local-global approach, described in Section 5.2.2, a whole track geometry of length $S^{\text{tot}} = N_{\mathcal{T}}S$, $\mathbf{X}^{\text{tot}}$, can therefore be constructed from the concatenation of $N_{\mathcal{T}}$ independent copies $\mathbf{X}^{(1)}, \dots, \mathbf{X}^{(N_{\mathcal{T}})}$ of the track irregularity vectors $\widetilde{\mathbf{X}}^{(A)}$, $\widetilde{\mathbf{X}}^{(EC)}$, $\widetilde{\mathbf{X}}^{(C)}$ or $\widetilde{\mathbf{X}}^{(SC)}$, with respect to the horizontal curvature of the considered track, such that $\mathbf{X}^{\text{tot}} = (\mathbf{X}^{(1)}, \dots, \mathbf{X}^{(N_{\mathcal{T}})})$.

Therefore, ν independent realizations $\{\mathbf{X}^{\text{tot}}(\theta_1), \dots, \mathbf{X}^{\text{tot}}(\theta_\nu)\}$ of $\mathbf{X}^{\text{tot}}$ can be generated from $\nu N_{\mathcal{T}}$ realizations of the local irregularity vector $\widetilde{\mathbf{X}}^{(A)}$, $\widetilde{\mathbf{X}}^{(EC)}$, $\widetilde{\mathbf{X}}^{(C)}$ or $\widetilde{\mathbf{X}}^{(SC)}$. However, for each realization $\mathbf{X}^{\text{tot}}(\theta_m)$ of $\mathbf{X}^{\text{tot}}$, a particular attention has to be paid to the junction between these different realizations. Indeed, these junctions have to guarantee the continuity of the track irregularity vector and at least the continuity of its first and second order spatial derivatives in order to avoid an artificial perturbation of the train dynamics. Spline interpolations on a length corresponding to the minimal wavelength of the measured irregularities are then used to fulfill these continuity conditions.

Hence, the proposed stochastic modeling allows us to generate realistic track geometries of length $S^{\text{tot}} = N_{\mathcal{T}}S$ that are representative of the whole network, and which take into account the spatial and statistical dependencies between the different track irregularities. As an illustration, a particular extract of length S of complete track geometry $\mathbf{X}^{\text{tot}}(\theta_1)$ is represented in Figure 5.16. This graph is centered at abscissa $s = 3S/2$, that is to say at a junction between the two first realizations of the track irregularity random fields. The four components of $\mathbf{X}^{\text{tot}}(\theta_1)$ are represented in the same graph, but their values are translated to allow a better visualization of the results.

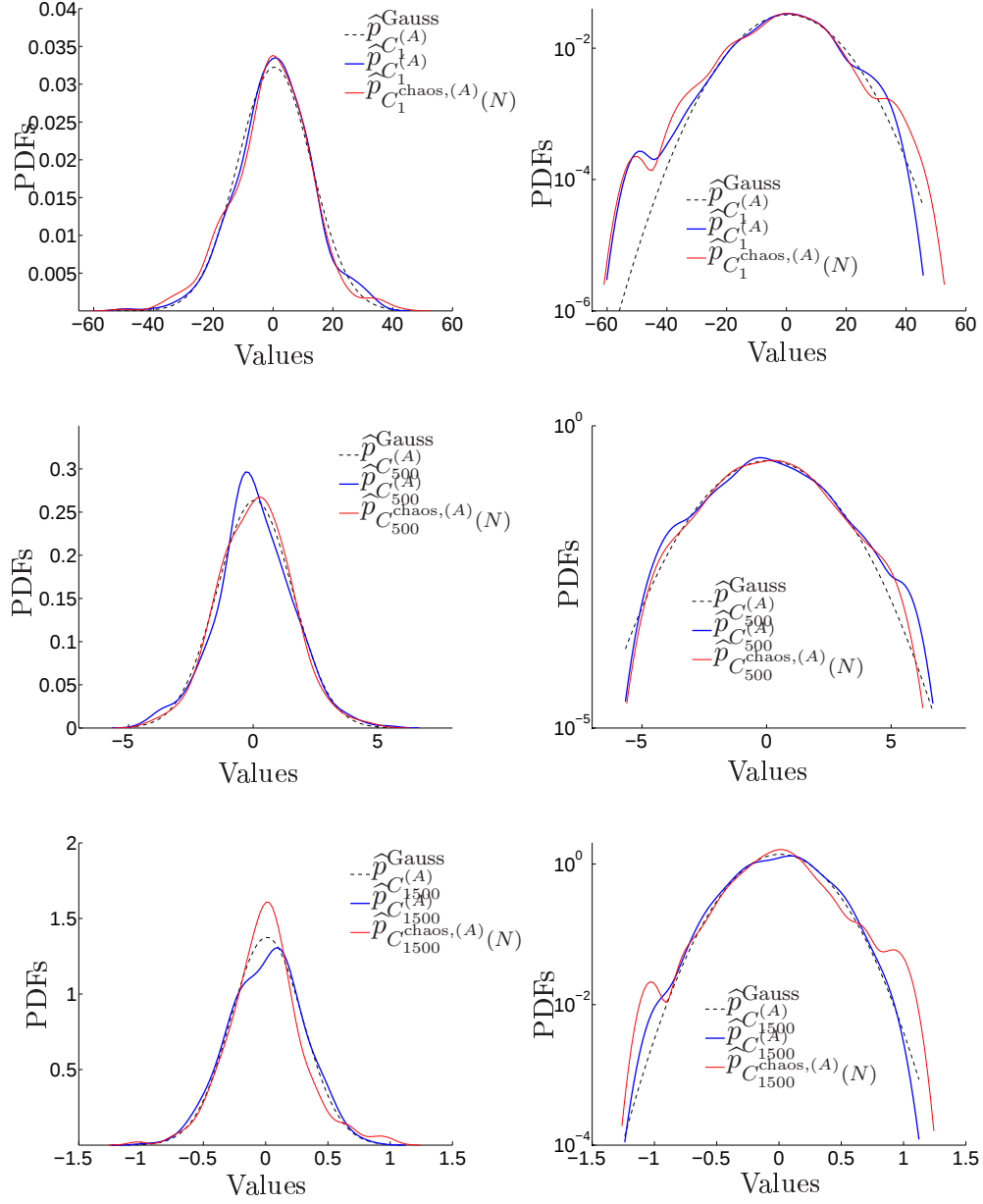


Figure 5.15: Comparisons of the PDFs of three particular components of $C^{(A)}$ and $C^{\text{chaos},(A)}(N)$

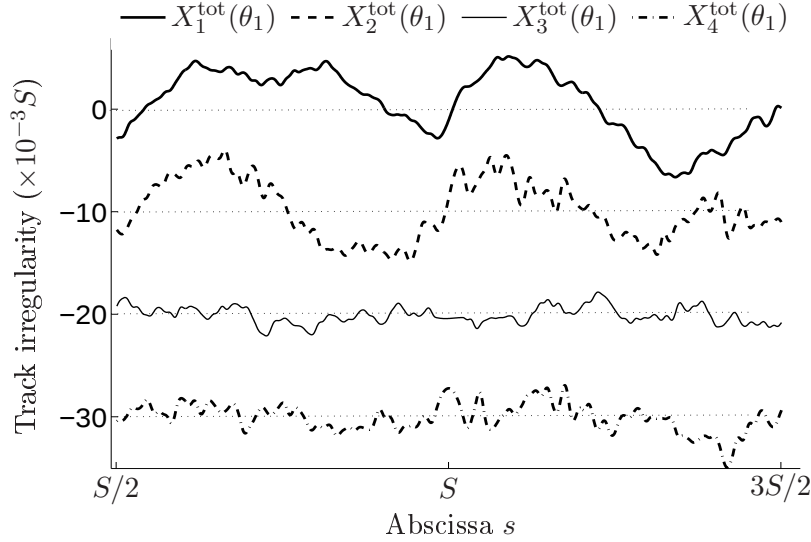


Figure 5.16: Extract of a simulated track geometry.

5.6 Statistical and frequency validations

As presented in Section 5.1, a complete parametrization of the physical and statistical properties of the track geometry are needed in conception, certification and maintenance prospects. Several validations of the proposed stochastic modeling are thus presented in this section, in order to allow a relevant investigation of the dynamic interaction between the train and the track.

In this section, 10^3 independent realizations of track irregularity random fields $\widetilde{\mathbf{X}}^{(A)}$, $\widetilde{\mathbf{X}}^{(EC)}$, $\widetilde{\mathbf{X}}^{(C)}$ and $\widetilde{\mathbf{X}}^{(SC)}$ are generated from the track stochastic modelings developed in Section 5.3. The notations of Section 3.4.3 are adopted again in this section. If X_q corresponds to one of the components of $\mathbf{X}^{(A)}$, $\mathbf{X}^{(EC)}$, $\mathbf{X}^{(C)}$, $\mathbf{X}^{(SC)}$, $\widetilde{\mathbf{X}}^{(A)}$, $\widetilde{\mathbf{X}}^{(EC)}$, $\widetilde{\mathbf{X}}^{(C)}$ or $\widetilde{\mathbf{X}}^{(SC)}$, for which ν (ν is equal to 10^3 , $\nu_A = 414$, $\nu_{EC} = 482$, $\nu_C = 522$ or $\nu_{SC} = 471$) independent realizations are known and denoted by $\{X_q(\theta_n), 1 \leq n \leq \nu\}$, we use $N_{\text{up}}(X_q(\theta_n), u, S)$, $1 \leq q \leq 4$, to denote the numbers of upcrossings of the level u by the n^{th} realization $X_q(\theta_n)$ of X_q over the length S , and we define $\mathcal{D}_i$, $1 \leq i \leq 10$ the domains such that for each level u , $\mathcal{D}_i$ gathers $i/10$ of the values of $\{N_{\text{up}}(X_q(\theta_1), u, S), \dots, N_{\text{up}}(X_p(\theta_\nu), u, S)\}$.

The domains for $\mathbf{X}^{(A)}$, $\mathbf{X}^{(EC)}$, $\mathbf{X}^{(C)}$, $\mathbf{X}^{(SC)}$ are thus compared to contour plots that corresponds to the equivalent domains for $\widetilde{\mathbf{X}}^{(A)}$, $\widetilde{\mathbf{X}}^{(EC)}$, $\widetilde{\mathbf{X}}^{(C)}$ and $\widetilde{\mathbf{X}}^{(SC)}$ in Figures 5.17, 5.18, 5.19 and 5.20. In addition, these results are compared to the cases when the random vectors $\mathbf{C}^{(A)}$, $\mathbf{C}^{(EC)}$, $\mathbf{C}^{(C)}$ and $\mathbf{C}^{(SC)}$ would have been modeled by Gaussian random vectors. To this end, we denote by $\mathbf{X}_{\text{gauss}}^{(A)}$, $\mathbf{X}_{\text{gauss}}^{(EC)}$, $\mathbf{X}_{\text{gauss}}^{(C)}$, $\mathbf{X}_{\text{gauss}}^{(SC)}$ the Gaussian approximations of $\mathbf{X}^{(A)}$, $\mathbf{X}^{(EC)}$, $\mathbf{X}^{(C)}$ and $\mathbf{X}^{(SC)}$ respectively.

In the same manner, for $1 \leq q \leq 4$, let $PSD(X_q)$ be the mean power spectral densities of X_q that have been computed from the available realizations of X_q of length S . The frequency characteristics of $\mathbf{X}^{(A)}$ and $\widetilde{\mathbf{X}}^{(A)}$ (the same results are obtained for the other curvature cases) are therefore compared in Figure 5.21.

The rather good agreement between the quantities corresponding to the measured and to the generated track geometries, especially compared to the Gaussian case, allows us to validate the stochastic modeling over a length S from a statistical and frequency point of view.

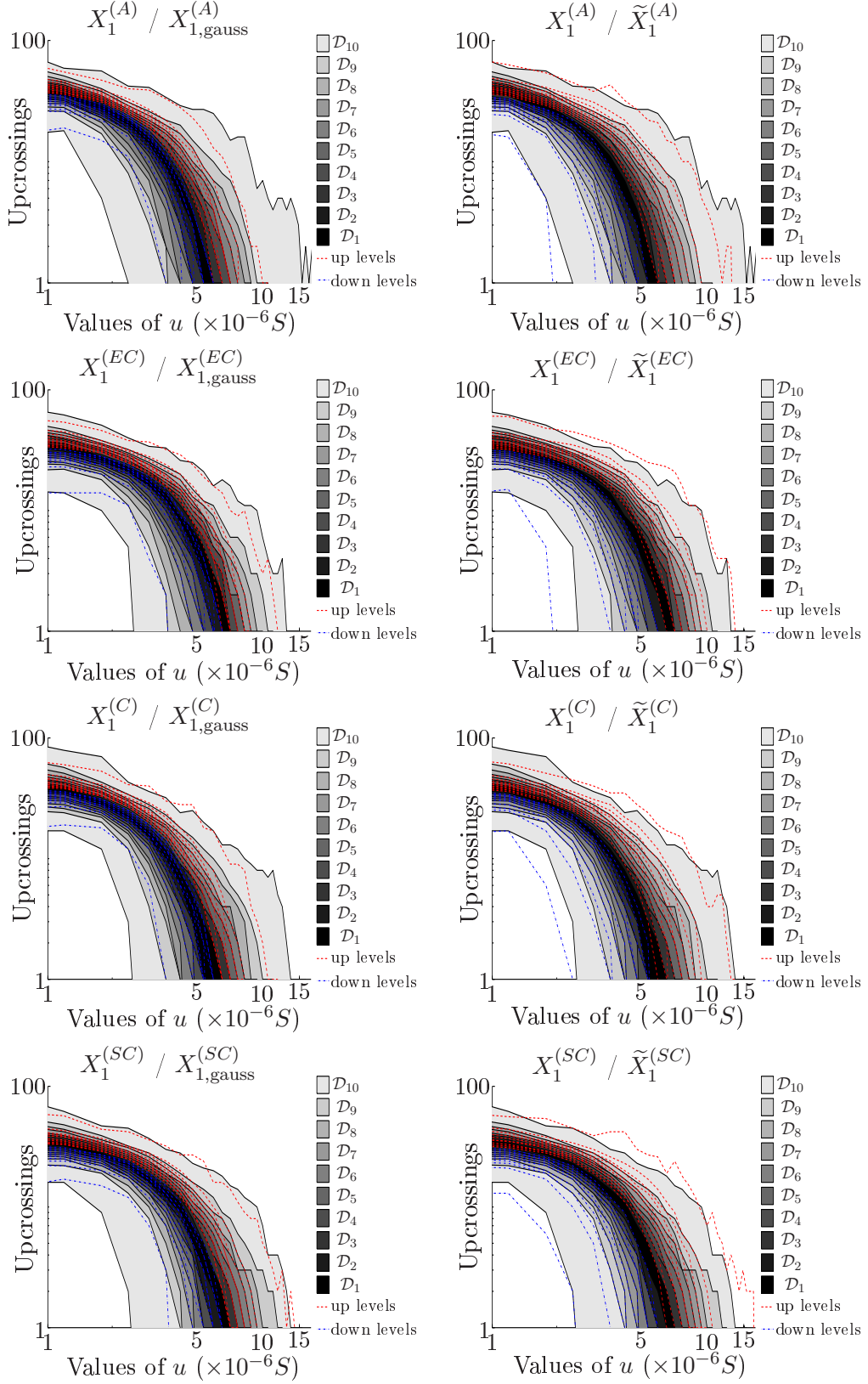


Figure 5.17: Statistical validation of the stochastic modeling of the horizontal alignment irregularity X_1

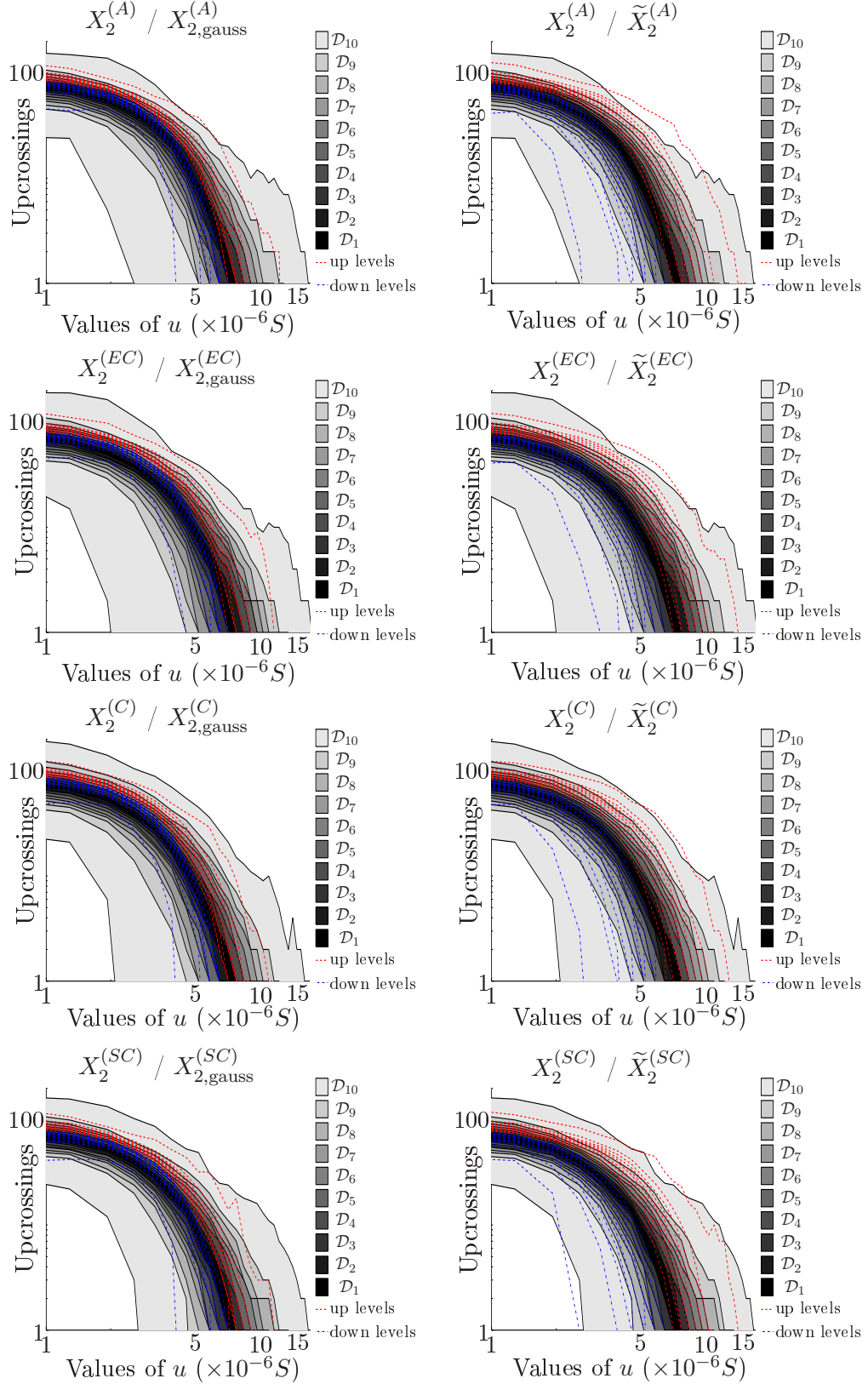


Figure 5.18: Statistical validation of the stochastic modeling of the vertical alignment irregularity X_2

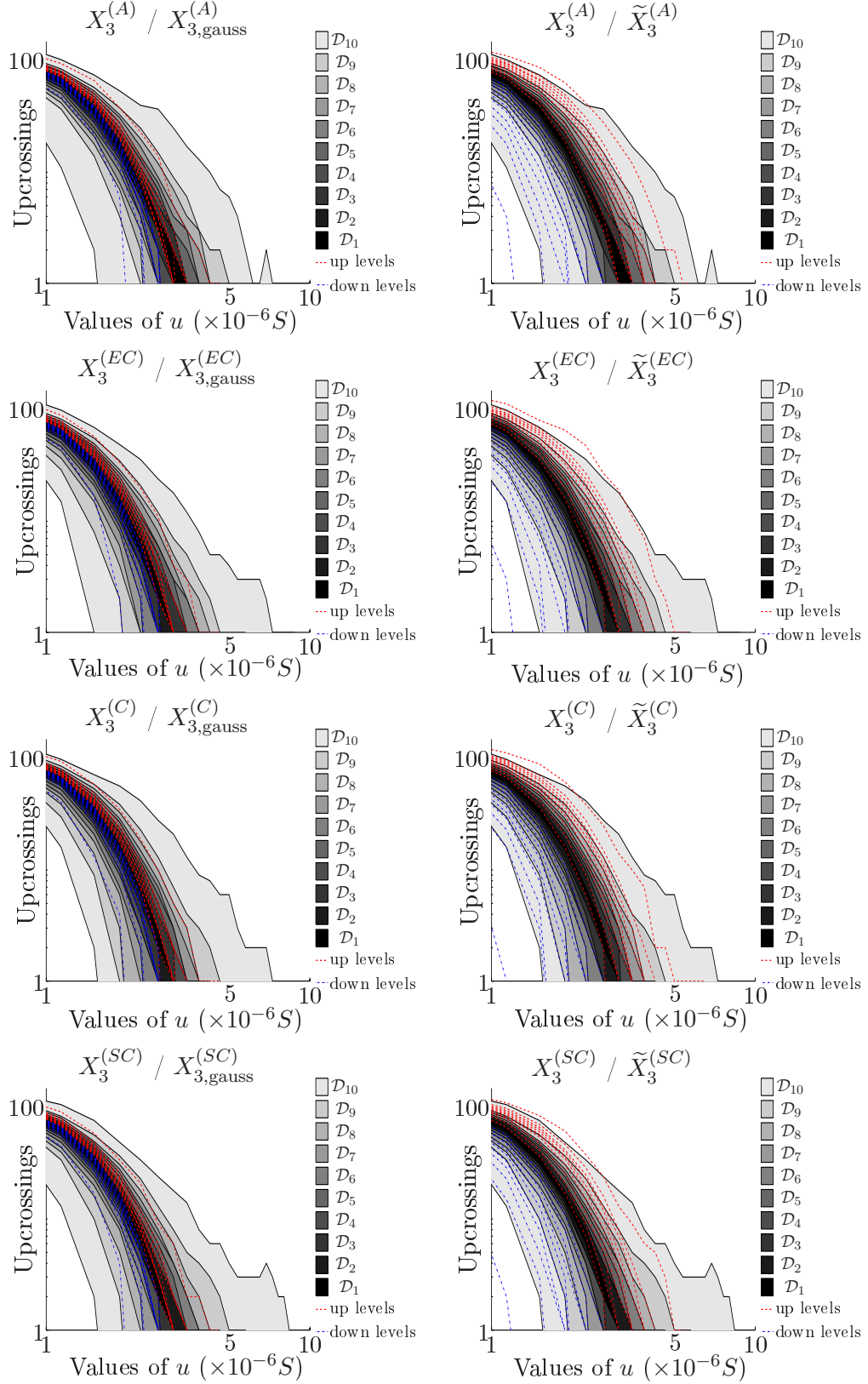


Figure 5.19: Statistical validation of the stochastic modeling of the cant deficiency irregularity X_3

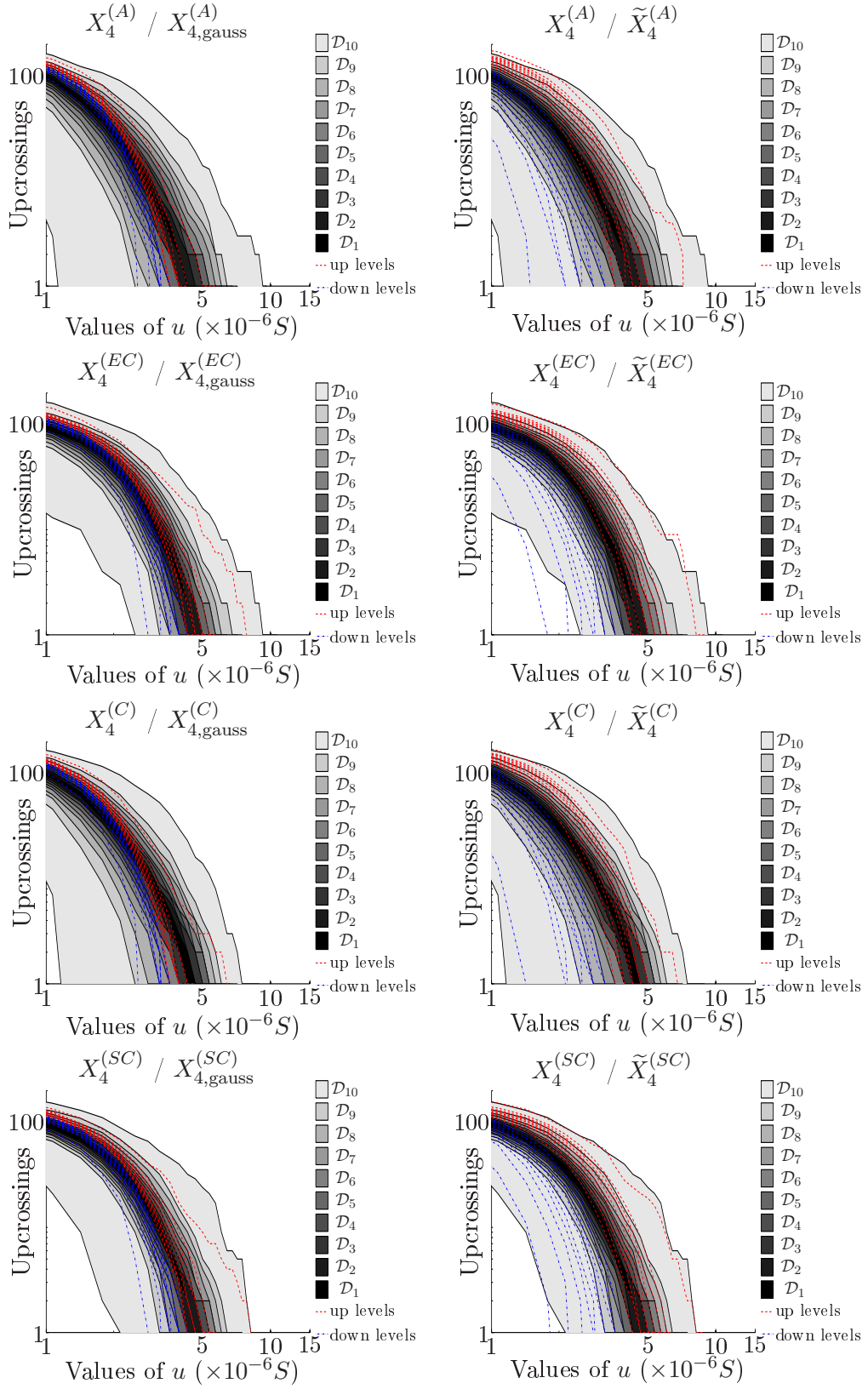


Figure 5.20: Statistical validation of the stochastic modeling of the gauge irregularity X_4

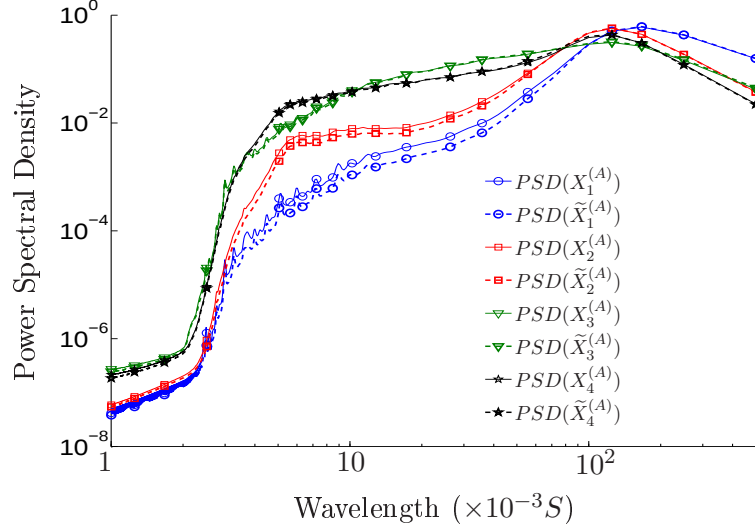


Figure 5.21: Validation of the generation of track geometries with respect to the frequency content.

5.7 Conclusions

A complete parametrization of the track geometry, which takes into account its physical properties and its variability are nowadays of great interest to be able to face always more challenging railway issues. In this prospect, this chapter has presented a general method to model a $\mathbb{R}^4$ -valued random field indexed by $s \in [0, S]$ thanks to a double projection, which can be applied to many other mechanical systems. First, an adapted Karhunen-Loève expansion is used to decompose the random field as a deterministic matrix-valued function and a high dimension random vector. The distribution of this high dimension random vector is then characterized thanks to a truncated PCE. This chapter moreover describes in detail how to control and justify the different truncation parameters. Then, complete track geometries that are realistic and representative of a whole railway network can be generated from a local-global approach. At last, a double validation of this stochastic model is presented, in order to make sure that the frequency and statistical contents of generated and measured track geometries are similar. These geometries can finally be used in any railway software to characterize the dynamic behavior of trains.

Chapter 6

Stochastic dynamics of high-speed trains and risk assessment

6.1 Introduction

For the track quality, the attention must be focused on two main issues. First, the safety of the track-vehicle system has to be guaranteed, and secondly, the maintenance costs have to be controlled and minimized. Safety being the main priority, trains and tracks have been designed with *a priori* high safety factors, such that the limit states of the railway system are not well known but railway accidents almost never happen. In a context of optimization of the maintenance, simulation has thus a big role to play, as it should be able to evaluate these limit situations when experiments cannot or would be too expensive.

In addition, these objectives have to be fulfilled in a context of increasing interoperability. Indeed, European high speed railway networks are meant to go to market. Hence, several high speed trains, such as ICE, TGV, ETR 500, ..., are likely to run on the same tracks, although they have been originally designed for specific and different railway networks. Due to different mechanical properties and structures, the dynamic behaviors, the aggressiveness of the vehicle on the track and the probabilities of exceeding security and comfort thresholds are thus different from one train to another one. From the infrastructure point of view, numerical methods are therefore needed to be able to evaluate and to compare the stability and the safety associated with each train that would apply to run on a particular railway network.

To this end, this chapter shows to what extent the stochastic modeling of the track geometry, which has been presented in Chapter 5, can be coupled with a multibody railway software to analyze the complex link between the track variability and the train dynamics.

6.2 Description of the railway dynamic problem

6.2.1 Deterministic railway problem

A railway simulation can be seen as the dynamic response of the train excited by the track geometry through the wheel/rail contact forces. Three kinds of inputs are thus used in these simulations.

- The vehicle model $\mathcal{V}$. Multibody simulations are usually employed to model the train dynamics. Carbodies, bogies and wheelsets are therefore modeled by rigid bodies linked with connections represented by rheologic models (dampers, springs, ...). For $1 \leq i \leq N_{\text{DoF}}$, and t in $[0, T]$, we denote by $u_i(t)$ the position at time t of the coordinate associated

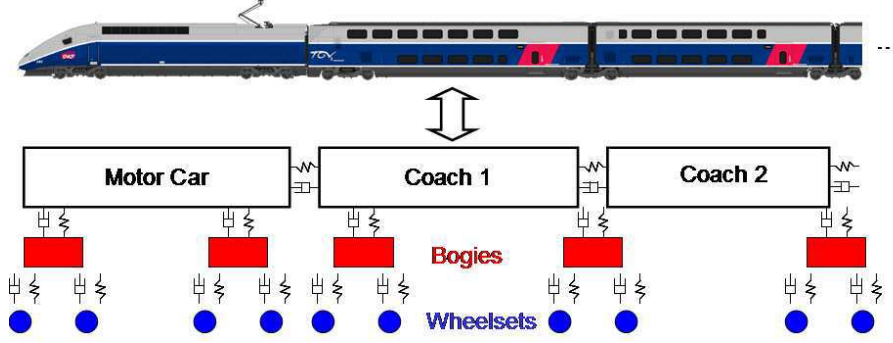


Figure 6.1: Simplified description of a multibody model of a TGV.

with each degree of freedom of the rigid bodies modeling of the train, and by $\dot{u}_i(t) = \frac{du_i}{dt}(t)$ its time derivative. For instance, for a classical one-carriage TGV, which is made of 10 coaches, 13 bogies and 52 wheelsets that are linked by a series of suspensions and bumpstops, N_{DoF} is about two hundreds (see Figure 6.1 for a simplified representation of the TGV).

- The track geometry $\mathcal{T}$. As presented in Chapter 5, this track characterization refers to a double scale description. On the first hand, the track design, which gathers the horizontal curvature c_H , the vertical curvature c_V and the cross level c_L , corresponds to the description of a perfect track without irregularities. On the other hand, four track irregularities, X_1 , X_2 , X_3 and X_4 have to be added to this description to define the real railway tracks. These are due to the train dynamics, the weather conditions and the track substructure evolutions.
- The contact model $\mathcal{C}$ allows the computation of the contact forces between the rails and the wheels. In the railway community, these contact forces are almost always computed from the wheel profile and the rail profile thanks to the Hertz and Kalker theories [7, 6].

Introducing the vector of the generalized coordinates,

$$\mathbf{U}(t) = (u_1(t), u_2(t), \dots, \dot{u}_1(t), \dot{u}_2(t), \dots), \quad (6.1)$$

the train dynamics can therefore be determined by solving the Euler-Lagrange equation, which is written as:

$$\frac{d}{dt} \left(\frac{\partial E_c}{\partial \dot{u}_i} \right) - \frac{\partial E_c}{\partial u_i} = L_i(\mathbf{U}, \mathcal{T}, \mathcal{C}), \quad 1 \leq i \leq N_{\text{DoF}}, \quad (6.2)$$

with E_c the total kinetic energy of the train, and $L_i(\mathbf{U}, \mathcal{T}, \mathcal{C})$ the general load that is applied to the degree of freedom i , which depends on the track geometry $\mathcal{T}$, on the wheel/rail contact $\mathcal{C}$ and on the generalized coordinated $\mathbf{U}$. Eq. (6.2) can be rewritten in a matrix form as:

$$[\mathbf{A}(\mathbf{U})]\dot{\mathbf{U}} = \mathbf{F}(\mathbf{U}, \mathcal{T}, \mathcal{C}), \quad (6.3)$$

with $[\mathbf{A}]$ and $\mathbf{F}$ two strongly **nonlinear** operators. This system is usually solved with an explicit time scheme. In the following, the commercial black-box software, Vampire (see [82, 83] for further details about this software), is used. The chosen time step of this explicit scheme was identified according to a convergence analysis and is generally taken equal to 10^{-4} second.

The generalized coordinates vector $\mathbf{U}$ is then post-treated to define the final comfort and safety criteria associated with the railway system. These outputs can be classified in two categories:

1. First, the maximal values of the vertical and lateral accelerations in the train coaches, $\ddot{z}_{\max}$ and $\ddot{y}_{\max}$, are controlled to guarantee the comfort of the passengers.
2. Secondly, the safety and maintenance criteria of the track-vehicle system are based on the analysis of the wheel/rail contact forces. In this prospect, three classical criteria are generally introduced to characterize the vehicle dynamics on a given track geometry of total length S^{tot} :

- a **shifting** criterion:

$$(Y_\ell + Y_r)_{\max} = \max_{\text{wheelset } w} \left\{ \max_{0 \leq s \leq S^{\text{tot}}} \{Y_\ell^w(s) + Y_r^w(s)\} \right\}, \quad (6.4)$$

- a **derailment** criterion:

$$(Y/Q)_{\max} = \max_{\text{wheel } q} \left\{ \max_{0 \leq s \leq S^{\text{tot}}} \{Y_q(s)/Q_q(s)\} \right\}, \quad (6.5)$$

- a **wear** criterion:

$$(T\gamma) = \sum_{\text{wheel } q} \left\{ \int_0^{S^{\text{tot}}} T_q(s) \gamma_q(s) ds \right\}, \quad (6.6)$$

where:

- Y_ℓ^w and Y_r^w are the left and right lateral forces of the same wheelset w , such that the higher $(Y_\ell + Y_r)_{\max}$ is, the more chance for a shifting of the track there is;
- Y_q and Q_q are the lateral and vertical components of the wheel/rail contact force at wheel q , such that the higher $(Y/Q)_{\max}$ is, the more on the flange a wheel of the train can be;
- T_q and γ_q are respectively the creep force and the slip at wheel q , such that the higher $(T\gamma)$ is, the higher the contact wear is likely to be for one run of the complete train.

Finally, given a model of the wheel/rail contact $\mathcal{C}$, the deterministic railway problem corresponding to the dynamics of a vehicle $\mathcal{V}$ on a track geometry $\mathcal{T}$ can be expressed as:

$$(\mathcal{V}, \mathcal{T}, \mathcal{C}) \mapsto \mathbf{c} = \mathbf{g}(\mathcal{V}, \mathcal{T}, \mathcal{C}), \quad \mathbf{c} = (\ddot{z}_{\max}, \ddot{y}_{\max}, (Y_\ell + Y_r)_{\max}, (Y/Q)_{\max}, (T\gamma)), \quad (6.7)$$

where it is reminded that $\mathbf{g}$ is a complex and nonlinear operator. These nonlinearities are mostly due to the train suspensions (especially the airsprings between the bogies and the coaches), to a series of bumpstops in the train description and to the wheel/rail contact forces.

Due to the train dynamics, to the track irregularities and to the specific wheel and rail profiles, the contact positions between each wheel of the train and the rails keep changing.

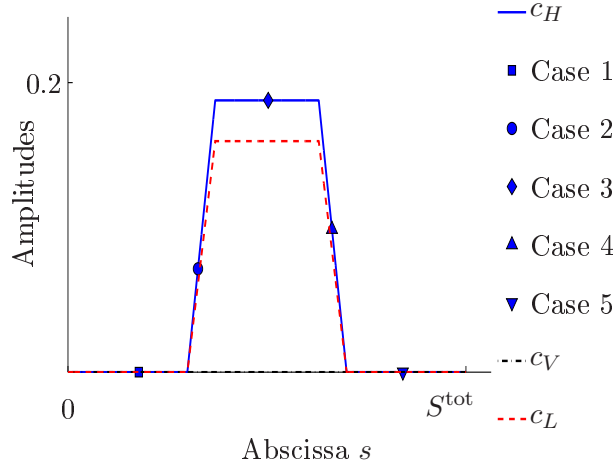


Figure 6.2: Evolution of the horizontal curvature c_H ($\times km^{-1}$), the vertical curvature c_V ($\times km^{-1}$) and of the cross level c_L ($\times m^{-1}$) with respect to the track curvilinear abscissa s .

The worst the track geometry is, the more discontinuous these changes are likely to be. The diversity of these contact positions and contact forces is illustrated in Figures 6.2 and 6.3. These figures are based on the run of a train on a measured track geometry around a curve.

For confidentiality reasons, very few numerical values are given in this work. Hence, only qualitative analysis will be presented in the following.

6.2.2 Domain of validity for the deterministic problem

As a first comment on the validity of the railway models, it is important to point out that all European railway reference standards and reference maintenance guides only consider the low-frequency content, $f \leq f_c$, of the train dynamic quantities of interest (either simulated or measured).

As presented in the former Section, the software Vampire is used to solve the railway deterministic problem. The train being constituted of rigid bodies, the simulated high-frequency response of the train cannot be physical. As an illustration, Figure 6.4 compares the measured and simulated frequency properties of a bogie of a TGV. As shown in Figure 6.5, although the transverse and vertical accelerations of the bogie are low-pass filtered at the reference cut-frequency $f = f_c$, it can be seen that the low-frequency response is well reproduced both in the time and frequency domains by the deterministic model.

As a consequence, in agreement with the work achieved in [84], it is assumed that the proposed railway deterministic model is valid on the frequency band $0 \leq f \leq f_c$. In the following, each output of the train dynamics (whether measured or simulated) will thus be low-pass filtered at frequency f_c before being analyzed.

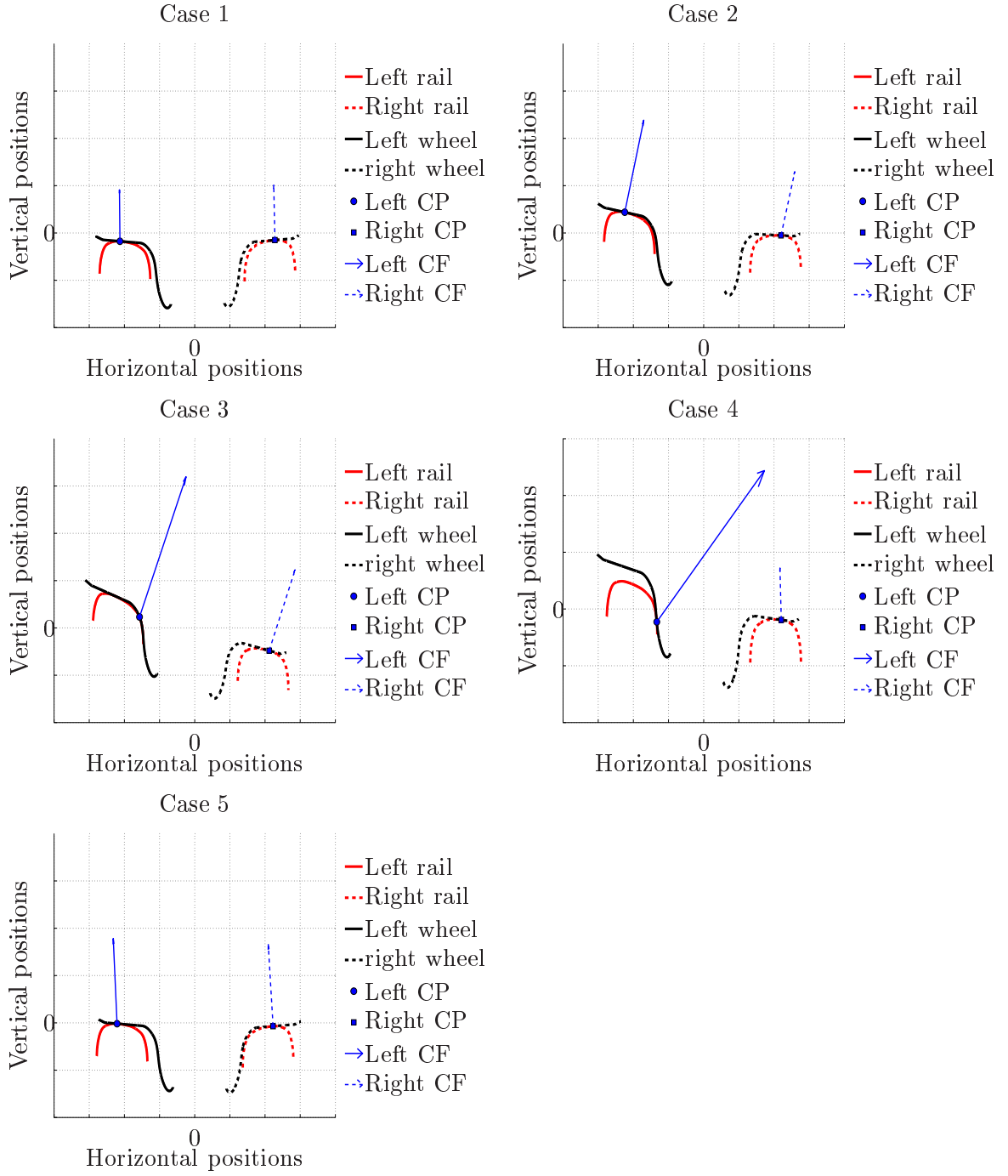


Figure 6.3: Evolution of the contact conditions with respect to the train dynamics and to the horizontal curvature (CP=contact point, CF=contact force).

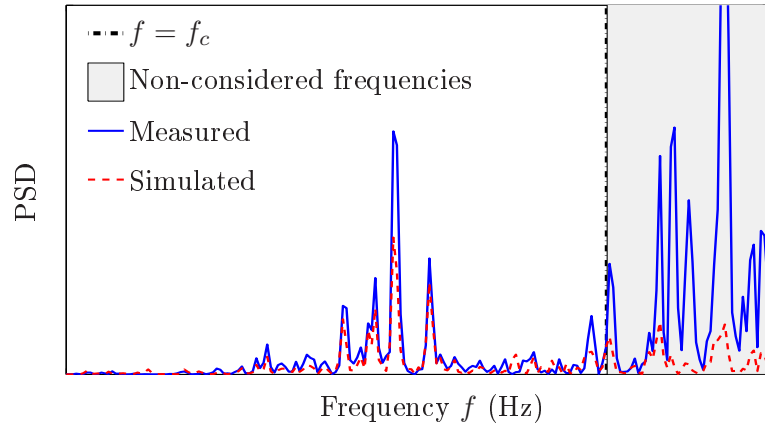


Figure 6.4: Frequency analysis of the transverse acceleration of the bogie of a TGV (figure extracted from [84], pp. 178-179).

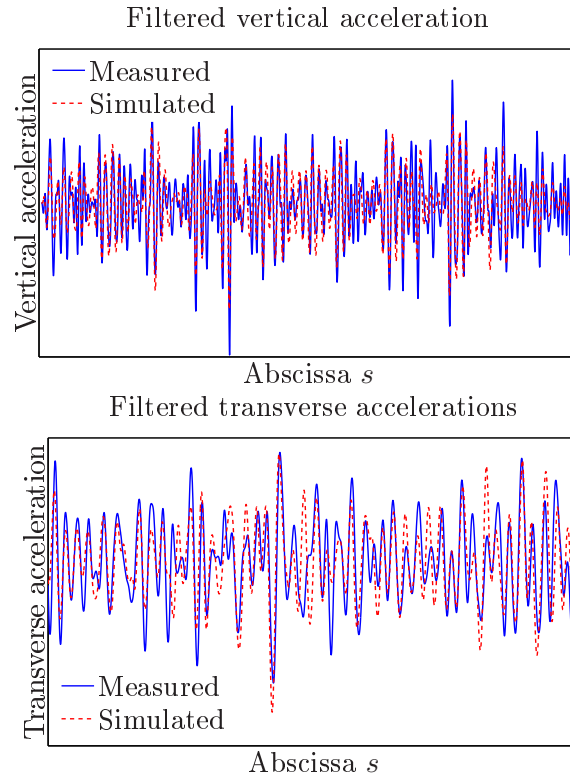


Figure 6.5: Validation of the low-frequency response of the deterministic model (figures extracted from [84], pp. 178-179).

6.3 Definition of the stochastic problem and validation of the modeling

6.3.1 Stochastic problem

The wheel and rail profiles of high speed trains and lines being checked and maintained very regularly, only perfect wheel and rail profiles will be considered in the following, such that the contact properties, $\mathcal{C}$, are chosen to be constant. As presented in Chapter 5, it is moreover supposed that the track irregularities can be separated from the track design. Hence, in the following, the track design is supposed to be constant, while the track irregularities can vary. As a consequence, vector $\mathbf{c}$, which is defined by Eq. (6.7), becomes a random vector that is denoted by $\mathbf{C} = (C_1, C_2, C_3, C_4, C_5)$. It is reminded that by definition of vector $\mathbf{c}$, C_1 and C_2 refer to the vertical and lateral maximal accelerations in the train coaches, C_3 is the maximal value of the sum of the transverse loads of the wheelsets, C_4 is the maximal value of the Y/Q ratio, and C_5 is the cumulated wear along the track.

At last, given a fixed description of the track design, (c_H, c_V, c_L) , and a normalized model of train, $\mathcal{V}$, for which mechanical parameters are also fixed and have been accurately identified, the railway stochastic problem can be written:

$$\mathbf{X}^{\text{tot}} \stackrel{\text{def}}{=} \{ \mathbf{X}^{\text{tot}}(s), s \in [0, S^{\text{tot}}] \} \mapsto \mathbf{C} = \mathbf{G}(\mathbf{X}^{\text{tot}} \mid c_H, c_V, c_L, \mathcal{V}, \mathcal{C}), \quad (6.8)$$

where $\mathbf{X}^{\text{tot}} = (X_1^{\text{tot}}, X_2^{\text{tot}}, X_3^{\text{tot}}, X_4^{\text{tot}})$ is the track irregularity random field computed from the local-global approach described in Chapter 5.

6.3.2 Validation of the stochastic problem

Two validations for the track generator presented in Chapter 5, based on the train dynamics, are proposed in this section. In a first step, it is shown that the track generator coupled with the Vampire software allows us to simulate train accelerations that are similar to accelerations that have been recorded on a real high speed train on a real track. In a second step, we show the relevance of the track stochastic modeling, to generate track conditions that are realistic and representative of the measured track geometries, for the analysis of the wheel/rail forces.

Relevance of the track stochastic modeling for the analysis of the train accelerations

Since 2007, the TGV IRIS-320 has been used to monitor the track geometry of the French high speed lines. This train has been modeled and simulations have been performed at constant speed $\mathbb{S}$ on $\nu = 500$ track geometries of total length S^{tot} . The chosen track design functions, c_H , c_V , c_L , are shown in Figure 6.2. The track irregularities of each track geometry are moreover characterized by independent realizations, $\mathbf{X}^{\text{tot}}(\Theta_n)$, $1 \leq n \leq \nu$, of $\mathbf{X}^{\text{tot}}$. For all s in $[0, S^{\text{tot}}]$, at position s , we respectively define $\hat{C}_z^{\text{sim}}(\Theta_n, s)$ and $\hat{C}_y^{\text{sim}}(\Theta_n, s)$ as the vertical and lateral maximal values of the accelerations in all the coaches of the train that is excited by the track irregularity $\mathbf{X}(\Theta_n)$.

Given these two sets of train responses, $\{\hat{C}_z^{\text{sim}}(\Theta_n), 1 \leq n \leq \nu\}$ and $\{\hat{C}_y^{\text{sim}}(\Theta_n), 1 \leq n \leq \nu\}$, let $\{\mathcal{D}_i^z(s), s \in [0, S^{\text{tot}}], 1 \leq i \leq 10\}$ and $\{\mathcal{D}_i^y(s), s \in [0, S^{\text{tot}}], 1 \leq i \leq 10\}$ be the decile functions, such that at each position s , $i/10$ of the values of $\hat{C}_z^{\text{sim}}(\Theta_n, s)$ and $\hat{C}_y^{\text{sim}}(\Theta_n, s)$ are in $\mathcal{D}_i^z(s)$ and $\mathcal{D}_i^y(s)$ respectively. These decile functions, whose representations are shown in Figure 6.6, allow us to evaluate the influence of the track irregularity variability on such maximal values.

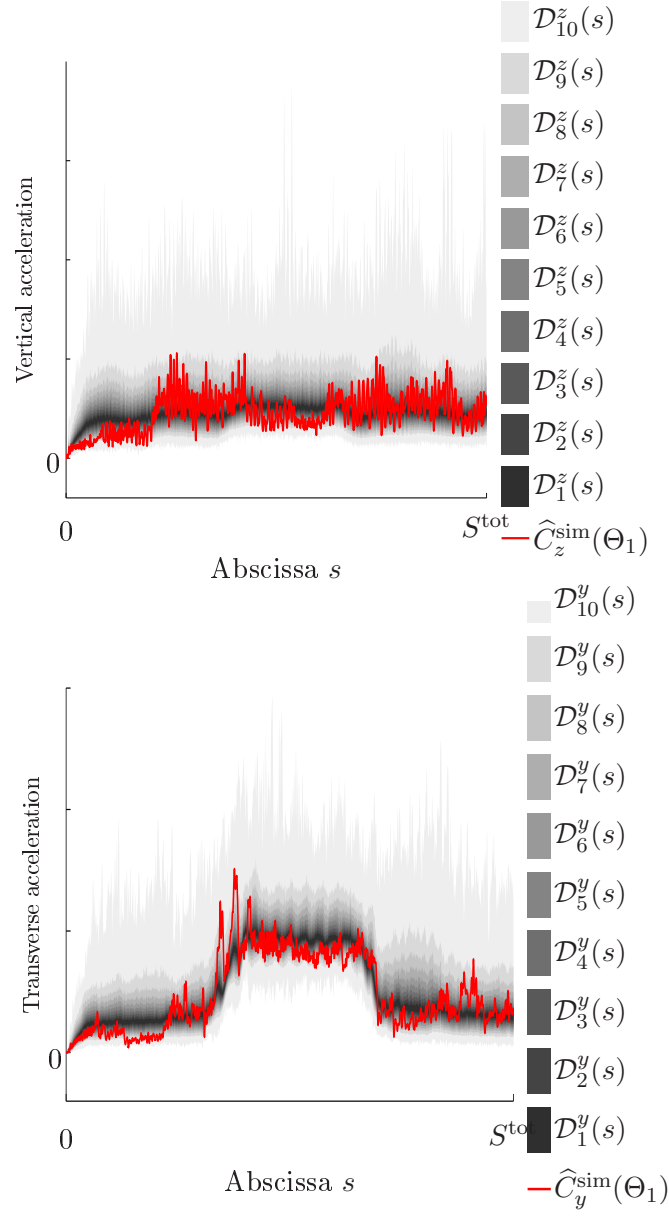


Figure 6.6: Influence of the track variability on the maximal values of the accelerations in the train coaches.

The IRIS-320 train is moreover equipped with accelerometers that record the vertical and transverse accelerations at three coaches,

$$\left\{ \ddot{y}_C^{(1)}, \ddot{y}_C^{(2)}, \ddot{y}_C^{(3)}, \ddot{z}_C^{(1)}, \ddot{z}_C^{(2)}, \ddot{z}_C^{(3)} \right\}.$$

In order to evaluate the relevance of the former results for the maximal accelerations in the train coaches, we define $\widehat{C}_z^{\text{exp}}$ and $\widehat{C}_y^{\text{exp}}$, such that for any value of the curvilinear abscissa of the track, s , we get:

$$\widehat{C}_z^{\text{exp}}(s) = \max_{i \in \{1,2,3\}} \left| \ddot{z}_C^{(i)}(s) \right|, \quad (6.9)$$

$$\widehat{C}_y^{\text{exp}}(s) = \max_{i \in \{1,2,3\}} \left| \ddot{y}_C^{(i)}(s) \right|. \quad (6.10)$$

Five particular evolutions for $\widehat{C}_z^{\text{exp}}$ and $\widehat{C}_y^{\text{exp}}$ over a length S^{tot} are then extracted from the experimental database, which are denoted by $\left\{ \widehat{C}_z^{\text{exp},(1)}, \dots, \widehat{C}_z^{\text{exp},(5)} \right\}$ and $\left\{ \widehat{C}_y^{\text{exp},(1)}, \dots, \widehat{C}_y^{\text{exp},(5)} \right\}$. These measurements were chosen as their dynamic characteristics were the most comparable to the simulated one, in terms of cross level, horizontal and vertical curvatures, speed of the train and length of the curve. If the chosen simulated dynamic characteristics were not similar to the extracted dynamic characteristics on the complete domain $[0, S^{\text{tot}}]$, non-valid domains were added to these figures. The evolutions of these measured accelerations are compared to the simulated ones in Figure 6.7.

In the light of these results, the track generator coupled with the Vampire software seems to be able to simulate realistic and representative runs of the IRIS-320 train to analyze the link between the two first quantities of interest of the stochastic modeling, C_1 and C_2 , and the track geometry variability.

Dynamic validation of the track generator for the analysis of the wheel/rail contact forces

No on-track measurements of the contact forces between the train and the track at high speed being available, an other approach is proposed to evaluate the relevance of the track generator to simulate realistic and representative values for C_3 , C_4 and C_5 .

To this end, the particular curve of total length S^{tot} shown in Figure 6.2, is once again considered. From the available measurements of the track geometry, $\nu^{\text{exp}} = 400$ different track conditions of total length S^{tot} , $\{\mathbf{X}_{\text{exp}}(\theta_1), \dots, \mathbf{X}_{\text{exp}}(\theta_{\nu^{\text{exp}}})\}$, are gathered. These track conditions stem from the random concatenation of measured track sections that are in alignment, in transition curve entrance or exit, or in curve, in order to suit the chosen track design.

The same normalized high-speed train $\mathcal{V}$, for which mechanical parameters are supposed to be accurately identified, is thus made run first on the ν^{exp} measured track conditions, and then on ν generated track conditions, $\{\mathbf{X}^{\text{tot}}(\Theta_1), \dots, \mathbf{X}^{\text{tot}}(\Theta_\nu)\}$, at the same speed S . Eight quantities of interest that are representative of the train dynamics are then compared:

- the left and right transverse contact forces at the first wheelset of the first bogie of the motor car, $Q_1 = Y_{MC}^\ell$ and $Q_2 = Y_{MC}^r$;
- the left and right transverse contact forces at the second wheelset of the second bogie of the second passenger car, $Q_3 = Y_{PC}^\ell$ and $Q_4 = Y_{PC}^r$;
- the left and right Y/Q ratio at the first wheelset of the first bogie of the motor car, $Q_5 = (Y/Q)_{MC}^\ell$ and $Q_6 = (Y/Q)_{MC}^r$;

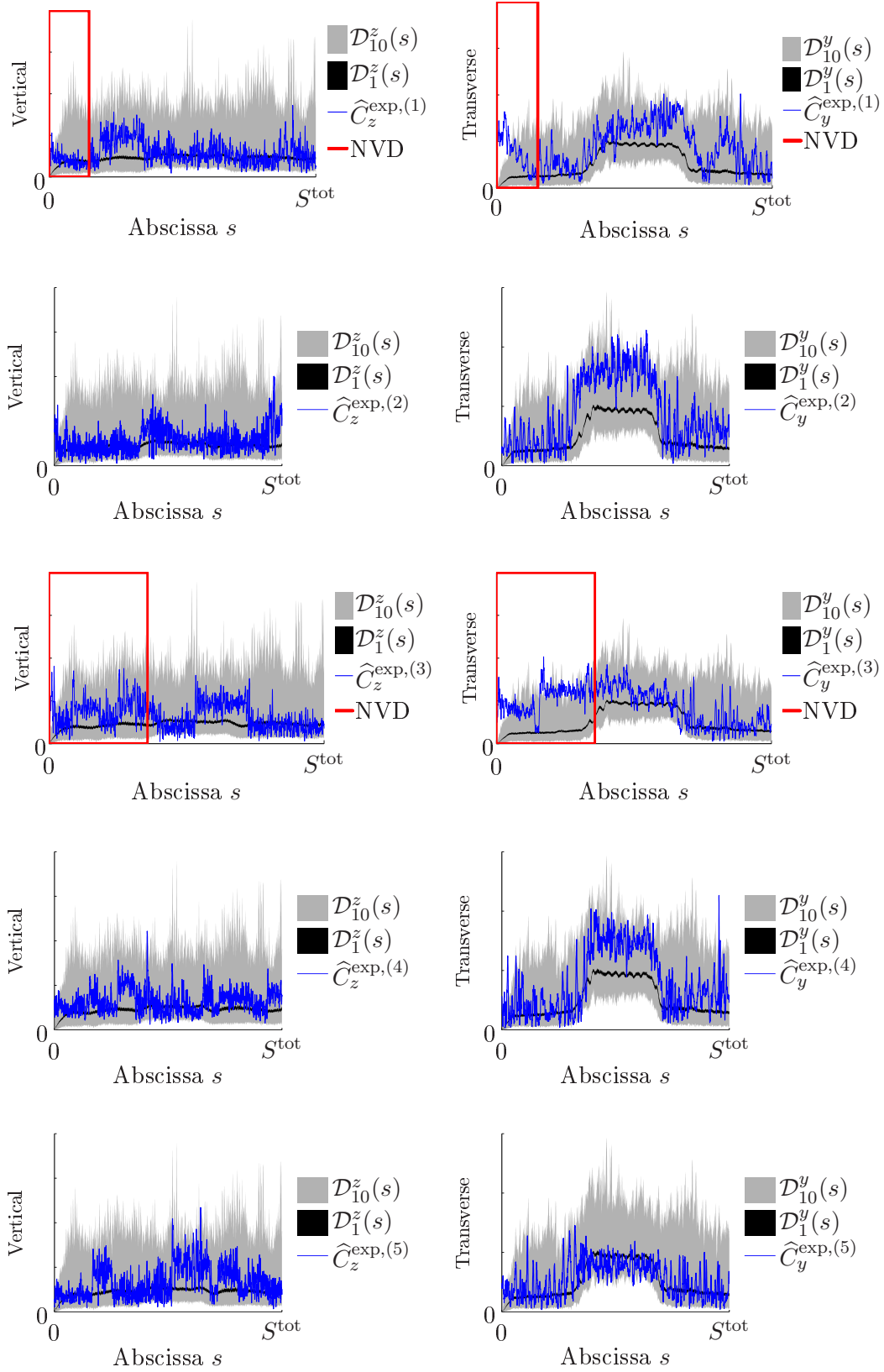


Figure 6.7: Comparison between simulated and measured maximal values of the vertical and transverse accelerations in the train coaches (NVD = non-valid domain).

- the left and right Y/Q ratio at the second wheelset of the second bogie of the second passenger car, $Q_7 = (Y/Q)_{PC}^\ell$ and $Q_8 = (Y/Q)_{PC}^r$.

In the same manner than in Section 5.6, for $1 \leq i \leq 8$, we are interested in the mean power spectral densities of Q_i and the mean numbers of upcrossings of the level u by Q_i over the length S^{tot} , which are respectively denoted by $PSD^{\text{mes}}(Q_i)$ and $N_{\text{up}}^{\text{mes}}(Q_i, u, S^{\text{tot}})$ when these quantities are computed from the measured track geometries and $PSD^{\text{gen}}(Q_i)$ and $N_{\text{up}}^{\text{gen}}(Q_i, u, S^{\text{tot}})$ when these quantities are computed from the generated track geometries. The comparisons between these quantities are represented in Figure 6.8. It can be seen that the fit is very good.

The stochastic modeling of the track geometry is thus relevant from the train response point of view. In the following, it is therefore supposed that the track stochastic modeling, coupled with the software Vampire is also relevant to investigate the relation between the track variability and the three quantities of interest C_3 , C_4 and C_5 .

6.4 Propagation of the variability

As explained in Section 6.1, a better understanding of the specific link between the track irregularities and the train response is needed to optimize the maintenance, and to better anticipate the consequences of modifications of the running conditions.

In this prospect, we denote by $P_C(d\mathbf{x}) = p_C(\mathbf{x})d\mathbf{x}$ the multidimensional distribution of random vector $\mathbf{C}$, where p_C is the associated density. This distribution is strongly related to the distribution of the track irregularity random field, $P_{\mathbf{X}^{\text{tot}}}$ (see Eq. (6.8)). Assuming that the latter distribution has been accurately identified from the local-global approach described in Chapter 5, the track variability has now to be propagated through the railway model to characterize P_C .

As the statistical dimension of $\mathbf{X}^{\text{tot}}$ is very high and as the relation between P_C and $P_{\mathbf{X}^{\text{tot}}}$ is very complex and strongly nonlinear, the Monte Carlo method appears to be the best approach to do so. Indeed the convergence properties associated with this method are independent of the statistical dimension of the input.

From ν independent realizations of $\mathbf{X}^{\text{tot}}$, $\{\mathbf{X}^{\text{tot}}(\Theta_1), \dots, \mathbf{X}^{\text{tot}}(\Theta_\nu)\}$, ν independent realizations of $\mathbf{C}$, $\{\mathbf{C}(\Theta_1), \dots, \mathbf{C}(\Theta_\nu)\}$, can be deduced as:

$$\mathbf{C}(\Theta_n) = \mathbf{G}(\mathbf{X}^{\text{tot}}(\Theta_n) \mid c_H, c_V, c_L, \mathcal{V}, \mathcal{C}), \quad 1 \leq n \leq \nu. \quad (6.11)$$

The statistical properties of $\mathbf{C}$ are finally deduced from the analysis of this ν -dimension set of independent realizations of $\mathbf{C}$.

Three applications of this stochastic modeling are now presented. These are based on the track design of total length S^{tot} shown in Figure 6.2, and on $\nu = 4,000$ track irregularity realizations. First, the influence of the track design and the track irregularities is illustrated. Then, it is shown to what extent such a method can be used to quantify the influence of an increase of the train speed on $\mathbf{C}$. At last, the method is used to compare the safety and the aggressiveness of three different high speed trains.

6.4.1 Influence of the track design

The idea of this section is to quantify the importance of the track irregularities and of the track design on vector $\mathbf{C}$. In this prospect, the response of a normalized high train $\mathcal{V}_1$ to the former

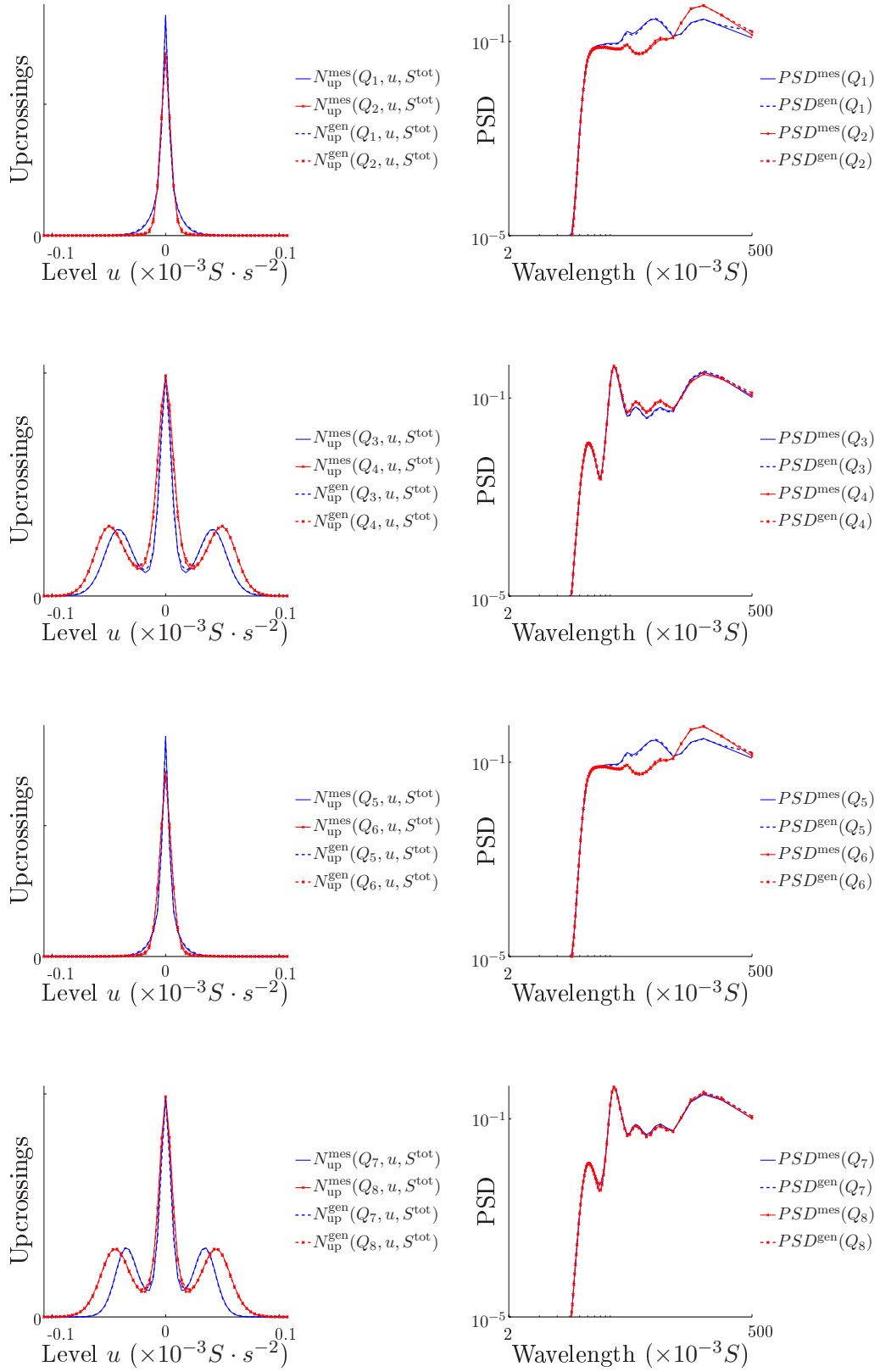


Figure 6.8: Spectral and statistical analysis of the dynamic quantities of interest $Q_1, Q_2, \dots, Q_8$.

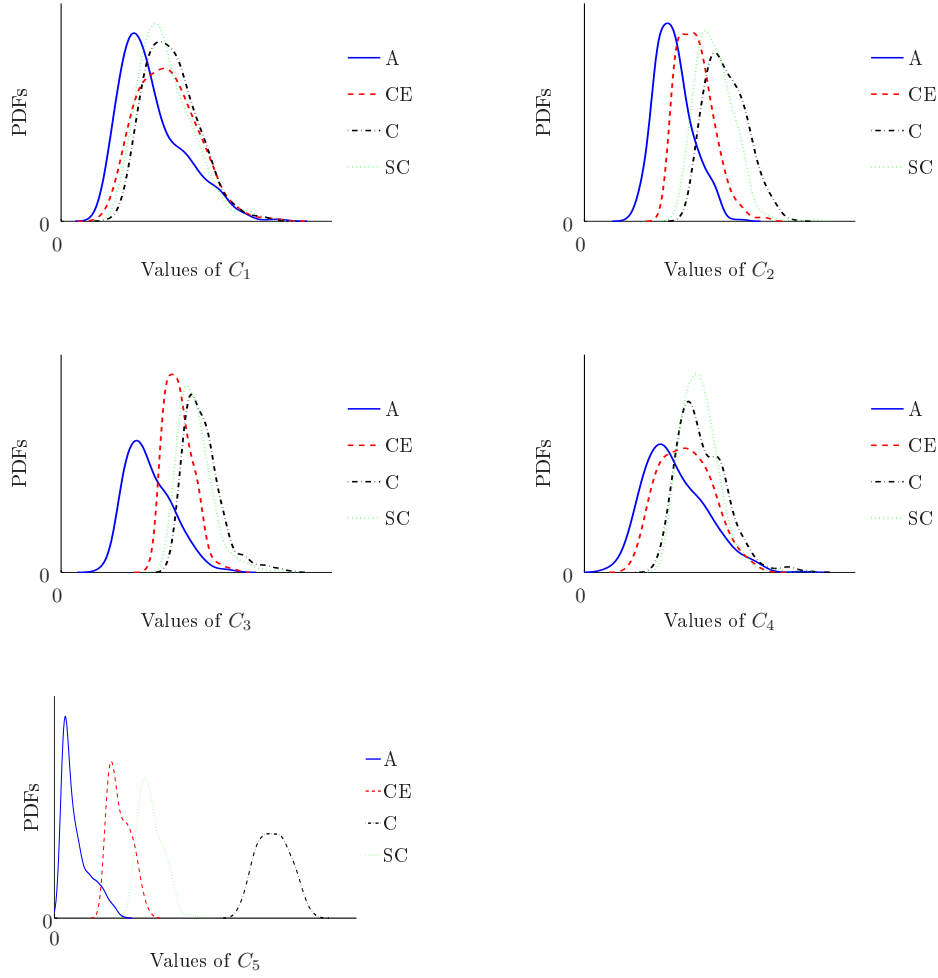


Figure 6.9: Influence of the track design on the marginal PDFs of vector $\mathbf{C}$.

ν track conditions of total length S^{tot} is analyzed. In the same manner as in Section 5.4.1, four categories are considered: the alignment (A), the curve entrance (CE), the established curve (C) and the curve exit (SC). The response of the train is therefore sorted with respect to these four curve categories, such that, for $1 \leq i \leq 5$, four values of the railway quantities of interest $C_i^A(\Theta_n)$, $C_i^{EC}(\Theta_n)$, $C_i^C(\Theta_n)$ and $C_i^{SC}(\Theta_n)$ can be computed. Based on these sets of ν independent realizations, the PDFs of the components of $\mathbf{C}$ are estimated from a kernel smoothing method, and are represented in Figure 6.9. From these graphs, it can be seen that the influence of the track design on the wear criterion, C_5 is very high. The other dynamic quantities, C_1 , C_2 , C_3 and C_4 seem however to be much more dependent on the track irregularities than on the track design.

6.4.2 Influence of an increase of the speed on the quantities of interest

The second application of the whole method deals with the influence of the speed on the PDFs of the five considered criteria. Only the established curve configuration case is shown.

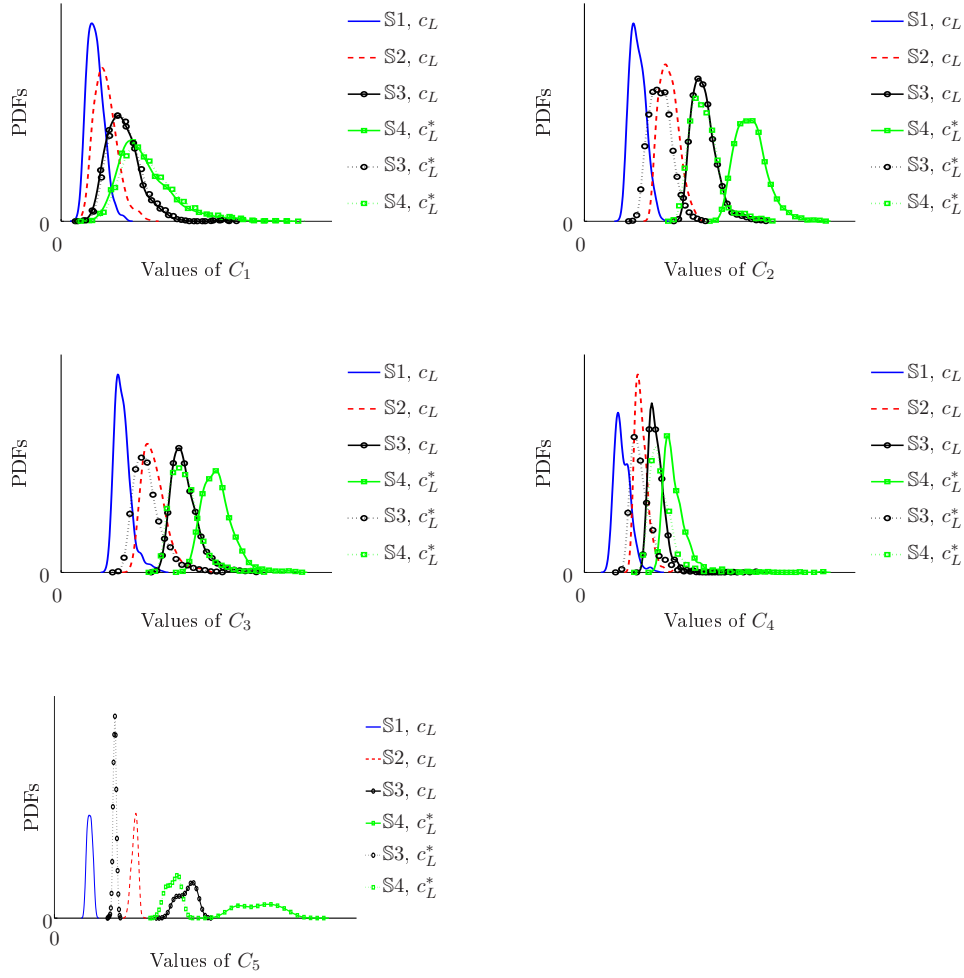


Figure 6.10: Influence of the train speed on the marginal PDFs of vector $\mathbf{C}$.

Railway simulations are therefore performed on the same ν realistic and representative track geometries, at the four speeds $\mathbb{S}1 = \mathbb{S}$, $\mathbb{S}2 = 1.1\mathbb{S}$, $\mathbb{S}3 = 1.2\mathbb{S}$ and $\mathbb{S}4 = 1.3\mathbb{S}$. Two other sets of simulations have then been carried out for a different value of the track superelevation, c_L^* , at speeds $\mathbb{S}3 = 1.2\mathbb{S}$ and $\mathbb{S}4 = 1.3\mathbb{S}$ in order to quantify the importance of this track design parameter with respect to the three criteria studied. In other words, whereas c_L is chosen to compensate the train inertial acceleration in curve at speed $\mathbb{S}1$, c_L^* allows the compensation of the train inertial acceleration in curve at speed $\mathbb{S}3$.

For each speed, the PDFs of each component of $\mathbf{C}$ are once again estimated using a kernel smoothing method based on the $\nu = 4,000$ independent railway simulations. These PDFs are represented in Figure 6.10. In this figure, the nonlinearity of the system can be noticed, as the consequences of an increase of the speed of 10% to 30% are much higher than 30% for each criterion. In particular, an increase of 30% of the speed of the train can yield an increase of more than 500% of the contact wear if the track superelevation is not adjusted. In addition, these figures emphasize the importance of the adjustment of the track superelevation to the speed, in terms of minimization of wear, of shifting and of risk of derailment.

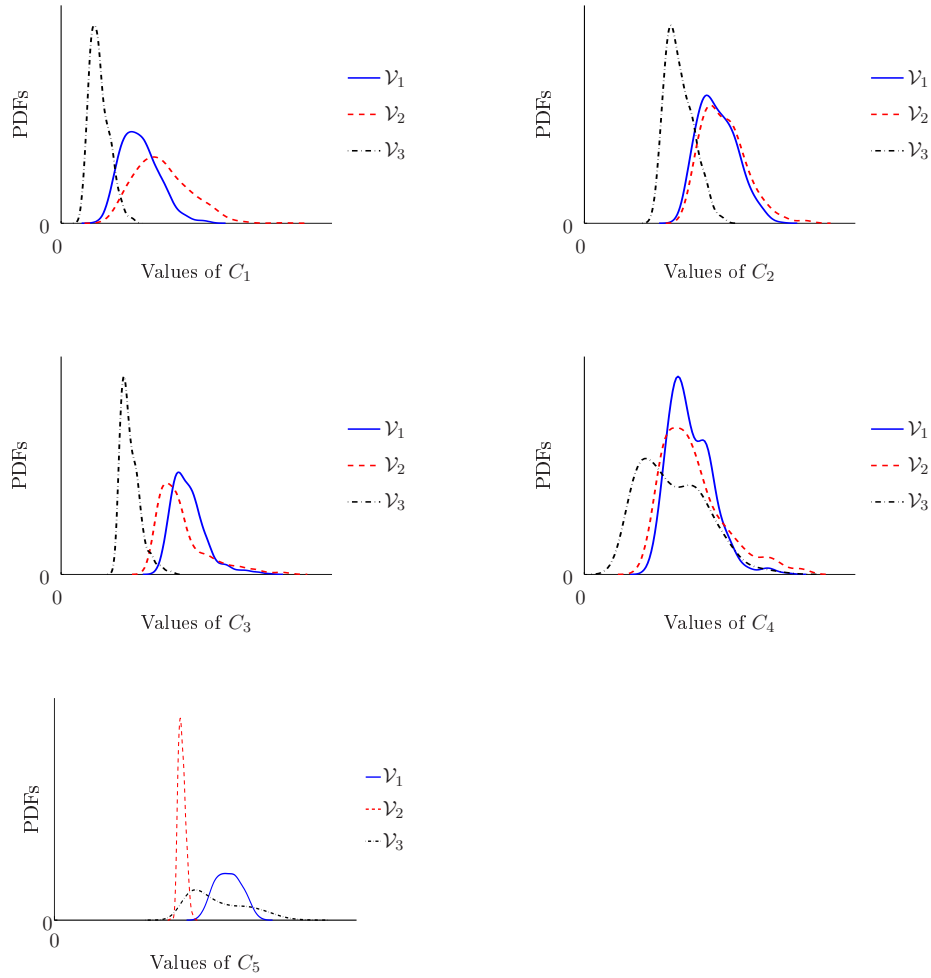


Figure 6.11: Influence of the train characteristics on the PDFs of $\mathbf{C}$.

6.4.3 Comparison of three high speed trains

In this section, it is supposed that three different models of three competitive high speed trains, $\mathcal{V}_1$, $\mathcal{V}_2$ and $\mathcal{V}_2$ are available. The mechanical parameters of these trains are very different and were carefully identified from experimental measurements. These three trains are thus made run on the same ν track geometries at the same speed S . The PDFs of each criterion C_i associated with each train are then shown in Figure 6.11. Hence, the stochastic modeling allows us to compare the dynamical response of these three trains when excited by a representative set of realistic track conditions. In particular, criteria C_3 and C_5 could be interesting indicators to compare the aggressiveness of each train.

6.5 Sensitivity analysis

For years, railway engineers have been working on the identification of the most dangerous and uncomfortable track irregularities for the train dynamics. These research have been mostly based on correlation analysis between the outputs of the train and the amplitudes, the wave-

lengths or the maximal values of the track irregularities.

Based on the ν former simulation of the normalized vehicle $\mathcal{V}_1$ at constant speed $\mathbb{S}$, the idea of this section is to perform an analysis of sensibility of the train response with respect to the track irregularities. At first, it will be shown that a direct analysis of the correlations between inputs and outputs has little chance of success, due to the high nonlinearities of the train suspensions and bumpstops and of the wheel/rail contact. Then, an original sensitivity method based on the scaled expansion developed in Chapter 4 will be presented.

6.5.1 Nonlinearities and importance of the conjunction of track irregularities

In this section, only four dynamic outputs, $\widehat{C}_1$, $\widehat{C}_2$, $\widehat{C}_3$ and $\widehat{C}_4$, are considered:

- $\widehat{C}_1$ the vertical acceleration at the center of gravity of the 5th coach of $\mathcal{T}$,
- $\widehat{C}_2$ the transverse acceleration at the center of gravity of the 5th coach of $\mathcal{T}$,
- $\widehat{C}_3$ the sum of the transverse loads of the first wheelset of the first bogie of the 5th coach of $\mathcal{T}$,
- $\widehat{C}_4$ the Y/Q ratio of the left wheel of the first wheelset of the first bogie of the 5th coach of $\mathcal{T}$.

Hence, we are interested in the identification of the **local** shapes of the track irregularities that bring about the highest values for these four quantities $\widehat{C}_1$, $\widehat{C}_2$, $\widehat{C}_3$ and $\widehat{C}_4$. To this end, for $1 \leq i \leq 4$ and $0 \leq S^{\text{por}} \leq S^{\text{tot}}$, we denote by

$$\mathcal{S}^i(S^{\text{por}}, T_i) = \left\{ \left(\mathbf{X}^{\text{por}, i, q}, \widehat{C}_i^q \right), 1 \leq q \leq Q_i \right\},$$

the sets gathering the Q_i track irregularities of length S^{por} , that are centered at the values of $\widehat{C}_i$ that are higher than the threshold T_i . Threshold T_i is chosen sufficiently high, such that at most one couple $\left(\mathbf{X}^{\text{por}, i, q}, \widehat{C}_i^q \right)$ can be extracted from each railway simulation. Hence, the elements of $\mathcal{S}^i(S^{\text{por}}, T_i)$ can therefore be considered as statistically independent.

For confidentiality reasons, the values of length S^{por} and threshold T_i , which are introduced to carry out a local analysis of the track irregularities, are not given in this work.

In addition, let $X_j^{\text{max}, i, q}$ and $\widehat{C}_i^{\text{max}, q}$ be the maximal values such that:

$$X_j^{\text{max}, i, q} = \max_{s \in [0, S^{\text{por}}]} \left\{ X_j^{\text{por}, i, q}(s) \right\}, \quad 1 \leq j \leq 4, \quad (6.12)$$

$$\widehat{C}_i^{\text{max}, q} = \max_{s \in [0, S^{\text{por}}]} \left\{ \widehat{C}_i^q \right\}. \quad (6.13)$$

For $1 \leq i \leq 4$, the evolutions of $X_j^{\text{max}, i, q}$ with respect to $\widehat{C}_i^{\text{max}, q}$ are then represented in Figure 6.12. From these scatter plots, it can therefore be noticed that no linear nor monotonous relation between $X_j^{\text{max}, i, q}$ and $\widehat{C}_i^{\text{max}, q}$ can be identified. In the same manner, from this direct approach, it is hard to tell if the increase of $\widehat{C}_i^{\text{max}}$ is mostly due to one track irregularity or another one.

In other words, for $1 \leq i \leq 4$, from the ν available simulations, it can easily be extracted track conditions with high track irregularities that would less excite the train than track conditions

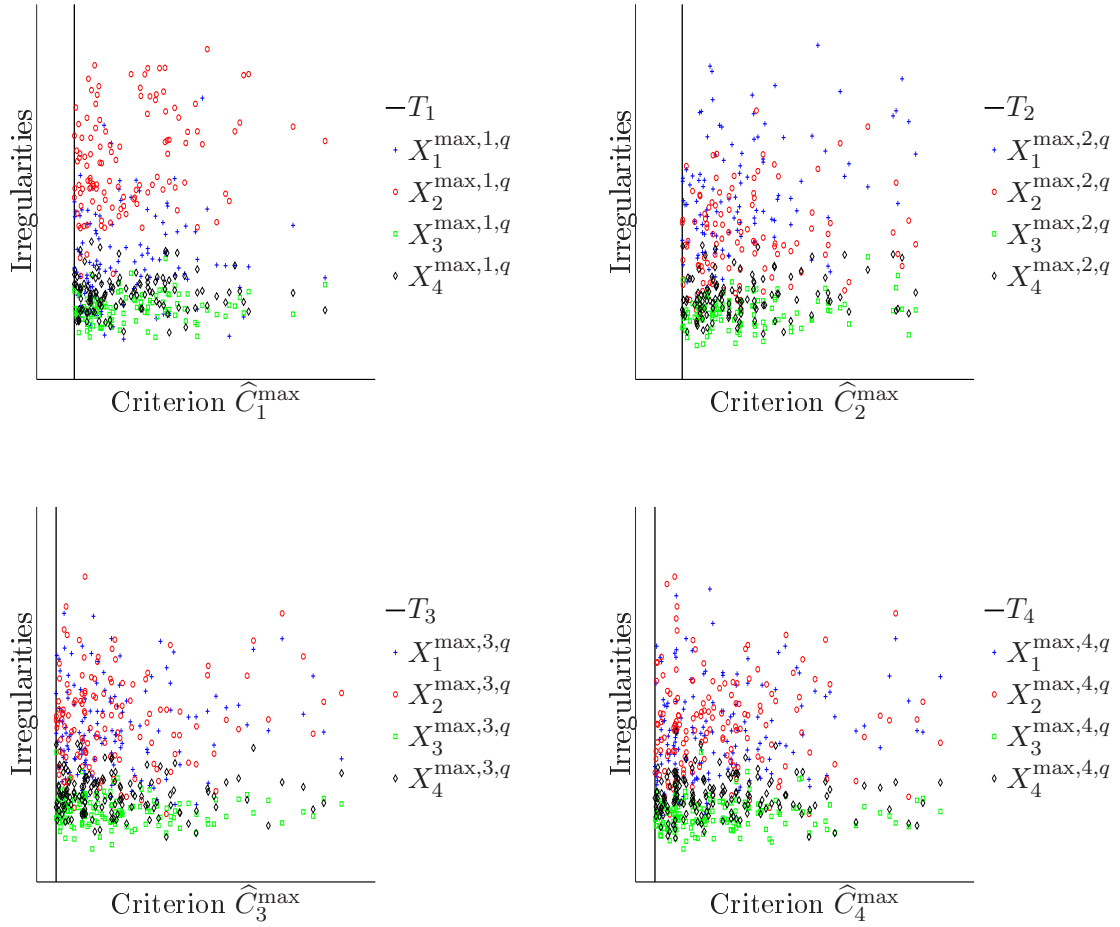


Figure 6.12: Analysis of the correlation between $\hat{C}^1$ and the four track irregularities.

with low track irregularities, as it is shown in Figure 6.13 (the values of the track irregularities on the first hand and of the dynamic quantities on the other hand were scaled to be shown in the same graphs). Even if the maximal amplitudes or the variances of the track irregularities seem to be representative quantities for the track quality, it can be seen from these results that they are not at all sufficient. New methods that would be able to better take into account the combination of the four track irregularities, as well as their specific shapes, are thus needed.

6.5.2 KL-based sensitivity analysis

The fact that no direct relation can be emphasized between the maximal values of $\hat{C}_1$, $\hat{C}_2$, $\hat{C}_3$, $\hat{C}_4$, and the maximal values of the track irregularities on a restricted length S^{por} , motivates the introduction of an alternative method to identify the most uncomfortable irregularity shapes.

To this end, the method we propose is based on the scaled expansion developed in Chapter 4. Using the same notations than in Chapter 2 and 4, for $1 \leq i \leq 4$, we define $\mathbf{Z}^{T_i}$ in $\mathcal{P}(\Omega^{\text{por}})$ as the second-order random field, indexed by s in $\Omega^{\text{por}} = [0, S^{\text{por}}]$ with values in $\mathbb{R}^5$, which is conditioned by $\mathbf{X}^{\text{tot}}$ in the sense that, for each independent realization $\{\mathbf{X}^{\text{tot}}(\Theta, s), s \in [0, S^{\text{tot}}]\}$ of $\mathbf{X}^{\text{tot}}$, if it exists s^* in $[S^{\text{por}}/2, S^{\text{tot}} - S^{\text{por}}/2]$ such that $|\hat{C}_i(s^*)| \geq T_i$, we get an independent

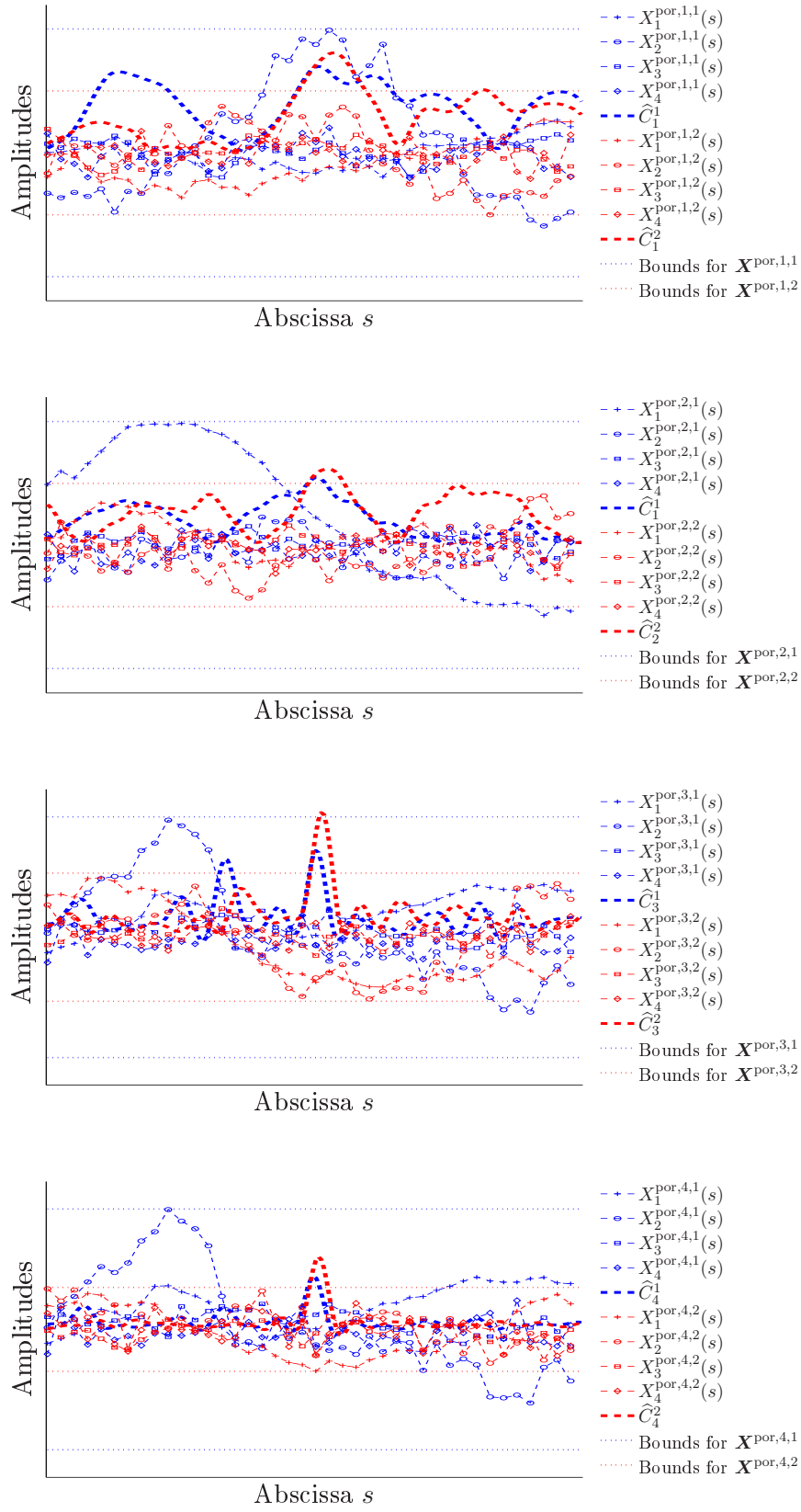


Figure 6.13: Influence of the maximal value of the track irregularities on the dynamic quantities $\hat{C}_1$, $\hat{C}_2$, $\hat{C}_3$, $\hat{C}_4$.

realization, $\mathbf{Z}^{T_i}(\Theta)$, of $\mathbf{Z}^{T_i}$:

$$s \in [0, S^{\text{por}}], \quad \begin{cases} Z_1^{T_i}(\Theta, s) = \widehat{C}_i(s + s^* - S^{\text{por}}/2), \\ Z_{j+1}^{T_i}(\Theta, s) = X_j^{\text{tot}}(\Theta, s + s^* - S^{\text{por}}/2), \quad 1 \leq j \leq 4. \end{cases} \quad (6.14)$$

Given this formalism, for $1 \leq i \leq 4$, it can be noticed that Q_i independent realizations, $\{\mathbf{Z}^{T_i}(\theta_q), 1 \leq q \leq Q_i\}$, of $\mathbf{Z}^{T_i}$ can be computed from the ν railway simulations computed in Section 6.5.1.

For any $\mathbf{O}$ in $]0, +\infty[^5$, let $\mathbf{Y}^{T_i}(\mathbf{O})$ be the scaled random field associated with $\mathbf{Z}^{T_i}$, such that:

$$\mathbf{Y}_j^{T_i}(\mathbf{O}) = O_j \mathbf{Z}_j^{T_i}, \quad 1 \leq j \leq 5. \quad (6.15)$$

For any fixed value of $\mathbf{O}$ in $]0, +\infty[^5$, $\mathbf{Y}_j^{T_i}(\mathbf{O})$ is also in $\mathcal{P}(\Omega^{\text{por}})$ and we can introduce $\boldsymbol{\mu}_{\mathbf{Y}^{T_i}(\mathbf{O})}$ and $[R_{\mathbf{Y}^{T_i}(\mathbf{O})}]$ as its mean value and its matrix-valued covariance function. Using the same notations than in Chapters 2 and 4, we moreover denote by $\{k^m(\widehat{C}_i), 1 \leq m\}$ and $\{\mathbf{k}^m(\mathbf{O}), 1 \leq m\}$ the KL projection basis associated with $\widehat{C}_i$ and $\mathbf{Y}^{T_i}(\mathbf{O})$ respectively, such that:

$$\begin{cases} \mathbf{Y}^{T_i}(\mathbf{O}) = \boldsymbol{\mu}_{\mathbf{Y}^{T_i}(\mathbf{O})} + \sum_{m \geq 1} \mathbf{k}^m(\mathbf{O}) (\mathbf{Y}^{T_i}(\mathbf{O}), \mathbf{k}^m(\mathbf{O})), \\ \widehat{C}_i = E[\widehat{C}_i] + \sum_{m \geq 1} k^m(\widehat{C}_i) (\widehat{C}_i, k^m(\widehat{C}_i)). \end{cases} \quad (6.16)$$

For any κ in $]0, +\infty[$, $\mathbf{O}$ is chosen such that:

$$\begin{cases} O_1 = \kappa, \\ O_j = 1 / \|\mathbf{Z}_j^{T_i}\|_{\mathcal{P}(\Omega^{\text{por}})}. \end{cases} \quad (6.17)$$

It is then assumed that, for $1 \leq m$, the functions

$$\kappa \mapsto (k_1^m(\mathbf{O}))^2 / \|k_1^m(\mathbf{O})\|_{L^2}^2, \quad (6.18)$$

$$\kappa \mapsto \left((k_2^m(\mathbf{O}))^2, (k_3^m(\mathbf{O}))^2, (k_4^m(\mathbf{O}))^2, (k_5^m(\mathbf{O}))^2 \right) / \|(k_2^m(\mathbf{O}), k_3^m(\mathbf{O}), k_4^m(\mathbf{O}), k_5^m(\mathbf{O}))\|_{L^2}^2, \quad (6.19)$$

converge to the limit functions $k^m(\widehat{C}_i)$ and $\mathbf{L}^m(T_i) = (L_1^m(T_i), L_2^m(T_i), L_3^m(T_i), L_4^m(T_i))$ respectively when κ tends to infinity. In other words, by making κ tend to infinity, we admit that it is possible to extract the KL expansion of $\widehat{C}_i$. Even if no numerical example has been found to contradict them, these convergence properties have not been proven yet in the general case.

Therefore, if $Z_1^{T_i}$ and $Z_{j+1}^{T_i}$ are uncorrelated, the components of $\mathbf{L}^m(T_i)$ are equal to zero. On the contrary, if $Z_1^{T_i}$ and $Z_{j+1}^{T_i}$ are correlated, these limit functions are not equal to zero, and it is assumed that the first elements of these limit functions allow us to identify the shapes of the track irregularities that are the most correlated to the component $k^m(\widehat{C}_i)$ of the KL expansion of $\widehat{C}_i$.

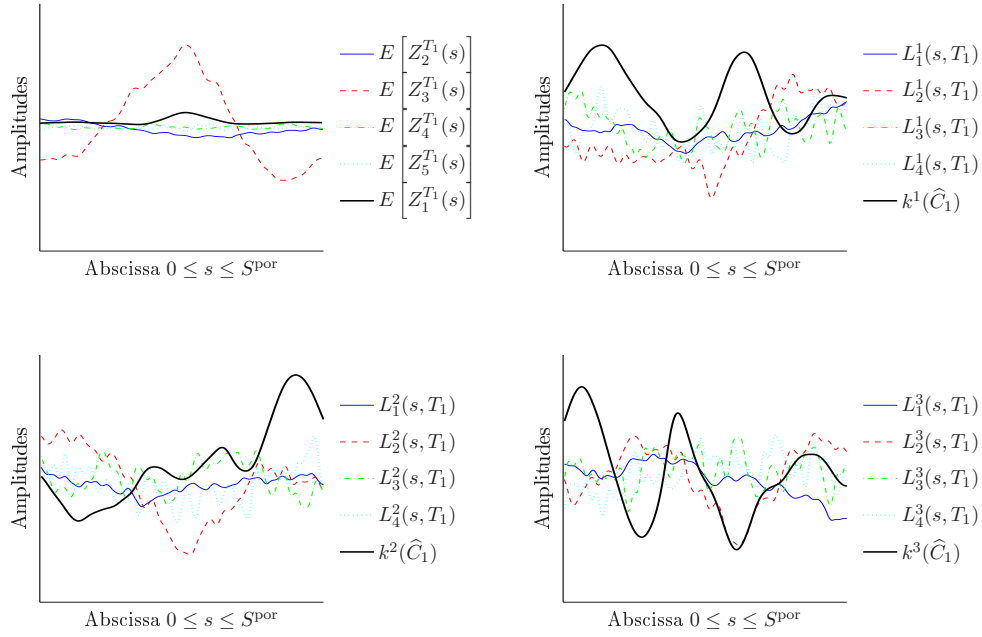


Figure 6.14: Correlation analysis between the shapes of the track irregularity and $\hat{C}_1$.

Based on the ν former simulations, for $1 \leq i \leq 4$, the mean values and the three first eigenfunctions associated with $\mathbf{Z}^{T_i}$ when κ tends to infinity are shown in Figures 6.14, 6.15, 6.16 and 6.17. From these graphs, as expected, it can be seen that the extreme values of the vertical and transverse accelerations of the train coaches are mostly correlated to the vertical and horizontal alignment irregularities respectively. More interesting, these figures show that the extreme values of the transverse wheel/rail forces and the Y/Q ratio are not due to a high value of one track irregularity but seem to be correlated to a combination of the four track irregularities. Indeed, from the mean value and the first eigenfunctions associated with $\mathbf{Z}^{T_3}$ and $\mathbf{Z}^{T_4}$, it appears that the high values of $\hat{C}_3$ and $\hat{C}_4$ coincides with a short-wavelength oscillation of the track irregularities, in which the sign of the gauge irregularity is opposite to the sign of the three other track irregularities. The fact that this wavelength corresponds to the frequency of highest energy for the transverse movement of the bogie lays stress on the strong dependencies between the track irregularities and the train responses. In the same manner, it could be interesting to find out the reasons of the presence of translated replica for the high values of $\hat{C}_3$ and $\hat{C}_4$ in the second and third limit functions.

6.6 Conclusions

A method to propagate the track geometry variability through railway mechanical simulations is nowadays of great interest. In this chapter, a stochastic model for the track-vehicle system has therefore been presented. Based on a given description of the track design, on a normalized model of a high speed train, on two rail and wheel profiles, and on the stochastic modeling of the track geometry developed in Chapter 5, this model allows the analysis to be carried out concerning the influence of the variability of the track irregularities on the train dynamics. The capability of this stochastic modeling to generate running conditions that are realistic and

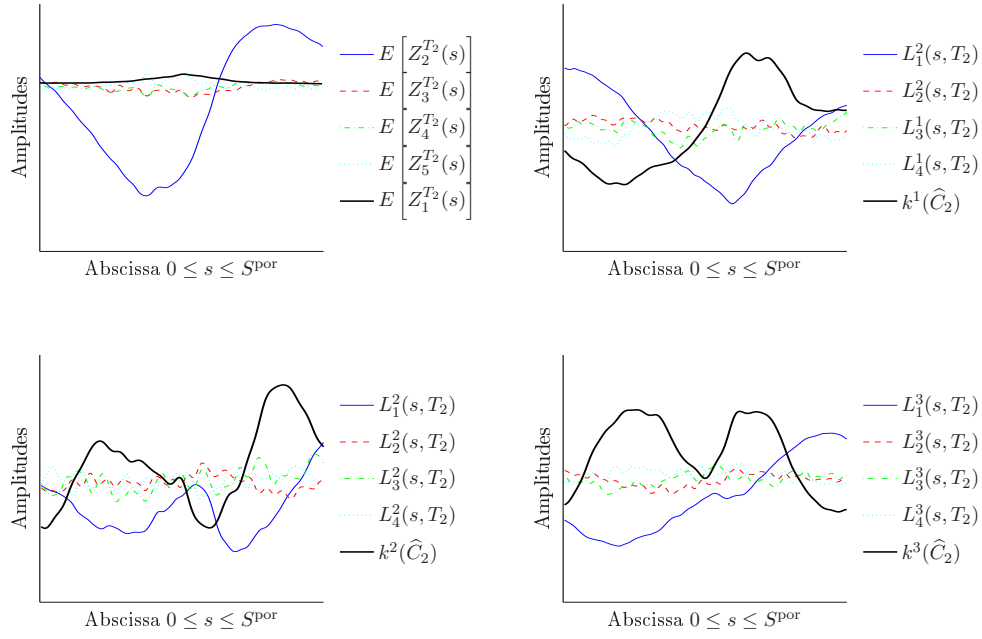


Figure 6.15: Correlation analysis between the shapes of the track irregularity and $\hat{C}_2$.

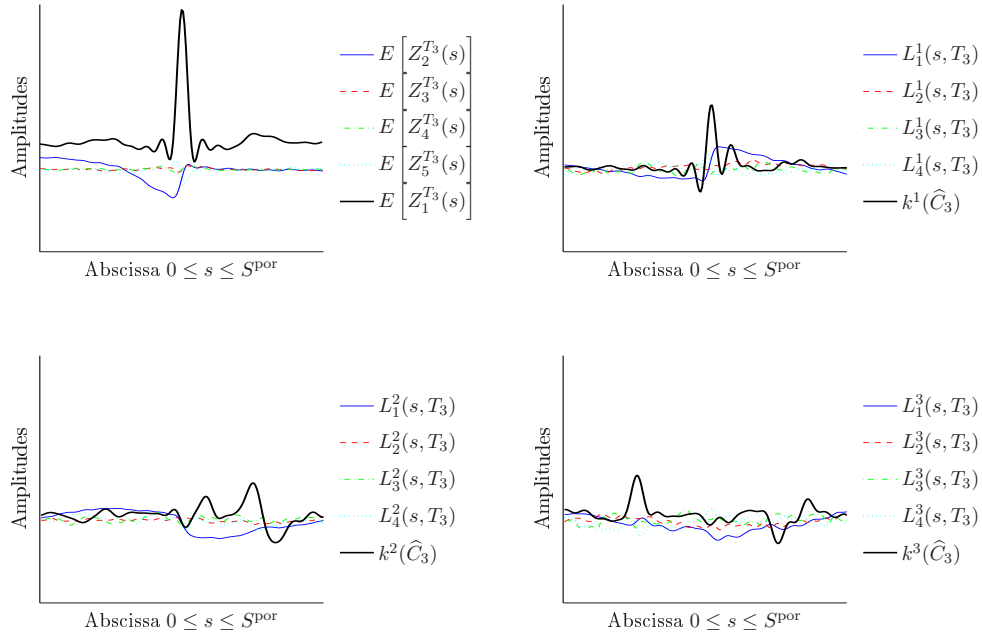


Figure 6.16: Correlation analysis between the shapes of the track irregularity and $\hat{C}_3$.

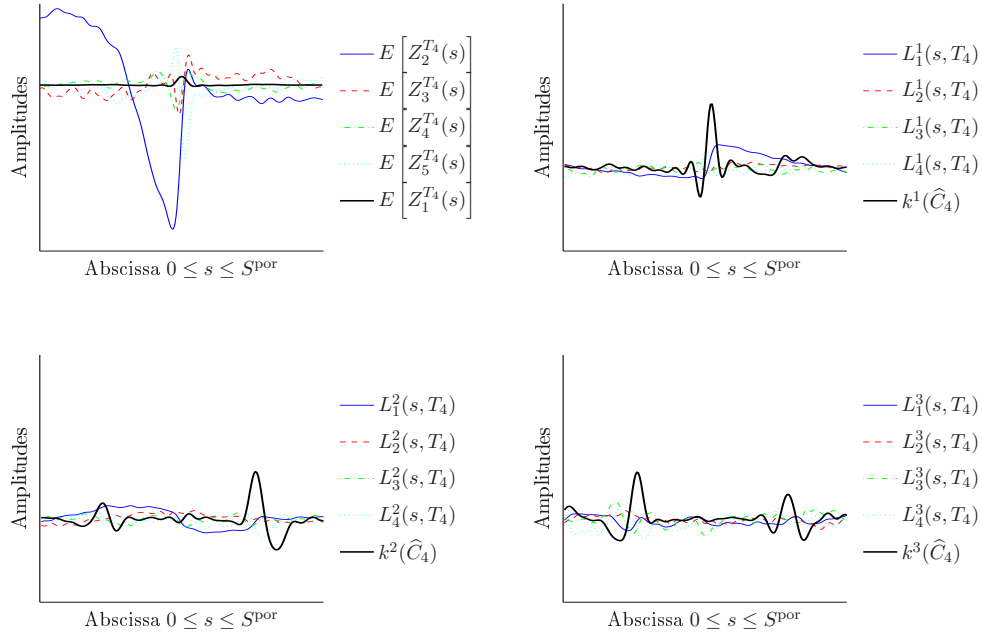


Figure 6.17: Correlation analysis between the shapes of the track irregularity and $\hat{C}_4$.

representative of the quality of a measured railway network has been validated from on-track measurements.

Five quantities of interest have then been introduced to characterize the train dynamic response, which correspond to classical railway comfort and safety criteria. The statistical properties of these dynamic criteria have moreover been identified using a Monte-Carlo approach. Three applications of the whole method have thus been presented. The first one compares the influence of the track design and the track irregularity variability. The second one analyzes the impact of an increase of the speed on the train stability, whereas the third one shows to what extent such an approach could be used to compare competitive high speed trains with respect to their response on a set of representative track conditions.

Finally, it has been underlined that the strong nonlinearity of the track-vehicle system and the high dependencies between the four track irregularities prevent us from identifying clear-cut relations between the five considered dynamic criteria and the track irregularities. In this context, an original method based on the scaled expansion has been presented to open new ways to identify the combined shapes of the track irregularities that could lead high values of the considered criterion to be obtained.

Conclusions and prospects

Summary of the industrial context

For years, the use of simulation in the railway community has been limited to a qualitative analysis approach. Numerical models have therefore been developed for allowing a better understanding of the physical phenomena. Due to increasing available computational resources, and to a series of breakthroughs in the solving of nonlinear equations and in the modeling of complex mechanical systems, simulation nowadays becomes more and more predictive. Hence, simulation cannot only be used to explain the experiments, but is expected to complete them, and sometimes to replace them.

The possibilities of a predictive simulation are huge. In a certification and conception prospect, it could indeed be used to quantify the stability and the safety associated with future trains. In a maintenance prospect, it could moreover allow us to evaluate the consequences of modifications of the running conditions, and to optimize the maintenance policies.

For a railway simulation to be predictive, the mechanical models of the train, of the wheel/rail contact and of the track geometry have to be fully validated from experimental measurements, and the simulations have to be raised on realistic and representative sets of excitations. For the last decades, increasing the modeling precision has been the main priority. Many efforts have therefore been made for the modeling and the identification of the parameters of real and complete trains. In the same manner, real rail and wheel profiles have been used to compute the wheel/rail contact properties. Hence, the comparison between simulated and on-track measured train responses is currently possible. Although not perfect, these deterministic models seem to give very promising results in a large band of frequencies. These models of the train and of the contact being strongly nonlinear, the dynamic behavior of trains has nevertheless to be characterized not from a single simulation but from a set of simulations that is representative of all the running conditions that the train is likely to be confronted to during its lifecycle. A particular attention has thus to be paid to the characterization of the track geometry variability, which represents the main source of excitation of the train dynamics.

From a general point of view, the track geometry can be seen as the sum of a mean line description (which is chosen once for all at the building of a new line) and a deviation from this mean position, which keeps evolving due to the train dynamics and to environmental stresses. Four track irregularities are generally introduced to characterize this deviation, which are the lateral and vertical offset irregularities, X_1 and X_2 , and the cross-level and the gauge irregularities, X_3 and X_4 . These four track irregularities can therefore be seen as a four-dimension random field, $\mathbf{X} = (X_1, X_2, X_3, X_4)$, for which components are strongly dependent. This thesis has therefore been motivated by the need for numerical methods to identify the statistical properties of this random field, as well as to propagate the track variability through the railway system model.

Scientific and industrial contributions

Since 2007, the track irregularities of the French high speed lines are regularly measured, to define a very useful database for the analysis of the track variability. The analysis of these experimental data has however emphasized that, due to the specific interaction between the train and the track irregularities, random field $\mathbf{X}$ is neither stationary nor Gaussian, which raises many difficulties. Under a local-global hypothesis, which has been justified from a convergence analysis, this database can however be decomposed as a finite set of independent realizations of track irregularity random field $\mathbf{X}$. Hence, this thesis deals with numerical methods to identify in inverse the statistical properties of non-stationary and non-Gaussian random fields from a finite set of independent realizations. The chosen methods are based on a double expansion presented hereinafter.

The first step of these methods is a truncated spatial expansion, such that random field $\mathbf{X}$ can be approximated as a finite sum of weighted spatial functions, where the weights are the components of the random vector $\boldsymbol{\eta}$, and are *a priori* dependent. This work has thus proposed contributions in the field of the identification of optimal projection families to condense and reduce the statistical dimension of $\mathbf{X}$ while guaranteeing an acceptable level of accuracy. Based on the classical Karhunen-Loève (KL) expansion, these developments have been motivated by two main reasons. First, as the maximal available information about the track irregularity random field is a finite set of independent realizations, the covariance function of $\mathbf{X}$, on which the KL expansion (and more precisely the Fredholm eigenvalue problem) is based, is unknown and can only be approximated. Although the KL projection basis is optimal in the sense that it minimizes the total mean-square error associated with $\mathbf{X}$, there is no reason for the KL basis associated with the approximation of the covariance function of $\mathbf{X}$ to be still optimal. When the number of available realizations is moreover very small compared to the stochastic dimension of the random field, as it is the case for the track irregularity random field, $\mathbf{X}$, it has been shown that the relevance of such projection basis can be very limited. In this prospect, an original method based on an optimization problem over the operator on which the Fredholm problem is solved has been introduced. Then, it has been underlined that minimizing the total mean-square error associated with $\mathbf{X}$ amounts to characterizing in priority the components of $\mathbf{X}$ that have the highest signal energy, even if their role on the train dynamics is low. An innovative scaled expansion has thus been proposed in this work, in order to reduce this bias and to minimize the maximal value of the errors associated with each component of $\mathbf{X}$. The interests brought by these two adaptations of the classical KL expansion in terms of error reduction have then been illustrated on simple examples as well as on the track irregularity case.

Once the optimal projection family for $\mathbf{X}$ has been identified, characterizing $\mathbf{X}$ amounts to identifying the multidimensional distribution of random vector $\boldsymbol{\eta}$, for which a set of independent realizations can be deduced from the realizations of $\mathbf{X}$. To this end, this work focused on the Polynomial Chaos Expansion (PCE) method, which is one of the currently most promising method to identify in inverse the distribution of non-Gaussian random vectors from independent realizations. This method is based on the projection of $\boldsymbol{\eta}$ on a polynomial basis of its probability space. In such a projection, the polynomial basis is random, but its distribution is chosen and known, whereas the projection coefficients are deterministic but unknown. This sum, which is infinite in theory, has then to be truncated, and the truncation has to be justified according to a convergence analysis. Once a given truncated projection family has been chosen, the finite set of coefficients has finally to be identified to completely characterize the distribution of $\boldsymbol{\eta}$. A good approach to identify such coefficients is to search them as the arguments that maximize

the likelihood of $\boldsymbol{\eta}$ at its available realizations. The likelihood function being not convex, it has been shown in this work that the solving of this optimization problem can be carried out using a random search algorithm based on the generation of rectangular matrices under orthogonality constraints.

In the case of a track, the train being very sensitive to the track irregularities on a large frequency band, a large number of projection functions are needed for the approximated projection of $\mathbf{X}$ to be accurate from the train response point of view, such that the dimension of $\boldsymbol{\eta}$ is very high. In such high dimensions, two adaptations of the classical formulation have therefore been presented to give relevant results. On the first hand, original iterative algorithms have been proposed to optimize the trials of these projection matrices under orthogonality constraints. On the second hand, a method to numerically stabilize the matrix of realizations of the statistical polynomial basis has been introduced, to allow more relevant convergence analysis. The possibilities opened by these two adaptations have also been illustrated on academical examples and on the track irregularity case.

Once the spatial projection family and the statistical projection coefficients have been identified from the available track database, the multidimensional distribution of $\mathbf{X}$ is characterized, and it is possible to generate quickly and easily independent realizations of $\mathbf{X}$. Coupled to a particular track design, these realizations allow the definition of realistic running conditions, which are representative of the quality of the measured railway network. The track generator can finally be used in any railway software to investigate the influence of the track variability on the train dynamics.

In this prospect, the multibody commercial software Vampire has been used to compute the dynamic responses of high speed trains on generated track geometries. After having shown that the simulated response of high speed trains on generated tracks was similar to the measured response of the same trains on measured tracks, this work analyzed the influence of the track variability on two comfort criteria, two safety criteria and a wear criterion. These studies underlined the high nonlinearity of the track/vehicle system, and quantify the influence of modifications of the running conditions and to evaluate and compare the stability and the aggressiveness of several high speed trains.

Discussion and perspectives

Scientific prospects

The application of the Karhunen-Loève expansion combined with the PCE approach to the modeling of the track irregularity random field revealed the important potential of such methods to identify in inverse the distribution of multivariate, non-Gaussian and non-stationary random fields in very high dimensions, but also emphasized some limitations that are listed hereinafter.

Reduction of the computational time associated with the identification of the optimal projection family. When confronted to multivariate random fields $\mathbf{X}$ that are characterized by a set of ν independent realizations, the method proposed to identify the optimal basis is established on three interlocked computational loops.

- The first one deals with the identification of the optimal value of the weight matrix $[\alpha]$, defined by Eq. (2.40). Let N_α be the number of values for $[\alpha]$ that are considered in this solving.

- Secondly, for each value of $[\alpha]$, the optimal value for the scaling vector $\mathbf{O}$, introduced in Section 4.2.2, is searched to minimize the maximal value of the errors associated with each components of $\mathbf{X}$. In the same manner, let N_O be the number of evaluations of $\mathbf{O}$ that are required to reach a targeted error.
- Then, for each value of $[\alpha]$ and $\mathbf{O}$, a method to evaluate the representativeness error associated with the considered projection family is needed. In this work, it has therefore been shown that this error can be computed from a Leave-One-Out (LOO) approach. As shown in Section 2.3.3, this method is nevertheless based on the solving of ν Fredholm eigenvalue problems.

Finally, the entire identification method requires the solving of $N_\alpha \times N_O \times \nu$ Fredholm problems. In this prospect, specific algorithms have been used in this work to reduce N_α and N_O , and to speed up the solving of the Fredholm eigenvalue problems associated with the evaluation of the LOO error. First, a dichotomy-based algorithm has been proposed for the identification of $[\alpha]$, as it allowed the identification of very accurate results in a very few number of iterations for the analyzed applications. The convexity of this problem over $[\alpha]$ has however not been proved in the general case, which could be an interesting prospect of the proposed work. If the convexity property is verified, it is expected that more advanced algorithms could be used to speed up this optimization step. In the same manner, an innovative iterative algorithm has been proposed to identify accurate solutions from a very limited number of iterations, which is denoted by N_O . Once again, the proof of the convergence of such an iterative algorithm in the general case is missing. This algorithm is moreover based on two parameters, τ and γ (see Eq. (2.51)), which play a major role on the convergence speed. In this manuscript, two values have been proposed, which stem from a quick parametric analysis. The optimization of these values constitutes another direction to minimize the total computational cost. At last, keeping in mind that computing the LOO error associated with any projection family amounts to solving a series of slightly modified eigenvalue problems, it was noticed that iterative methods, such as the subspace iteration methods [85], helped us to reduce drastically the computational time, especially when confronted to very high dimensional cases. Further developments in the field of these efficient identification methods would thus be of a great interest.

More gains can also be expected from the definition of rejection procedures in all these loops (for instance, if the solving of the first Fredholm problems, associated with given values for $[\alpha]$ and $\mathbf{O}$, seems to indicate non-satisfactory results, we directly move to other values for them), but also in the definition of efficient identification methods for the couple $([\alpha], \mathbf{O})$ instead of one for $[\alpha]$ and one for $\mathbf{O}$.

Application of the decomposition method to the PCE identification. The identification of the PCE for very high dimensional random vectors is also very time consuming. As a random search method has been proposed to identify these PCE coefficients, the precision of the results is completely driven by the available computational resources, as the more elements we try, the more chance we have to find accurate values. Even if such random-search algorithm can easily be distributed on several computers, it is important to optimize the trials. In this context, the advantages of a method based on a line-by-line identification has been shown in Section 3.2. The most time-consuming step of the numerical application of this method is the computation of the multidimensional likelihood at the available experimental points. Indeed, the complexity of the evaluation of this quantity, whose expression is given by Eq. (1.60), is $M \times \nu \times \nu^{\text{chaos}}$, where ν is the number of available measurements, M is the dimension of the considered random vector, that we denote by $\boldsymbol{\eta} = (\eta_1, \dots, \eta_M)$, and ν^{chaos} is the number of

independent realizations on which the computation of the multidimensional PDF of the PCE approximation of $\boldsymbol{\eta}$, $\boldsymbol{\eta}^{\text{chaos}}$, is based. The choice for ν^{chaos} is thus very dependent on the value of M , as the higher M is, the more independent realizations ν^{chaos} we need to evaluate the PDF of $\boldsymbol{\eta}^{\text{chaos}}$ from a nonparametric approach. In this prospect, we believe that the application of decomposition methods for the PCE identification could lead to considerable improvements in terms of cost-efficiency. In other words, if $\boldsymbol{\eta}^{(q)}$, $1 \leq q \leq Q$, refer to M/Q -dimension random vectors such that $\boldsymbol{\eta} = (\boldsymbol{\eta}^{(1)}, \dots, \boldsymbol{\eta}^{(Q)})$, we think that it could be interesting to search the PCE of $\boldsymbol{\eta}$ from the aggregation of the Q PCE of $\boldsymbol{\eta}^{(q)}$. Indeed, the dimension of $\boldsymbol{\eta}^{(q)}$ being much smaller, the associated number ν^{chaos} of generated realizations at each evaluation step could be drastically reduced.

In such an approach, the definition of methods to perform such an aggregation is however an opened subject.

Need for innovative methods to compare the statistical properties of two sets of independent realizations in very high dimension. To characterize the distribution of the track irregularity random field, which was the initial goal of the thesis, we finally had to identify the distribution $P_{\boldsymbol{\eta}}$ of a 2,000-dimension random vector, $\boldsymbol{\eta}$, while the maximal available information about this random vector was a finite set of almost 500 independent realizations. More precisely, the ability of generating independent realizations of $\boldsymbol{\eta}$ was as important as the identification of $P_{\boldsymbol{\eta}}$.

Hence, being confronted to such very high dimensional problems with so little information, we do not pretend to be able to identify exactly $P_{\boldsymbol{\eta}}$, but propose a method to search its best **reachable** approximation. In the same manner, we don't claim to be able to generate new realizations of $\boldsymbol{\eta}$, but try to generate sets of independent realizations that have the closest statistical properties to the available set of measurements.

In this context, the PCE approach presents many advantages to face this challenge of the high dimension. First, it is very general, in the sense that whatever the dimension is, no subjective assumption is needed. It seems moreover particularly able to take advantage of the increasing computational resources, as the relevance of the results increases with the number of tested trials. At last, once the projection coefficients are identified, the generation of independent realizations of the PCE approximation, $\boldsymbol{\eta}^{\text{chaos}}(N)$, of $\boldsymbol{\eta}$ is quick and very easy.

However, even in this favorable case, where it is possible to generate as many realizations of $\boldsymbol{\eta}^{\text{chaos}}(N)$ as needed, the relevance of $\boldsymbol{\eta}^{\text{chaos}}(N)$ remains difficult to evaluate. The number of realizations of $\boldsymbol{\eta}$ being still small, it is still difficult to compare the dependencies associated with the components of $\boldsymbol{\eta}$ and the ones associated with the components of $\boldsymbol{\eta}^{\text{chaos}}(N)$. Statistical tests and likelihood-based methods to compare these sets are indeed completely useless in such high dimensions.

The fact that this random vector is to be used in a very complex and nonlinear mechanical problem is however an interesting opportunity to compare $\boldsymbol{\eta}$ and $\boldsymbol{\eta}^{\text{chaos}}(N)$. Indeed, if the mechanical problem makes use of the dependencies between the components of $\boldsymbol{\eta}$, it should be possible to quantify the distance between the multidimensional distributions of $\boldsymbol{\eta}$ and $\boldsymbol{\eta}^{\text{chaos}}(N)$, by comparing the outputs of this problem that correspond to $\boldsymbol{\eta}$ on the first hand, and to $\boldsymbol{\eta}^{\text{chaos}}(N)$ on the second hand. In this context, in Section 3.4.3, an original method to compare the dependencies between $\boldsymbol{\eta}$ and $\boldsymbol{\eta}^{\text{chaos}}(N)$ in very high dimension has been proposed. This method is based on the generation of a series of random fields, which are written as weighted sums of randomly chosen spatial functions, while the weights are the components of $\boldsymbol{\eta}$ and $\boldsymbol{\eta}^{\text{chaos}}(N)$. By generating large sets of these functions, and by comparing these random fields on quantities that can actually be evaluated (the number of upcrossings for instance), it is possible to investigate the capability of $\boldsymbol{\eta}^{\text{chaos}}(N)$ to represent the dependencies between the

components of $\boldsymbol{\eta}$.

At last, it is believed that the definition of problems that are more specific and more adapted to the dependency structure of $\boldsymbol{\eta}$ could help us to construct methods to precisely evaluate the quality of PCE approximations in very high dimension.

Scaled expansion and shape correlations. In Section 6.5.2, an innovative sensitivity method based on the KL expansion has been proposed to analyze the correlation between an output function and a multivariate input function. Applied to a series of mechanical system, this method seems to give very promising results. In a validation prospect, analyzing the theoretical basis of this method is however an open topic.

Model updating. As the measurement train IRIS 320 continuously monitors the track geometry, the number of experimental data for the track geometry modeling progressively increases with respect to time. As a consequence, the number of independent realizations of the track-geometry random field $\mathbf{X}$ is likely to increase. From a prior stochastic modeling of $\mathbf{X}$ that is based on an original set of realizations, methods to identify updated modelings of $\mathbf{X}$ that takes into account new available realizations of $\mathbf{X}$ would be very interesting. To this end, the recent adaptations of the Bayes theorem to the PCE seem to be very promising (see [50] and [86] for further details about these adaptations), and it would be worth applying them on the track geometry, whose dimension is very high.

Industrial prospects

The stochastic modeling of the track geometry opens many opportunities in terms of certification, optimization of the railway system and minimization of the maintenance costs. In order to extend the domain of application of this modeling, a series of complementary developments could be carried out as explained below.

Evaluation of the global quality of several railway networks with respect to the train dynamics. The stochastic modeling of the track geometry we propose is only based on the measurements of the track geometry of a given railway network. From the local-global approach, these measurements can then be sorted with respect to the track design in sets of track portions of same length S . For each of these sets, once the spatial and statistical expansions have been achieved, it is possible to generate as many track conditions as needed to evaluate the stability and the aggressiveness associated with each train that could run on this network.

Reciprocally, such a method can directly be used to compare the global quality of several high speed lines from the train response point of view. Indeed, once the track generator associated with a series of railway networks have been computed, sets of track conditions that are representative of their qualities can be generated. Railway simulations can then be carried out from these realistic running conditions, such that the quality of the different networks can finally be compared by analyzing the associated distributions of comfort and safety criteria.

Quantification of the degradation of the track quality due to the train dynamics. In Section 6.5.2, the strong dependencies between the train dynamics and the track irregularities were pointed out. Many informations about the train mechanical properties could indeed be found back by analyzing the wavelengths of the track irregularities, or by analyzing the most frequent positions for the damaged track sections. Modeling the coupling between the train dynamics and the time evolution of the track irregularities is however still an open issue that

was not treated in this work but that would require a particular attention. Two different points of view can be analyzed.

On the first hand, if the track geometry of a whole network is measured between two maintenance operations at different time steps $t_1, \dots, t_N$, a stochastic modeling of the track geometry for each of these time step can be computed. Therefore, it should be possible to evaluate the time evolutions of the distribution of the comfort and safety criteria, and therefore to quantify the influence of the train dynamics on the global quality of the considered network. In the same manner, by continuously monitoring the evolution of the track geometries, the influence of the runs of trains on the track irregularities could be evaluated.

But much more gain could be expected from a predictive coupling model. Indeed, if it is possible to predict (from physical and/or statistical models) the future evolution of the track geometry due to the train dynamics, the continuous monitoring of the track can not only be used to identify the track sections that would cause the highest train responses, but should also indicate the track portions that are still not dangerous but which are the most likely to become critical if no maintenance operation is planned.

Robust optimization. Based on the track generator that was proposed in this work, it is possible to construct huge sets of realistic track geometries, which would correspond to track conditions that the train can be confronted to during its lifecycle. This track stochastic modeling opens therefore new possibilities for the train manufacturers and train operators to optimize the mechanical properties of trains with respect to their stochastic dynamic response on the variable track geometry.

External communications

The work achieved in this thesis has given rise to a series of communications, which are listed hereinafter.

Referred journal publications

- C. Funfschilling, G. Perrin, S. Kraft. Propagation of variability in railway dynamic simulations : application to virtual homologation. *Vehicle System Dynamics*, 2012, 50, Pages: 245-261.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Identification of Polynomial Chaos Representations in High Dimension from a Set of Realisations. *SIAM - Journal on Scientific Computing*, 2012, 34 (6), Pages: 2917-2945.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Karhunen-Loève expansion revisited for vector-valued random fields: scaling, errors and optimal basis. *Journal of Computational Physics*, 2013, 242 (1), Pages: 607-622.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Track Geometry Stochastic Modeling. *Probabilistic Engineering Mechanics*, accepted in July 2013.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. A posteriori error and optimal reduced basis for stochastic processes defined by a finite set of realizations. *SIAM/ASA - Journal on Uncertainty Quantification*, submitted in January 2013.

- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Quantification of the influence of the track geometry on the train dynamics and stability from experimental measurements. Mechanical Systems and Signal Processing, in preparation.

Conferences with proceedings

- N. Lestaille, G. Perrin, C. Funfschilling. Characterization of the influence of rail wear on train dynamics. COMPDYN 2013, 4th ECCOMAS Thematic Conference on Computational Methods in Structural Dynamics and Earthquake Engineering, June 2013, Kos Island, Greece. Proceedings of COMPDYN 2013.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Dynamical behavior of trains excited by a non-Gaussian vector-valued random field. COMPDYN 2013, 4th ECCOMAS Thematic Conference on Computational Methods in Structural Dynamics and Earthquake Engineering, June 2013, Kos Island, Greece. Proceedings of COMPDYN 2013.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Statistical inverse problems for non-Gaussian vector valued random fields with a set of experimental realizations. ICOSSAR 2013, 11th International Conference on Structural Safety and Reliability, June 2013, New-York, United States. Proceedings of ICOSSAR 2013.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Influence of the track geometry variability on the train behavior. Vienna University of Technology, Vienna. Congress on Computational Methods in Applied Sciences and Engineering (ECCOMAS 2012), September 2012, Vienna, Austria. Vienna University of Technology, Vienna, Proceedings of long abstracts of ECCOMAS 2012.
- G. Perrin, C. Soize, D. Duhamel, C. Funfschilling. Modeling the track geometry variability. University of Sao Paulo. 10th World Congress on Computational Mechanics(WCCM), July 2012, Sao Paulo, Brazil. Proceedings of WCCM 2012.
- G. Perrin, D. Duhamel, C. Soize, C. Funfschilling. Track Irregularities Stochastic Modelling. Thirteenth International Conference on Civil, Structural and Environmental Engineering Computing, September 2011, Chania, Greece. Proceedings of the Thirteenth International Conference on Civil, Structural and Environmental Engineering Computing.
- G. Perrin, C. Funfschilling, B. Sudret. Propagation of Variability in Railway Dynamics Simulations. 11th International Conference on Applications of Statistics and Probability in Civil Engineering, August 2011, Zurich, Switzerland. Proceedings of the 11th International Conference on Applications of Statistics and Probability in Civil Engineering.

Conferences without proceedings

- C. Funfschilling, G. Perrin, C. Soize, D. Duhamel. Dynamical Behavior of Trains Excited by a Non-Gaussian Vector Valued Random Fields. Multibody Dynamics 2013, July 2013, Zagreb, Croatia.
- G. Perrin, D. Duhamel, C. Soize, C. Funfschilling. Modeling and Identification of non-Gaussian multivariate random fields and application to the excitation of trains by the track irregularities. Premières Journées des Jeunes Chercheurs en Vibrations, April 2013, Marne-la-Vallée, France.

- G. Perrin, D. Duhamel, C. Soize, C. Funfschilling. Identification of polynomial chaos representations in high dimension. SIAM Conference on Uncertainty Quantification, April 2012, Raleigh, North Carolina, United States.

Appendix

A Proof of Lemma 2

Using the notations of Section 4.2, $\{\mathbf{k}^i(\mathbf{O}), i \geq 1\}$ defines a spatially orthonormal basis of $\mathcal{P}^{(Q)}(\Omega)$. Autocorrelation function $[R_{\mathbf{Y}\mathbf{Y}}(\mathbf{O})]$ can therefore be projected on this basis, such that, by construction of the Karhunen-Loève basis:

$$[R_{\mathbf{Y}\mathbf{Y}}(\mathbf{O})] = \sum_{i \geq 1} \lambda_i(\mathbf{O}) \mathbf{k}^i(\mathbf{O}) \otimes \mathbf{k}^i(\mathbf{O}). \quad (6.20)$$

Let $\mathcal{B} = \{\mathbf{b}^i, 1 \leq i\}$ be another countable basis of Hilbertian space $\mathcal{P}^{(Q)}(\Omega)$, and $\mathcal{F}^{(M)} = \{\mathbf{b}^i, 1 \leq i \leq M\}$ be a M -dimension subset of $\mathcal{B}$. For all $i \geq 1$, $\mathbf{f}^i$ can then be projected on $\{\mathbf{k}^j(\mathbf{O}), j \geq 1\}$:

$$\mathbf{f}^i = \sum_{j \geq 1} P_{ij} \mathbf{k}^j(\mathbf{O}), \quad P_{ij} = (\mathbf{f}^i, \mathbf{k}^j(\mathbf{O})). \quad (6.21)$$

Without loss of generality, family $\mathcal{F}$ can be supposed to be spatially orthonormal, as it can be orthonormalized a posteriori without modifying the corresponding projection error. From Eqs. (4.19), this yields:

$$1 = (\mathbf{f}^i, \mathbf{f}^i) = \sum_{j \geq 1} \sum_{\ell \geq 1} P_{ij} P_{i\ell} (\mathbf{k}^j(\mathbf{O}), \mathbf{k}^\ell(\mathbf{O})) = \sum_{j \geq 1} P_{ij}^2. \quad (6.22)$$

Let $\tilde{\mathbf{Y}}^{(M)}$ be the projection of random field $\mathbf{Y} = [\text{Diag}(\mathbf{O})]\mathbf{X}$ on $\mathcal{F}^{(M)}$:

$$\tilde{\mathbf{Y}}^{(M)} = \sum_{i=1}^M \mathbf{f}^i C_i, \quad C_i = (\mathbf{Y}, \mathbf{f}^i). \quad (6.23)$$

Random field $\widetilde{\mathbf{X}}^{(M)}$ is thus introduced as:

$$\widetilde{\mathbf{X}}^{(M)} = [\text{Diag}(\mathbf{O})]^{-1} \tilde{\mathbf{Y}}^{(M)}. \quad (6.24)$$

From Eqs. (6.20) and (6.23), for all $M \geq i \geq 1$, we get:

$$E \{C_i^2\} = \int_{\Omega^2} (\mathbf{f}^i(s))^T [R_{\mathbf{Y}\mathbf{Y}}(\mathbf{O}, s, s')] \mathbf{f}^i(s') ds ds' = \sum_{j \geq 1} \lambda_j(\mathbf{O}) P_{ij}^2. \quad (6.25)$$

Therefore, from Eqs (6.22) and (6.25):

$$\begin{aligned}
\sum_{i=1}^M (\lambda_i(\mathbf{O}) - E\{C_i^2\}) &= \sum_{i=1}^M \lambda_i(\mathbf{O}) \sum_{j \geq 1} P_{ij}^2 - \sum_{i=1}^M \sum_{j \geq 1} \lambda_j(\mathbf{O}) P_{ij}^2 \\
&= \sum_{i=1}^M \sum_{j \geq 1} P_{ij}^2 (\lambda_i(\mathbf{O}) - \lambda_j(\mathbf{O})) \\
&= \sum_{i=1}^M \sum_{j \geq M+1} P_{ij}^2 (\lambda_i(\mathbf{O}) - \lambda_j(\mathbf{O})) \\
&\geq (\lambda_M(\mathbf{O}) - \lambda_{M+1}(\mathbf{O})) \sum_{i=1}^M \sum_{j \geq M+1} P_{ij}^2 \geq 0,
\end{aligned} \tag{6.26}$$

as by construction, for all $j \geq i$, $\lambda_j(\mathbf{O}) \leq \lambda_i(\mathbf{O})$. Moreover, it can be noticed that, by definition of matrix $[\text{Diag}(\mathbf{O})]$, for $1 \leq q \leq Q$:

$$\begin{aligned}
E\left\{\left(X_q - \widehat{X}_q^{(M)}, X_q - \widehat{X}_q^{(M)}\right)\right\} &= E\left\{\left(O_q^{-1}\left(Y_q - \widehat{Y}_q^{(M)}\right), O_q^{-1}\left(Y_q - \widehat{Y}_q^{(M)}\right)\right)\right\} \\
&= O_q^{-2} \sum_{M+1 \leq i} \lambda_i(\mathbf{O}) (k_q^i(\mathbf{O}), k_q^i(\mathbf{O})),
\end{aligned} \tag{6.27}$$

where it is reminded that $\widehat{\mathbf{Y}}^{(M)}$ is the projection of $\mathbf{Y} = [\text{Diag}(\mathbf{O})]\mathbf{X}$ on $\mathcal{K}^{(M)}(\mathbf{O})$ and $\widehat{\mathbf{X}}^{(M)} = [\text{Diag}(\mathbf{O})]^{-1} \widehat{\mathbf{Y}}^{(M)}$. In the same manner:

$$\begin{aligned}
E\left\{\left(X_q - \widetilde{X}_q^{(M)}, X_q - \widetilde{X}_q^{(M)}\right)\right\} &= E\left\{\left(O_q^{-1}\left(Y_q - \widetilde{Y}_q^{(M)}\right), O_q^{-1}\left(Y_q - \widetilde{Y}_q^{(M)}\right)\right)\right\} \\
&= O_q^{-2} \sum_{M+1 \leq i} E\{C_i^2\} (f_q^i, f_q^i).
\end{aligned} \tag{6.28}$$

It can finally be deduced from Eqs. (6.26), (6.27) and (6.28) that:

$$\begin{aligned}
&\sum_{q=1}^Q O_q^2 \mathcal{N}^2(X_q) \varepsilon_q^2 \left(\mathcal{K}^{(M)}(\mathbf{O})\right) - \sum_{q=1}^Q O_q^2 \mathcal{N}^2(X_q) \varepsilon_q^2 (\mathcal{F}^{(M)}) \\
&= \sum_{q=1}^Q O_q^2 \left[E\left\{\left(X_q - \widehat{X}_q^{(M)}, X_q - \widehat{X}_q^{(M)}\right)\right\} - E\left\{\left(X_q - \widetilde{X}_q^{(M)}, X_q - \widetilde{X}_q^{(M)}\right)\right\} \right] \\
&= \sum_{M+1 \leq i} \left[\lambda_i(\mathbf{O}) \sum_{q=1}^Q (k_q^i(\mathbf{O}), k_q^i(\mathbf{O})) - E\{C_i^2\} \sum_{q=1}^Q (f_q^i, f_q^i) \right] \\
&= \sum_{M+1 \leq i} [\lambda_i(\mathbf{O}) - E\{C_i^2\}] \\
&= \sum_{i=1}^M [E\{C_i^2\} - \lambda_i(\mathbf{O})] \\
&\leq 0.
\end{aligned} \tag{6.29}$$

This result being true for all family $\mathcal{F}^{(M)}$ in $\mathbb{H}^M$, family $\mathcal{K}^{(M)}(\mathbf{O})$ is thus M -optimal for $\mathbf{X}$ regarding error $\sum_{q=1}^Q O_q^2 \mathcal{N}^2(X_q) \varepsilon_q^2$.

q	c_q	$\omega_q S/(2\pi)$	ℓ_q/S	$T_q S/(2\pi)$
1	1	20%	20%	5
2	0.5	30%	25%	7
3	0.25	20%	35%	8
4	0.1	30%	40%	10

Figure 6.18: Numerical values used in the definition of autocorrelation matrix $[R_{\mathbf{X}\mathbf{X}}]$.

B Generation of the matrix-valued autocorrelation matrix

For $1 \leq p, q \leq 4$, matrix-valued autocorrelation function $[R_{\mathbf{X}\mathbf{X}}]$ is chosen such that:

$$[R_{\mathbf{X}\mathbf{X}}(s, s')]_{pq} = \frac{c_p c_q (1 + \delta_{pq})}{2} \sum_{k=1}^{200} \sqrt{\lambda_k^{(p)} \lambda_k^{(q)}} d_k^{(p)}(s) d_k^{(q)}(s'), \quad \forall (s, s') \in [0, 1]^2, \quad (6.30)$$

where for all $1 \leq k \leq 200$:

$$\int_0^1 h_p(s, s') d_k^{(p)}(s') ds' = \lambda_k^{(p)} d_k^{(p)}(s), \quad (6.31)$$

$$h_p(s, s') = \exp(-|s - s'|/\ell_p) \cos(\omega_p |s - s'|) \cos(T_p s), \quad (6.32)$$

$$\lambda_k^{(p)} \geq \lambda_{k+1}^{(p)} > 0, \quad (6.33)$$

$$(d_k^{(p)}, d_k^{(q)}) = \delta_{pq}. \quad (6.34)$$

The numerical values of vectors $\mathbf{c} = (c_1, \dots, c_4)$, $\boldsymbol{\omega} = (\omega_1, \dots, \omega_4)$, $\boldsymbol{\ell} = (\ell_1, \dots, \ell_4)$, $\mathbf{T} = (T_1, \dots, T_4)$ are gathered in Figure 6.18. Several comments can be made about this formalism.

- Application $(s, s') \mapsto h_q(s, s')$ is not necessary positive-definite regarding the chosen numerical parameters, but only its 200 highest strictly positive eigenvalues, $\{\lambda_k^{(q)}, 1 \leq k \leq 200\}$, are considered.
- Couples $\{\lambda_k^{(q)}, d_k^{(q)}\}$ are solutions of the Fredholm problem associated with h_q , but are not solutions of the Fredholm problem associated with $[R_{\mathbf{X}\mathbf{X}}]$.
- Coefficient c_q^2 can be related to the signal energy of X_q , such that if $c_p > c_q$, $\mathcal{N}^2(X_p) > \mathcal{N}^2(X_q)$.
- Coefficient $2\pi/\omega_q$ can be considered as a pseudo-wavelength for the mean-squared stationary part of $[R_{\mathbf{X}\mathbf{X}}]_{pq}$.
- Coefficient ℓ_q can be seen as the auto-correlation length of X_q .
- Coefficient T_q is introduced as a perturbation for $[R_{\mathbf{X}\mathbf{X}}]_{pq}$, such that the smaller T_q is, the less mean-squared stationary $[R_{\mathbf{X}\mathbf{X}}]_{pq}$ is.

C Definition of the local-global error functions

It is assumed that ν track portions of same length L , $\{\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(\nu)}\}$, have been collected from the available measurements of the railway network of interest. For any value for S , ν new track geometries, $\{\mathbf{y}^{(1)}(S), \dots, \mathbf{y}^{(\nu)}(S)\}$, of total length L , are then built from the concatenation of track subsections of length S that have randomly been chosen in $\{\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(\nu)}\}$.

For (s, s') in $[0, L]^2$ and $f \geq 1/L$, let $(s, s') \mapsto [R_{zz}(s, s')]$, $(s, s') \mapsto [R_{yy}(s, s', S)]$, $f \mapsto \Sigma_z(f)$ and $f \mapsto \Sigma_y(f, S)$ be the following quantities:

$$[R_{zz}(s, s')] = \frac{1}{\nu} \sum_{n=1}^{\nu} \mathbf{z}^{(n)}(s) \mathbf{z}^{(n)}(s')^T, \quad (6.35)$$

$$[R_{yy}(s, s', S)] = \frac{1}{\nu} \sum_{n=1}^{\nu} \mathbf{y}^{(n)}(s, S) \mathbf{y}^{(n)}(s', S)^T, \quad (6.36)$$

$$\Sigma_z = \sqrt{\frac{1}{\nu} \sum_{n=1}^{\nu} PSD(\mathbf{z}^{(n)})}, \quad \Sigma_y(S) = \sqrt{\frac{1}{\nu} \sum_{n=1}^{\nu} PSD(\mathbf{y}^{(n)}(S))}, \quad (6.37)$$

where $PSD(\mathbf{z}) = (PSD(z_1), \dots, PSD(z_P))$ is the power spectral density estimation of any $\mathbb{R}^P$ -valued function $\mathbf{z} = (z_1, \dots, z_P)$. For any value of S in $[0, L]$, errors $err_{\text{cov}}^2(S)$ and $err_{\text{spect}}^2(S)$, which have been introduced in Section 5.2.2 are then defined by:

$$err_{\text{cov}}^2(S) = \|[R_{zz}] - [R_{yy}(S)]\|_M^2 / \|[R_{zz}]\|_M^2, \quad (6.38)$$

$$err_{\text{spect}}^2(S) = \|\Sigma_z - \Sigma_y(S)\|_V^2 / \|\Sigma_z\|_V^2, \quad (6.39)$$

where, for all $(P \times P)$ matrix-valued function $[R]$, and for all $\mathbb{R}^P$ -valued function Σ ,

$$\|[R]\|_M^2 = \int_0^L \int_0^L \text{Tr}([R(s, s')][R(s, s')]^T) ds ds', \quad (6.40)$$

$$\|\Sigma\|_V^2 = \int_{1/L}^{+\infty} \Sigma(f)^T \Sigma(f) df. \quad (6.41)$$

Hence, on the first hand, $err_{\text{cov}}^2(S)$ corresponds to a covariance error. On the other hand, $err_{\text{spect}}^2(S)$ can be seen as a spectral error, which characterizes the impact of S on the frequency content of the track irregularities.

Bibliography

- [1] C. Soize. *Stochastic Models of Uncertainties in Computational Mechanics*. American Society of Civil Engineers, Reston, 2013.
- [2] A. Dutfoy. Reference guide. Tutorial, Open TURNS (www.openturns.org), 2013.
- [3] C. Funfschilling, G. Perrin, and S. Kraft. Propagation of variability in railway dynamic simulations: application to virtual homologation. *Vehicle System Dynamics*, 50:245–261, 2012.
- [4] S. Iwnicki. *Handbook of Railway Vehicle Dynamics*. Taylor and Francis, 2006.
- [5] P. Aknin. Outils de description de la géométrie des voies et déconvolution des relevés expérimentaux, application aux tracés ferroviaires. Technical report, INRETS-LTN, décembre 1995.
- [6] J.J. Kalker. *Three-Dimensional elastic bodies in rolling contact*. KLUWER ACADEMIC PUBLISHER, 1990.
- [7] H. Hertz. Ueber die berührung fester elastischer körper. *Journal für reine und angewandte Mathematik*, 92, 1882.
- [8] E. T. Jaynes. Information theory and statistical mechanics. *The Physical Review*, 106(4):620–630, 1963.
- [9] C. Soize. Construction of probability distributions in high dimension using the maximum entropy principle. applications to stochastic processes, random fields and random matrices. *International Journal for Numerical Methods in Engineering*, 76(10):1583–1611, 2008.
- [10] J. M. Bernardo and A. F. M. Smith. *Bayesian Theory*. John Wiley and Sons, Chichester, 2000.
- [11] B. P. Carlin and T. A. Louis. *Bayesian Methods for Data Analysis, Thrid Edition*. CRC Press, Boca Raton, 2009.
- [12] P. Congdon. *Bayesian Statistical Modelling, Second Edition*. John Wiley and Sons, Chichester, 2007.
- [13] C. Soize. Random matrix theory for modeling unvertainties in computational mechanics. *Comput. Methods Appl. Mech. Engrg.*, 194:1333–1366, 2005.
- [14] C. Soize. A comprehensive overview of a non-parametric probabilistic approach of model uncertainties for predictive models in structural dynamics. *Journal of sound and vibration*, 288:623–652, 2005.

- [15] R. Y. Rubinstein and D. P. Kroese. *Simulation and the Monte Carlo Method, Second Edition*. Wiley, 2007.
- [16] P. Krée and C. Soize. *Mathematics of random phenomena*. D. Reidel Publishing Company, Dordrecht, 1986.
- [17] P. Whittle. *Hypothesis testing in time series, PhD thesis*. University of Uppsala, 1951.
- [18] P. Whittle. *Prediction and Regulation by Linear Least-Square Methods*. University of Minnesota Press, 1983.
- [19] G. Box and G.M. Jenkinsl. *Time Series Analysis: Forecasting and Control*. Holden-Day, San Francisco, 1970.
- [20] Grigoriu M. Parametric translation models for stationary non-gaussian processes and fields. *Journal of Sound and Vibration*, 303:3–5, 2007.
- [21] O.P. Le Maître and O.M. Knio. *Spectral Methods for Uncertainty Quantification*. Springer, 2010.
- [22] R. Ghanem and P. D. Spanos. *Stochastic Finite Elements: A Spectral Approach, rev. ed.* Dover Publications, New York, 2003.
- [23] C. Allery, A. Hambouni, D. Ryckelynck, and N. Verdon. A priori reduction method for solving the two-dimensional burgers' equations. *Applied Mathematics and Computation*, 217:6671–6679, 2011.
- [24] J.A. Atwell and B.B. King. Proper orthogonal decomposition for reduced basis feedback controllers for parabolic equations. *Math. Comput. Modell.*, 33 (1-3):1–19, 2001.
- [25] G. Berkooz, P. Holmes, and J.L. Lumley. The proper orthogonal decomposition in the analysis of turbulent flows. *Annu. Rev. Fluid Mech.*, 25:539–575, 1993.
- [26] G. P. Brooks and J. M. Powers. A karhunen-loève least-squares technique for optimization of geometry of a blunt body in supersonic flow. *Journal of Computational Physics*, 195:387–412, 2004.
- [27] E.A. Christensen, M. Brons, and J.M. Sorensen. Evaluation of proper orthogonal decomposition-based decomposition techniques applied to parameter dependent nonturbulent flows. *SIAM J. Sci. Comput*, 21 (4):1419–1434, 2000.
- [28] S.P. Huang, S.T. Quek, and K.K. Phoon. Convergence study of the truncated karhunen-loève expansion for simulation of stochastic processes. *Int J Num Meth Engng*, 52 (9):1029–43, 2001.
- [29] K. Kunisch and S. Volkwein. Galerkin proper orthogonal decomposition methods for parabolic problems. *Numer. Math.*, 90 (1):117–148, 2001.
- [30] L. Li, K. Phoon, and S. Quek. Comparison between karhunen-loève expansion and translation-based simulation of non-gaussian processes. *Computers and Structures*, 85:264–76, 2007.
- [31] X. Ma and N. Zabaras. Kernel principal component analysis for stochastic input model generation. *Comptes rendus de l'Académie des sciences de Paris*, 220, 1945.

- [32] H.G. Matthies. Stochastic finite elements: Computational approaches to stochastic partial differential equations. *Zamm-Zeitschrift für Angewandte Mathematik und Mechanik*, 88 (11):849–873, 2008.
- [33] A. Nouy and O.P. Le Maître. Generalized spectral decomposition method for stochastic non-linear problems. *J. Comput. Phys.*, 228 (1):202–235, 2009.
- [34] G. Perrin, C. Soize, D. Duhamel, and C. Funfschilling. Identification of polynomial chaos representations in high dimension from a set of realizations. *SIAM J. Sci. Comput.*, 34(6):2917–2945, 2012.
- [35] K.K. Phoon, S.P. Huang, and S.T. Quek. Implementation of karhunen-loeve expansion for simulation using a wavelet-galerkin scheme. *Probabilistic Engineering Mechanics*, 17:293–303, 2002.
- [36] K.K. Phoon, S.P. Huang, and S.T. Quek. Simulation of strongly non-gaussian processes using karhunen-loeve expansion. *Probabilistic Engineering Mechanics*, 20:188–198, 2005.
- [37] C. Schwab and R. A. Todor. Karhunen-loeve approximation of random fields by generalized fast multipole methods. *Journal of Computational Physics*, 217:100–122, 2006.
- [38] C. Soize. Identification of high-dimension polynomial chaos expansions with random coefficients for non-gaussian tensor-valued random fields using partial and limited experimental data. *Computer Methods in Applied Mechanics and Engineering*, 199:2150–2164, 2010.
- [39] P.D. Spanos and B.A. Zeldin. Galerkin sampling method for stochastic mechanics problems. *Journal of Engineering Mechanics*, 120 (5):1091–1106, 1994.
- [40] P.D. Spanos, M. Beer, and J. Red-Horse. Karhunen -loève expansion of stochastic processes with a modified exponential covariance kernel. *Journal of Engineering Mechanics*, 133 (7):773–779, 2007.
- [41] B. Wen and N. Zabaras. A multiscale approach for model reduction of random microstructures. *Computational Materials Science*, 63:269–285, 2012.
- [42] M.M.R. Williams. The eigenfunctions of the karhunen-loeve integral equation for a spherical system. *Probabilistic Engineering Mechanics*, 26:202–207, 2011.
- [43] S. Q. Wu and S. S. Law. Statistical moving load identification including uncertainty. *Probabilistic Engineering Mechanics*, 29:70–78, 2012.
- [44] P.C. Hansen. Numerical tools for analysis and solution of fredholm integral-equations of the 1st kind. *Inverse problems*, 8 (6):849–872, 1992.
- [45] J. Weese. A reliable and fast method for the solution of fredholm integral-equations of the 1st kind based on tikhonov regularization. *Computer physics communications*, 69:99–111, 1992.
- [46] A. Nataf. Détermination des distributions de probabilité dont les marges sont données. *Comptes Rendus de l’Académie des Sciences*, 225:42–43, 1986.
- [47] M. Rosenblatt. Remarks on a multivariate transformation. *Annals of Mathematical Statistics*, 23:470–472, 1952.
- [48] N. Wiener. The homogeneous chaos. *American Journal of Mathematics*, 60:897–936, 1938.

- [49] R. Ghanem and P.D. Spanos. Polynomial chaos in stochastic finite elements. *Journal of Applied Mechanics*, Transactions of the ASME 57:197–202, 1990.
- [50] M. Arnst, R. Ghanem, and C. Soize. Identification of bayesian posteriors for coefficients of chaos expansions. *Journal of Computational Physics*, 229 (9):3134–3154, 2010.
- [51] S. Das, R. Ghanem, and S. Finette. Polynomial chaos representation of spatio-temporal random field from experimental measurements. *J. Comput. Phys.*, 228:8726–8751, 2009.
- [52] B.J. Debuschere, H.N. Najm, P.P. Pébay, O. M. Knio, R. G. Ghanem, and O. P. Le Maître. Numerical challenges in the use of polynomial chaos representations for stochastic processes. *SIAM J. Sci. Comput.*, 26:698–719, 2004.
- [53] C. Desceliers, R. Ghanem, and C. Soize. Maximum likelihood estimation of stochastic chaos representations from experimental data. *Internat. J. Numer. Methods Engrg.*, 66:978–1001, 2006.
- [54] C. Desceliers, C. Soize, and R. Ghanem. Identification of chaos representations of elastic properties of random media using experimental vibration tests. *Comput. Mech.*, 39:831–838, 2007.
- [55] R.G. Ghanem and A. Doostan. On the construction and analysis of stochastic models: Characterization and propagation of the errors associated with limited data. *J. Comput. Phys.*, 217:63–81, 2006.
- [56] R. Ghanem, S. Masri, M. Pellissetti, and R. Wolfe. Identification and prediction of stochastic dynamical systems in a polynomial chaos basis. *Comput. Methods Appl. Mech. Engrg.*, 194:1641–1654, 2005.
- [57] R. Ghanem and J. Red-Horse. Propagation of uncertainty in complex physical systems using a stochastic finite element approach. *Phys. D*, 133:137–144, 1999.
- [58] D. Ghosh and R. Ghanem. Stochastic convergence acceleration through basis enrichment of polynomial chaos expansions. *Internat. J. Numer. Methods Engrg.*, 73:162–184, 2008.
- [59] O. M. Knio and O. P. Le Maître. Uncertainty propagation in cfd using polynomial chaos decomposition. *Fluid Dynam. Res.*, 38:616–640, 2006.
- [60] D. Lucor, C. H. Su, and G. E. Karniadakis. Generalized polynomial chaos and random oscillators. *Internat. J. Numer. Methods Engrg.*, 60:571–596, 2004.
- [61] Y. M. Marzouk, H. N. Najm, and L. A. Rahn. spectral methods for efficient bayesian solution of inverse problems. *J. Comput. Phys.*, 224:560–586, 2007.
- [62] Y. M. Marzouk and H. N. Najm. Dimensionality reduction and polynomial chaos acceleration of bayesian inference in inverse problems. *J. Comput. Phys.*, 228:1862–1902, 2009.
- [63] A. Nouy, A. Clement, F. Schoefs, and N. Moes. An extended stochastic finite element method for solving stochastic partial differential equations on random domains. *Methods Appl. Mech. Engrg.*, 197:4663–4682, 2008.
- [64] A. Nouy. Identification of multi-modal random variables through mixtures of polynomial chaos expansions. *Comptes Rendus Mécanique*, 338 (12):698–703, 2010.

- [65] J. R. Red-Horse and A. S. Benjamin. A probabilistic approach to uncertainty quantification with limited information. *Reliability Engrg. System Safety*, 85:183–190, 2004.
- [66] S. Sakamoto and R. Ghanem. Polynomial chaos decomposition for the simulation of non-gaussian non stationary stochastic processes. *J. Engrg. Mechanics*, 128:190–201, 2002.
- [67] G. I. Schueller. On the treatment of uncertainties in structural mechanics and analysis. *Computational and Structures*, 85:235–243, 2007.
- [68] C. Soize. Generalized probabilistic approach of uncertainties in computational dynamics using random matrices and polynomial chaos decompositions. *Internat. J. Numer. Methods Engrg.*, 81:939–970, 2010.
- [69] C. Soize and R. Ghanem. Physical systems with random uncertainties: Chaos representations with arbitrary probability measure. *SIAM J. Sci. Comput.*, 26, 2004.
- [70] G. Stefanou, A. Nouy, and A. Clement. Identification of random shapes from images through polynomial chaos expansion of random level set functions. *Internat. J. Numer. Methods Engrg.*, 79:127–155, 2009.
- [71] S. Sarkar, S. Gupta, and I. Rychlik. Wiener chaos expansions for estimating rain-flow fatigue damage in randomly vibrating structures with uncertain parameters. *Probabilistic Engineering Mechanics*, 26:387–398, 2011.
- [72] D. Xiu and G. E. Karniadakis. The wiener-askey polynomial chaos for stochastic differential equations. *SIAM Journal on Scientific Computing*, 24, No. 2:619–644, 2002.
- [73] R. G. Miller. The jackknife - a review. *Biometrika*, 61:1–15, 1974.
- [74] N.T. Quan. The prediction sum of squares as a general measure for regression diagnostics. *Journal of Business and Economic Statistics*, 6 (4):501–504, 1988.
- [75] G. Blatman and B. Sudret. Adaptive sparse polynomial chaos expansion based on least angle regression. *J. Comput. Phys.*, 230 pages=2345–2367,, 2011.
- [76] C. Soize and C. Desceliers. Computational aspects for constructing realizations of polynomial chaos in high dimension. *SIAM Journal on Scientific Computing*, 32(5):2820–2831, 2010.
- [77] A. Edelman, T.A. Arias, and S.T. Smith. The geometry of algorithms with orthogonality constraints. *SIAM J. Matrix Anal. Appl.*, 20 (2):303–353, 1998.
- [78] H. G. Matthies and C. Bucher. Finite elements for stochastic media problems. *Comput. Methods Appl. Mech. Engrg.*, 168:3–17, 1999.
- [79] H. G. Matthies and A. Keese. Galerkin methods for linear and nonlinear elliptic stochastic partial differential equations. *Comput. Methods Appl. Mech. Engrg.*, 194:1295–1331, 2005.
- [80] S. Sakamoto and R. Ghanem. Simulation of multi-dimensional non-gaussian non-stationary random fields. *Probabilistic Engineering Mechanics*, 17:167–176, 2002.
- [81] D.X. Zhang and Z.M. Lu. An efficient, high-order perturbation approach for flow in random porous media via karhunen-loeve and polynomial expansions. *Journal of Computational Physics*, 194 (2):773–794, 2004.

- [82] J. Evans. Rail vehicle dynamic simulation using vampire. *Vehicle System Dynamics*, Supplement 31:119–140, 1999.
- [83] O. Evans. *Vampire User guide*. AEA Technology, 2006.
- [84] S. Kraft. *Parameter identification for a TGV model*. PhD thesis, Ecole Centrale de Paris, 2012.
- [85] K. J. Bathe and E.L. Wilson. *Numerical Methods in Finite Element Analysis*. Prentice-Hall, 1976.
- [86] C. Soize. A computational inverse method for identification of non-gaussian random fields using the bayesian approach in very high dimension. *Computer Methods in Applied Mechanics and Engineering*, 200 (45-46):3083–3099, 2011.