



Parameter control in the presence of uncertainties

Victor Trappler

► To cite this version:

Victor Trappler. Parameter control in the presence of uncertainties. Modeling and Simulation. Université Grenoble Alpes [2020-..], 2021. English. NNT : 2021GRALM017 . tel-03275015v2

HAL Id: tel-03275015

<https://theses.hal.science/tel-03275015v2>

Submitted on 27 Sep 2021

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

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

THÈSE

Pour obtenir le grade de

DOCTEUR DE L'UNIVERSITÉ GRENOBLE ALPES

Spécialité : Mathématiques Appliquées

Arrêté ministériel : 25 mai 2016

Présentée par

Victor TRAPPLER

Thèse dirigée par **Arthur VIDARD**, Chargé de recherche,
Université Grenoble Alpes
et codirigée par **Elise ARNAUD**, MCF, UGA
et **Laurent DEBREU**, Directeur de recherche, Inria
préparée au sein du **Laboratoire Laboratoire Jean Kuntzmann**
dans l'**École Doctorale Mathématiques, Sciences et**
technologies de l'information, Informatique

Contrôle de paramètre en présence d'incertitudes

Parameter control in the presence of uncertainties

Thèse soutenue publiquement le **11 juin 2021**,
devant le jury composé de :

Monsieur YOUSSEF MARZOUK

PROFESSEUR, Massachusetts Institute of Technology, Rapporteur

Monsieur PIETRO CONGEDO

DIRECTEUR DE RECHERCHE, INRIA CENTRE DE SACLAY-ILE-DE-FRANCE, Rapporteur

Monsieur OLIVIER ROUSTANT

PROFESSEUR DES UNIVERSITÉS, INST NAT SC APPLIQ TOULOUSE, Président

Monsieur REMY BARAILLE

INGENIEUR DOCTEUR, SERVICE HYDRO ET OCEANOGRAPHIQUE MARINE, Examineur



REMERCIEMENTS

Je voudrais tout d’abord remercier mes directeur·trice·s Élise, Arthur et Laurent pour m’avoir accueilli, guidé et soutenu durant toute la durée de la thèse.

Je voudrais aussi remercier les membres du jury, qui ont accepté de se pencher sur mon travail, et qui par leurs remarques m’ont ouvert des nouvelles perspectives, et de nouveaux axes de recherche.

Je voudrais ensuite remercier tous les collègues du LJK et de l’équipe AIRSEA. Que ce soit sur un plan scientifique, administratif ou le plus important humain, j’ai pu trouver des personnes brillantes et formidables qui m’ont aidé et encouragé, et ont contribué à rendre mon passage au laboratoire agréable.

Merci à tous mes amis de Grenoble, de Marlenheim, du lycée et d’ailleurs pour avoir toujours fait l’effort de m’écouter parler de mon travail, mais aussi pour me rappeler qu’il y a pleins d’autres choses à faire. Par crainte d’en oublier, je préfère ne citer personne, mais l’intention est là.

Enfin et surtout, je voudrais remercier ma famille: Juliette, Papa, Maman, Lucie, Papy, Mamie qui m’ont toujours soutenu durant ces 3 4 années riches en émotions. Tout ce travail est pour vous.

* * *

RÉSUMÉ FRANÇAIS

Présentation générale et calibration de modèles

De nombreux phénomènes naturels sont modélisés afin de mieux connaître leurs comportements et de pouvoir les prévoir. Cependant, lors du processus de modélisation, de nombreuses sources d'erreurs sont introduites. Elles proviennent par exemple des paramétrisations qui rendent compte des phénomènes sous-mailles, ou bien de l'ignorance des conditions environnementales réelles dans lesquelles le phénomène est observé.

De manière plus formelle, on peut distinguer grossièrement deux types d'incertitudes dans ces modèles, comme évoqué dans [Walker et al. \(2003\)](#).

- les incertitudes dites *épistémiques*, qui proviennent d'un manque de connaissance sur des caractéristiques du phénomène étudié, mais qui pourraient être réduites
- les incertitudes dites *aléatoires*, qui proviennent directement de la variabilité intrinsèque du phénomène étudié.

Dans le cadre de cette thèse, les incertitudes épistémiques prennent la forme de la méconnaissance de la valeur d'un paramètre $\theta \in \Theta$, que l'on va chercher à calibrer. Un exemple de ce genre de problèmes est l'estimation de la friction dans les modèles d'océan qui va donc nous servir de cas d'application. En effet, la friction de fond est due à la rugosité du plancher océanique, provoquant de la dissipation d'énergie à cause des turbulences engendrées. L'estimation de la friction de fond est un problème à forts enjeux, notamment dans les régions côtières, du fait de son influence sur les courants et de son interaction avec la marée ([Sinha and Pingree, 1997](#); [Boutet, 2015](#)).

Cette estimation peut être traitée dans un cadre d'assimilation de données avec des méthodes variationnelles comme dans [Das and Lardner \(1991, 1992\)](#) sur un cas simplifié, ou dans un cas plus réaliste dans [Boutet \(2015\)](#), avec une méthode de gradient stochastique, permettant de se passer du calcul exact du gradient.

Les incertitudes aléatoires, quant à elles, représentent des conditions environnementales, comme le forçage d'un modèle ou les conditions aux bords. Ces conditions ne

sont pas directement contrôlées par le modèle, donc l'on subit leurs fluctuations, ou leur imprécision. Ces variables environnementales vont être modélisées à l'aide d'une variable aléatoire U , de réalisation notée $u \in \mathbb{U}$.

Comme le modèle que l'on cherche à calibrer vise à représenter la réalité, il est souhaitable que les prédictions du modèle soient le plus proche possible des observations dont on dispose. Cette notion est retranscrite en définissant une fonction J , dite fonction coût ou fonction objectif qui mesure l'écart entre la sortie du modèle et les observations disponibles. Cette fonction prendra donc en entrée le paramètre à estimer θ , que l'on nommera paramètre de contrôle, ainsi que u , le paramètre environnemental:

$$\begin{aligned} J : \Theta \times \mathbb{U} &\rightarrow \mathbb{R}_+ \\ (\theta, u) &\mapsto J(\theta, u) \end{aligned}$$

La définition de la fonction coût dans un problème de calibration sera abordé dans le [Chapitre 1](#), en lien notamment avec l'inférence fréquentiste et Bayésienne.

Notions de robustesse

Ne pas prendre en compte les incertitudes aléatoires dans l'estimation de θ peut amener à compenser de manière artificielle certains aspects physiques dûs à la variable environnementale, et donc amener à un comportement analogue au *sur-apprentissage* (overfitting), ou *optimisation localisée* (terme introduit dans [Huyse et al. \(2002\)](#)): des situations où le paramètre estimé n'est optimal que pour la valeur de u supposée, et pour une autre réalisation de la variable aléatoire sous-jacente, le modèle ainsi calibré donne des prédictions potentiellement aberrantes ([Kuczera et al., 2010](#)).

On cherche donc à définir une valeur de θ , notée $\hat{\theta}$ de manière à ce que $J(\hat{\theta}, u)$ reste *acceptable* lorsque l'on prend en compte la variabilité intrinsèque de u . Comme U est une variable aléatoire, pour un θ donné, la fonction coût peut être elle aussi vue comme une variable aléatoire: $J(\theta, U)$, que l'on va chercher à "minimiser" dans un sens qui reste à définir. Cette problématique porte différents noms, comme "optimisation robuste", où cette robustesse doit être comprise comme l'insensibilité aux variations de U , "optimisation sous incertitudes" (*Optimisation under Uncertainty* ou *OUU*), ou encore "optimisation stochastique", selon la configuration du problème. Une nomenclature prenant en compte les différences, notamment sur les contraintes potentiellement présentes, peut être trouvé dans [Lelièvre et al. \(2016\)](#).

L'objectif de la thèse est donc de proposer différents critères de robustesse, et d'appliquer des méthodes adaptées permettant d'estimer un paramètre en présence d'incertitudes. Cette estimation se réalise dans un premier temps dans des cas simples (fonctions analytiques, problèmes simplifiés de faibles dimensions), puis sur des problèmes plus complexes d'estimation de la friction de fond océanique.

Critères basés sur le regret additif et relatif

Dans le [Chapitre 2](#), nous abordons le problème de calibration en présence d'incertitudes. Un certain nombre des méthodes d'optimisation sous incertitude se basent sur la minimi-

sation des moments de la variable aléatoire $J(\cdot, U)$ comme dans [Lehman et al. \(2004\)](#); [Janusevskis and Le Riche \(2010\)](#), ou les incorporent dans un problème d'optimisation multiobjectif ([Baudoui, 2012](#); [Ribaud, 2018](#)).

Dans le cadre de cette thèse, nous proposons une approche basée sur le regret, qui consiste à comparer les valeurs de la fonction J avec le *minimum conditionnel*, qui est le minimum de la fonction $J(\cdot, u)$, où u est une réalisation de la variable aléatoire U . Le minimum conditionnel est donc défini par

$$J^*(u) = \min_{\theta \in \Theta} J(\theta, u)$$

et le *minimiseur conditionnel* associé est

$$\theta^*(u) = \arg \min_{\theta \in \Theta} J(\theta, u)$$

Ceci nous permet de définir différentes notions de regret: le regret additif $J - J^*$ et étant donné la stricte positivité de J , le regret relatif J/J^* . Ceci nous permet d'introduire une notion d'*acceptabilité*, à entendre dans le sens d'écart par rapport au minimum conditionnel.

Pour un $u \in \mathbb{U}$ donné, $\theta \in \Theta$ est dit β -acceptable si $J(\theta, u) \leq J^*(u) + \beta$, pour $\beta \geq 0$. La notion de β -acceptabilité est donc associée au regret additif: $J(\theta, u) - J^*(u)$. Similairement, on définit la notion de α -acceptabilité: θ est dit α -acceptable si $J(\theta, u) \leq \alpha J^*(u)$, pour $\alpha > 1$. Dans la suite, sous nous intéresserons plus particulièrement au regret relatif, qui permet de mieux prendre en compte les variations de magnitude de la fonction objectif, mais les définitions suivantes peuvent être adaptée au regret additif.

En prenant en compte le caractère aléatoire de U , nous pouvons donc étudier la probabilité pour un point θ , d'être α -acceptable:

$$\Gamma_\alpha(\theta) = \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)]$$

Cette probabilité peut ensuite être optimisée, pour donner

$$\theta_{\text{RR}, \alpha} = \arg \max_{\theta \in \Theta} \Gamma_\alpha(\theta)$$

L'optimum atteint est donc la probabilité maximale avec laquelle le regret-relatif est borné par α .

Si, au lieu de choisir un seuil α pour la minimisation, nous cherchons plutôt à atteindre une certaine probabilité d'acceptabilité p , nous pouvons définir la fonction quantile du regret relatif comme

$$q_p(\theta) = Q_U \left(\frac{J(\theta, U)}{J^*(U)}; p \right)$$

où $Q_U(\cdot; p)$ est la fonction quantile à l'ordre p de la variable aléatoire en argument. $q_p(\theta)$ représente donc la valeur qui borne le regret au point θ avec une probabilité donnée p . Ce quantile peut aussi être minimisé, donnant

$$\theta_{\text{RR}, \alpha_p} = \arg \min_{\theta \in \Theta} q_p(\theta)$$

et le minimum atteint est par conséquent α_p , qui vérifie $\Gamma_{\alpha_p}(\theta_{RR, \alpha_p}) = p$.

D'après ces deux formulations nous pouvons donc soit chercher à maximiser la probabilité Γ_α pour $\alpha > 1$ bien choisi, soit chercher à minimiser le quantile q_p , au niveau de confiance p .

Ces critères dépendent donc d'un paramètre additionnel, α , ou p selon la formulation choisie, qui va permettre d'ajuster le caractère *conservatif* de la solution. En effet, choisir une grande valeur de α (ou p très proche de 1) permet de se prévenir des hautes déviations de la fonction objectif avec un grande probabilité. Si à l'inverse, α est choisi plus faible, on favorisera les solutions qui donnent des valeurs de la fonction objectif proches du minimum atteignable, mais potentiellement avec une probabilité plus faible. Ce travail a mené à la publication de [Trappler et al. \(2020\)](#).

Optimisation robuste et processus Gaussiens

D'un point de vue pratique, ces notions de minimiseur conditionnel et de minimum conditionnel peuvent s'avérer difficiles et coûteuses à calculer, car nécessitant une procédure d'optimisation. De plus, la connaissance de la fonction objectif doit être suffisante afin de calculer assez précisément les quantités Γ_α et q_p . Dans le [Chapitre 3](#), nous proposons d'utiliser des processus Gaussiens (GP), afin de créer un modèle de substitution, bien moins coûteux à évaluer, permettant de se passer d'une connaissance exhaustive de la fonction J .

Soit Z le GP construit avec un plan d'expérience $\mathcal{X} = \{(\theta_i, u_i), J(\theta_i, u_i)\}_{1 \leq i \leq n}$, comprenant donc n points. Le métamodèle associé à Z et construit d'après $\mathcal{X}$ sera noté $m_Z : \Theta \times \mathbb{U} \rightarrow \mathbb{R}$, et utilisé en lieu et en place de J pour estimer Γ_α ou q_p , dans une approche dite *plug-in*.

Les propriétés des GP nous permettrons aussi d'établir des stratégies d'enrichissement. En effet, des méthodes existantes dites *adaptatives* permettent d'améliorer l'estimation de diverses quantités, comme la probabilité de défaillance ([Razaaly, 2019](#); [Moustapha et al., 2016](#); [Bect et al., 2012](#)), ou les minimiseurs et minimums conditionnels dans [Ginsbourger et al. \(2014\)](#). Ces méthodes, parfois appelées méthodes SUR (*Stepwise Uncertainty Reduction*, réduction d'incertitude séquentielle) sont basées sur la définition d'un critère κ qui va ensuite être optimisé, et dont le maximiseur va ensuite être évalué par la fonction (supposée coûteuse) J :

$$(\theta_{n+1}, u_{n+1}) = \arg \max_{(\theta, u) \in \Theta \times \mathbb{U}} \kappa((\theta, u); Z)$$

puis le plan d'expérience est enrichi avec ce nouveau point et son évaluation:

$$\mathcal{X}_{n+1} = \mathcal{X}_n \cup \{(\theta_{n+1}, u_{n+1}), J(\theta_{n+1}, u_{n+1})\}$$

et enfin, Z est mis à jour avec le nouveau plan d'expérience $\mathcal{X}_{n+1}$. Ce critère va donc représenter une mesure de l'incertitude sur l'estimation, que l'on va chercher à réduire. Nous allons ainsi proposer plusieurs méthodes permettant d'améliorer l'estimation de Γ_α ou de q_p . Nous proposons aussi des méthodes basées sur l'échantillonnage d'une variable aléatoire à support dans $\Theta \times \mathbb{U}$, dont les échantillons sont des points à forte incertitudes

par rapport à l’objectif final, comme dans [Echard et al. \(2011\)](#); [Razaaly \(2019\)](#). Après une procédure de réduction statistique, comme le partitionnement (ou *clustering* en anglais), on peut donc évaluer et ajouter au plan d’expérience un *lot* de points, et ainsi tirer parti du parallélisme quand une telle architecture est disponible.

Ceci sera fait en définissant notamment $Z^*(u) = Z(\theta^*(u), u)$, et

$$\Delta_{\alpha,\beta}(\theta, u) = Z(\theta, u) - \alpha Z^*(u) - \beta$$

et

$$\Xi(\theta, u) = \log \left(\frac{Z(\theta, u)}{Z^*(u)} \right)$$

qui sont deux processus stochastiques dont les distributions, exactes sinon approchées, pourront être déduites à partir de la loi de Z . Nous pourrions donc établir des stratégies d’enrichissement de plans d’expériences par rapport à ces deux processus.

Application au code de calcul CROCO

Dans le [Chapitre 4](#), nous nous intéressons à la calibration robuste d’un modèle réaliste d’océan, basé sur le code de calcul CROCO. Les incertitudes introduites dans ce cadre portent sur l’amplitude de différentes composantes de marée.

Comme mentionné plus tôt, nous cherchons à estimer un paramètre régissant la friction de fond. Cette étude sera effectuée dans un cadre d’expériences jumelles, c’est-à-dire que les observations seront obtenues grâce au code de calcul.

Nous effectuerons tout d’abord une optimisation de la fonction objectif, sans introduire d’incertitudes. Ensuite, dans le but de réduire la dimension du problème, nous segmenterons le domaine océanique étudié selon le type de sédiments qui se trouve au fond. Afin de quantifier l’influence de chacune des régions délimitées par la classe de sédiments, une analyse de sensibilité globale sera effectuée, afin de calculer les indices de Sobol’ correspondants ([Sobol, 2001](#); [Iooss, 2011](#)). Une étude similaire sera menée pour les différentes composantes du paramètre représentant les incertitudes u . Enfin, une fois la dimension du problème de calibration réduite significativement, nous appliquerons des méthodes présentées au chapitre précédent, afin d’estimer de manière robuste le paramètre de friction de fond dans ce problème académique.

* * *

TABLE OF CONTENTS

Remerciements	i
Résumé Français	iii
List of Figures	xiii
List of Tables	xv
Introduction	1
1 Inverse Problem and calibration	5
1.1 Introduction	7
1.2 Forward, inverse problems and probability theory	7
1.2.1 Model space data space and forward problem	7
1.2.2 Forward problem	8
1.2.3 Inverse Problem	8
1.2.4 Notions of probability theory	9
1.3 Parameter inference	18
1.3.1 From the physical experiment to the model	18
1.3.2 Frequentist inference, MLE	19
1.3.3 Bayesian Inference	21
1.4 Calibration using adjoint-based optimisation	24
1.5 Model selection	26
1.5.1 Likelihood ratio test and relative likelihood	26
1.5.2 Criteria for non-nested model comparison	29
1.6 Parametric model misspecification	30
1.7 Partial conclusion	32
2 Robust estimators in the presence of uncertainties	33
2.1 Defining robustness	35

2.1.1	Classifying the uncertainties	35
2.1.2	Robustness and/or reliability	35
2.1.3	Robustness under parameteric misspecification	36
2.2	Probabilistic inference	37
2.2.1	Frequentist approach	37
2.2.2	Bayesian approach	38
2.3	Variational approach	41
2.3.1	Decision under deterministic uncertainty set	41
2.3.2	Robustness based on the moments of an objective function	43
2.4	Regret-based families of estimators	48
2.4.1	Conditional minimum and minimiser	49
2.4.2	Regret and model selection	51
2.4.3	Relative-regret	54
2.4.4	The choice of the threshold	57
2.5	Partial Conclusion	58
3	Adaptive strategies for calibration using Gaussian Processes	61
3.1	Introduction	63
3.2	Gaussian process regression	63
3.2.1	Random processes	64
3.2.2	Kriging equations	65
3.2.3	Covariance functions	67
3.2.4	Initial design and validation	67
3.3	General enrichment strategies for Gaussian Processes	68
3.3.1	1-step lookahead strategies	69
3.3.2	Batch selection of points: sampling-based methods	70
3.4	Criteria of enrichment	70
3.4.1	Criteria for exploration of the input space	71
3.4.2	Criteria for optimisation of the objective function	72
3.4.3	Contour and volume estimation	75
3.5	Adaptive strategies for robust optimisation using GP	78
3.5.1	PEI for the conditional minimisers	79
3.5.2	Gaussian formulations for the relative and additive regret families of estimators	80
3.5.3	GP-based methods for the estimation of Γ_α	83
3.5.4	Optimisation of the quantile of the relative regret	92
3.6	Partial conclusion	96
4	Application to the numerical coastal model CROCO	99
4.1	Introduction	101
4.2	CROCO and bottom friction modelling	101
4.2.1	Parameters and configuration of the model	102
4.2.2	Modelling of the bottom friction	103
4.2.3	Definition of the control and environmental parameters	103
4.3	Deterministic calibration of the bottom friction	106
4.3.1	Twin experiment setup	107
4.3.2	Cost function definition	107

4.3.3	Gradient-descent optimisation	107
4.4	Sensitivity analysis of the objective function	112
4.4.1	Global Sensitivity Analysis: Sobol' indices	112
4.4.2	SA of the objective function for the calibration of CROCO	113
4.5	Robust Calibration of the bottom friction	114
4.5.1	Objective function and global minimum	116
4.5.2	Conditional minimums and conditional minimisers	118
4.5.3	Relative-regret based estimators	121
4.6	Partial conclusion	128
	Conclusion and perspectives	131
A	Appendix	135
A.1	Lognormal approximation of the ratio of normal random variables	135
A.2	Full sediment repartition in the Bay of Biscay and the English Channel	137
A.3	Misspecified deterministic calibration	137
	Bibliography	149

LIST OF FIGURES

1.1	Forward and Inverse problem diagram	9
1.2	Example of cdf and of pdf	12
1.3	Probability Density functions of 1D and 2D Gaussian r.v.	16
1.4	Probability density functions of χ^2 r.v.	17
1.5	Forward and inverse problem using models as defined Definition 1.2.1 . .	18
1.6	Example of overfitting phenomenon	27
1.7	Example of relative likelihood, and associated likelihood interval, for $\gamma = 0.15$	29
1.8	Effect of the misspecification on the minimiser.	31
2.1	Sources of uncertainties and errors in the modelling	36
2.2	Joint likelihood and posterior distribution	39
2.3	Robust optimisation under uncertainty set	43
2.4	Difference between integrated likelihood and mean loss	45
2.5	Conditional mean and standard deviation	46
2.6	Pareto frontier	47
2.7	Influence of the skewness	48
2.8	Density estimation of the minimisers of J	51
2.9	Different acceptable regions corresponding to different $u \in \mathbb{U}$	53
2.10	Regions of β -acceptability	54
2.11	Regions of α -acceptability	57
2.12	Comparison of the regions of acceptability for additive and relative regret	58
3.1	Illustration of GP regression	66
3.2	Common covariance functions for GP	68
3.3	Optimisation criteria for GP	74
3.4	Estimation of $f^{-1}(B)$ using GP	75
3.5	Samples in the margin of uncertainties and centroids	78
3.6	Illustration of enrichment using the PEI criterion	81
3.7	Augmented IMSE for the estimation of Γ_α	85
3.8	Enriching the design according to the criterion of Eq. (3.77)	86

3.9	Probability of coverage and margin of uncertainty	87
3.10	Principle of the two consecutive adjustments for batch selection	88
3.11	Full batch iteration with double adjustment	90
3.12	Error in the estimation of Γ_α using a sampling-based method	90
3.13	GP prediction, final experimental design and margin of uncertainty	91
3.14	Relation between Γ_α and q_p	92
3.15	Error of the estimation when reducing the augmented ISME of Ξ	94
3.16	One iteration of the QeAK-MCS procedure, with $K_q = 3$ and $K_{\mathbb{M}} = 4$	96
3.17	Evolution of the estimation error for $K = K_q K_{\mathbb{M}} = 12$	97
4.1	Bathymetry chart of the domain modelled	102
4.2	Drag coefficient C_d as a function of the height and the roughness	104
4.3	Repartition of the sediments on the ocean floor	105
4.4	Calibration of the bottom friction using gradient-descent with well-specified environmental variables	109
4.5	Final values of the optimisation procedure, based on the sediment type	110
4.6	Optimisation of z_b , misspecified case	111
4.7	SA on the sediments-based regions	114
4.8	SA on the tide components	115
4.9	Representation of the different steps for the robust calibration of the numerical model using relative-regret estimates	116
4.10	Conditional minimum $m_{Z^*}(u)$ estimated using the GP Z	119
4.11	Distribution of the minimisers $\theta^*(U)$	120
4.12	Evolution of the IMSE during the enrichment strategy	123
4.13	Evolution of the maximal probability of acceptability $\max \Gamma_\alpha^{\text{PI}}$	124
4.14	Components of $\hat{\theta}_{\text{RR},\alpha}$ after augmented IMSE reduction	124
4.15	Volume of the margin of uncertainty associated with $\mathbf{q}_2 = \hat{\alpha}_p$	126
4.16	Estimation of the threshold α_p , depending on the level p	127
4.17	Relative-regret based estimates $\hat{\theta}_{\text{RR},\alpha_p}$, depending on the level p	128
A.1	Full sediments repartition	137
A.2	Optimisation on the whole space, $u^b = (0, 0)$	139
A.3	Optimisation on the whole space, $u^b = (0, 0.5)$	140
A.4	Optimisation on the whole space, $u^b = (0, 1)$	141
A.5	Optimisation on the whole space, $u^b = (0.5, 0)$	142
A.6	Optimisation on the whole space, $u^b = (0.5, 0.5) = u^{\text{truth}}$	143
A.7	Optimisation on the whole space, $u^b = (0.5, 1)$	144
A.8	Optimisation on the whole space, $u^b = (1, 0)$	145
A.9	Optimisation on the whole space, $u^b = (1, 0.5)$	146
A.10	Optimisation on the whole space, $u^b = (1, 1)$	147

LIST OF TABLES

2.1	Nomenclature of robustness proposed in Lelièvre et al. (2016)	36
2.2	Objective function, expected loss, additive regret and relative error . . .	55
2.3	Summary of single objective robust estimators	59
3.1	Common stationary covariance functions	67
4.1	Types and sizes of each sediment class	104
4.2	Harmonic constituents used in the configuration	106
4.3	Values of the θ component of the global optimiser, and truth value . . .	117
4.4	Comparison of methods and numerical results for the robust calibration of CROCO	128
A.1	Correspondence between the full sediments repartition classes and the segmentation used in Chapter 4	138

INTRODUCTION

To understand and to be able to forecast natural phenomena is crucial for many applications with high social, environmental and economic stakes. In earth sciences especially, the modelling of the ocean and the atmosphere is important for day to day weather forecasts, hurricanes tracking, or pollutant dispersion for instance.

Those natural phenomena are then modelled mathematically, usually by representing the physical reality with some general equations (Navier-Stokes equations in Computational Fluid Dynamics typically), and by making successive reasonable assumptions, simplifications and discretisations in order to be able to implement appropriate solvers.

Models are then only a partial representation of the reality, which aim at representing complex processes that occur across a large range of scales, scales that interact with each other. By essence, no modelling system would be able to take all of those into account but instead, their effects are incorporated in the modelling by overly simplifying them and by parametrizing them.

In ocean modelling, and especially in coastal regions, a telling example of this is the parametrization of the bottom friction. This phenomenon occurs as the asperities of the ocean bed dissipates energy through turbulences, and thus affects the circulation at the surface. Since this happens at a subgrid level, *i.e.* at a scale usually several order of magnitudes below the scale of the domain after discretization, modelling all those turbulences is completely unfeasible in practice: the knowledge of the ocean bed is too limited for such applications and the computational power required would be unthinkable. Instead, the effect of the bottom friction is accounted for through parametrization, so by introducing a new parameter which is defined at every point of the mesh. This parametrisation, or more precisely, the estimation of this parameter will motivate the work carried in this thesis.

As this modelling is supposed to represent the reality, the prediction should be compared with some data acquired through observations. This comparison usually takes the form of the definition of a misfit function J that measures the error between the

forecast and the observations. This objective function is then minimised with respect to some chosen parameters θ (Das and Lardner, 1991, 1992; Boutet, 2015) in order to get a calibrated model. Those parameters will be called the *control parameters*.

However, the parameters introduced are not the only source of errors in the modelling. For such complex systems, some additional inputs are subject to unrepresented statistical fluctuations (Zanna, 2011), manifesting themselves at the boundary conditions, or in the forcing of the model for instance. Such *intrinsic* uncertainties are often subject to variability, and neglecting this can lead to further errors (McWilliams, 2007), or aberrant predictions (Kuczera et al., 2010). We chose to model this additional source of uncontrollable uncertainty with a random variable U . Because of this, the objective function is then a function which takes two arguments: the parameter to calibrate θ , and some other parameter u , which can be thought as a realisation of the random variable U , that we shall call *environmental* parameter.

Due to the presence of this random source of uncertainty, we wish to calibrate the model, *i.e.* to select a value of the control parameter θ , in a manner that guarantees that the model represents accurately enough the reality, despite the random nature of the value of the environmental parameter. In other words, as the objective function measures in some sense the quality of the calibration, we wish that this function exhibits *acceptable* values as often as possible, when θ is fixed. This defines intuitively the underlying notion of *robustness* with respect to the variability of the uncertain variable.

In this thesis, we will study an aspect of the calibration of a numerical model under uncertainty, by discussing the notions of robustness, and by proposing a new family of criteria. Specific methods will also be introduced, and applied to the calibration of a numerical model of the ocean. The thesis is organised as follows:

- in [Chapter 1](#), we introduce notions of statistics and probabilities that we will use to define the calibration problem. More specifically, the statistical and Bayesian inference problems will be broached, as well as some aspects of nested model selection using the likelihood ratio test. The link between probabilistic formulations of the inference problem, and variational approach based on the optimisation of a specified objective function will be emphasized.
- in [Chapter 2](#), we are going to discuss some of the notions of robustness that can be found in the literature, either from a probabilistic inference aspect, or through the prism of optimisation under uncertainties. Most existing methods rely on the optimisation of the moments of $\theta \mapsto J(\theta, U)$ (in [Lehman et al. \(2004\)](#); [Janusevskis and Le Riche \(2010\)](#)), while other methods are based on multiobjective problems, such as in [Baudoui \(2012\)](#); [Ribaud \(2018\)](#).

We propose a new family of criteria, which are based on the comparison between the objective function at a couple (θ, u) and its optimal value given the same environmental variable. This notion of regret, either relative or additive, is then optimised in the sense of minimising the probability of exceeding a specified threshold, or to minimise one of its quantile, in order to control with high enough probability its variations. This work has led to the publication of an article ([Trappler et al., 2020](#)).

- The family of criteria introduced in Chapter 2 can be quite expensive to evaluate, that is why in Chapter 3, we will discuss the use of metamodels, Gaussian Processes especially, in order to choose iteratively the new points to evaluate. The process of selection, called *SUR* method (Bect et al., 2012) (Stepwise Uncertainty Reduction) will depend on the type of robust estimation we wish to carry. Different methods will be proposed in order to estimate members of the regret-based family of robust estimators. These approaches differ by the measure of uncertainty on the function we wish to optimise. We will also introduce methods in order to select a batch of points, in order to take advantage of parallelism when available.
- Finally, in Chapter 4, we will study the calibration of a regional coastal model based on CROCO¹. This study will focus on the estimation of the bottom friction parameter, where some uncertainties are introduced in the form of small perturbations of the amplitude of some tidal constituents, that force the model. This problem will be treated using twin experiments, where the observations will in fact be generated using the model.

The definition of the problem will require first to segment the input space, and to quantify the influence of each input variable, using *global sensitivity analysis* (Iooss, 2011). Based on this analysis, The input space will be reduced, in order to carry a tractable robust estimation, using some of the methods proposed in Chapter 3.

* * *

¹<https://www.croco-ocean.org/>

CHAPTER 1

INVERSE PROBLEM AND CALIBRATION

Contents

1.1	Introduction	7
1.2	Forward, inverse problems and probability theory	7
1.2.1	Model space data space and forward problem	7
1.2.2	Forward problem	8
1.2.3	Inverse Problem	8
1.2.4	Notions of probability theory	9
1.2.4.a	Probability measure, and random variables	9
1.2.4.b	Real-valued random variables	10
1.2.4.c	Real-valued random vectors	12
1.2.4.d	Bayes' Theorem	13
1.2.4.e	Important examples of real random variables	15
1.3	Parameter inference	18
1.3.1	From the physical experiment to the model	18
1.3.2	Frequentist inference, MLE	19
1.3.3	Bayesian Inference	21
1.3.3.a	Posterior inference	21
1.3.3.b	Bayesian Point estimates	22
1.3.3.c	Choice of a prior distribution	23
1.4	Calibration using adjoint-based optimisation	24
1.5	Model selection	26
1.5.1	Likelihood ratio test and relative likelihood	26
1.5.2	Criteria for non-nested model comparison	29

1.6	Parametric model misspecification	30
1.7	Partial conclusion	32

1.1 Introduction

In this chapter we will first lay the ground for developing the general ideas behind calibration, by introducing the notions of models, and forward and inverse problems in [Section 1.2](#). This implies also a short review of notions of probability theory. Calibration will be defined in [Section 1.3](#) as the optimisation of an objective function: Maximum likelihood estimation in a frequentist setting, or posterior maximisation using Bayes' theorem. In practice, for large-scale applications, the optimisation is performed using gradient-descent, and the computational cost of gradient computation can be overcome by adjoint method, as described in [Section 1.4](#). Finally, we are going to discuss two aspects related to calibration, namely model selection in [Section 1.5](#) and the influence of nuisance parameters and model misspecification in calibration in [Section 1.6](#).

1.2 Forward, inverse problems and probability theory

Running a simulation using numerical tools is useful to grasp a better understanding of the physical phenomena, or to forecast them. On the other hand, when observing and comparing the measurements and the output of the numerical simulation, we can quantify the mismatch between the two and tune some parameters involved in the computations. Indeed, these parameters represent different physical quantities or processes that are for example unresolved at the model's scale (such as the modelling of the turbulences), or ill-known. A proper estimation of these parameters has to be performed in order to guarantee a meaningful output when evaluating the model.

Model calibration or parameter estimation has been widely treated in the literature, either from a statistical and probabilistic point of view using likelihood-based methods and Bayesian inference, or from a *variational* point of view by defining proper objective functions. To match those two approaches, we will first review the problem from a probabilistic point of view, in order to define properly some appropriate objective functions and introduce tools from optimal control theory to optimise them.

1.2.1 Model space data space and forward problem

In order to describe accurately a physical system, we have to define the notion of models, and will be following [Tarantola \(2005\)](#) approach to define inverse problems. A model represents the link between some parameters and some observable quantities. An example is a model that takes the form of a system of ODEs or PDEs, maybe discretized, while the parameters are the initial conditions and the output is one or several time series, describing the time evolution of a quantity at one or several spatial points. An important point is that a model is not only the *forward operator*, but must also include the parameter space.

Definition 1.2.1 – Model: A model $\mathfrak{M}$ is defined as a pair composed of a *forward operator* $\mathcal{M}$, and a *parameter space* Θ

$$\mathfrak{M} = (\mathcal{M}, \Theta) \tag{1.1}$$

The forward operator is the mathematical representation of the physical system, while the parameter space is chosen here to be a subset of a finite dimensional space, so usually, Θ will be a subset of $\mathbb{R}^p$.

As we will usually choose Θ as a subset of $\mathbb{R}^p$, for $p \geq 1$, we can define the dimensionality of the model, based on the number of *degrees of freedom* available for the parameters.

Remark 1.2.2: The dimension of a model $\mathfrak{M} = (\mathcal{M}, \Theta)$ is the number of parameters not reduced to a singleton, so if $\Theta \subset \mathbb{R}^p$, the dimension of $\mathfrak{M}$ is $d \leq p$. The dimension of a model $\mathfrak{M}$ is sometimes called the degrees of freedom of $\mathfrak{M}$.

Example 1.2.3: A model with parameter space $\Theta = \mathbb{R}^2 \times [0, 1]$ has dimension 3, while $\Theta = \mathbb{R}^2 \times \{1\}$ has dimension 2.

Now that we have introduced the forward operator and the parameter space, we will focus on the output of the model. Ideally, the data space $\mathbb{Y}$ consists in all the physically acceptable results of the physical experiment. Then, the forward operator $\mathcal{M}$ maps the parameter space $\Theta \subset \mathbb{R}^p$ to the data space $\mathbb{Y}$, as one can expect that all models provide physically acceptable outputs.

1.2.2 Forward problem

Given a model $(\mathcal{M}, \Theta)$, the *forward problem* consists in applying the forward operator to a given $\theta \in \Theta$, in order to get the *model prediction*. The forward problem is then to obtain information on the result of the experiment based on the parameters we chose as input, so deriving a satisfying forward operator $\mathcal{M}$.

$$\begin{aligned} \mathcal{M} : \Theta &\longrightarrow \mathbb{Y} \\ \theta &\longmapsto \mathcal{M}(\theta) \end{aligned} \tag{1.2}$$

As said earlier, the forward operator can be a set of ODEs or PDEs, discretized or not. The forward problem is then the attempt to link the causes, so the parameters, to the consequences, *i.e.* the output in the data space.

1.2.3 Inverse Problem

The inverse problem is the counterpart of the forward problem, and consists in trying to gather more information on the parameters, based on: first the result of the experiment (the observation of the physical process), and secondly on the knowledge of the forward operator, as illustrated [Fig. 1.1](#).

This is done by directly comparing the output of the forward operator, and trying to reduce the mismatch between the observed data and the model prediction.

However, a purely deterministic approach for the inverse problem is doomed to under-perform: as most physical processes are not perfectly known, some uncertainties remain in the whole modelling process. Those uncertainties are ubiquitous: the observations available may be corrupted by a random noise coming from the measurement devices and

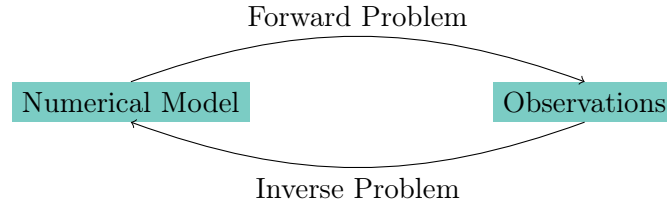


Figure 1.1 – Forward and Inverse problem diagram

the model may not represent perfectly the reality, thus introducing a systematic bias for instance. Taking into account those uncertainties is crucial to solve the inverse problem.

In that perspective we are going to introduce briefly the usual probabilistic framework, along with common notations that we will use throughout this manuscript. Those notions are well established in the scientific literature, and one can read [Billingsley \(2008\)](#) for a more thorough description.

1.2.4 Notions of probability theory

1.2.4.a Probability measure, and random variables

We are first going through some usual notions of probability theory.

Definition 1.2.4 – Event probability and conditioning: Let us consider the usual probabilistic space $(\Omega, \mathcal{F}, \mathbb{P})$. We call an outcome of a random experiment ω an element of the sample space Ω , and an event A is an element of the σ -algebra $\mathcal{F}$ (σ -algebra on the set Ω). The probability of an event $A \in \mathcal{F}$ is defined as the Lebesgue integral

$$\mathbb{P}[A] = \int_A d\mathbb{P}(\omega) = \mathbb{P}[\{\omega; \omega \in A\}] \quad (1.3)$$

Observing an event $B \in \mathcal{F}$ can bring information upon another event $A \in \mathcal{F}$. In that sense, we introduce the conditional probability of A given B . Let $A, B \in \mathcal{F}$. The event A given B is written $A | B$ and its probability is

$$\mathbb{P}[A | B] = \frac{\mathbb{P}[A \cap B]}{\mathbb{P}[B]} \quad (1.4)$$

Formally, an event can be seen as an outcome of some uncertain experiment, and its probability is “how likely” this event will happen.

Let us now introduce a measurable state (or sample) space S , that is the set of all possible outcomes we can observe (and upon which we can assign a probability).

Definition 1.2.5 – Random Variable, Expectation: A random variable (abbreviated as r.v.) X is a measurable function from $\Omega \rightarrow S$. A random variable will usually be written with an upper case letter. A realisation or observation x of the r.v. X is the actual image of $\omega \in \Omega$ under X : $x = X(\omega)$. If S is countable, the

random variable is said to be *discrete*. When $S \subseteq \mathbb{R}^p$ for $p \geq 1$, X is sometimes called a random vector

The expectation of a r.v. $X : \Omega \rightarrow S$ is defined as

$$\mathbb{E}[X] = \int_{\Omega} X(\omega) d\mathbb{P}(\omega) \quad (1.5)$$

Using the [Definition 1.2.5](#), the probability of an event A can be seen as the expectation of the indicator function of A :

$$\begin{aligned} \mathbb{1}_A : \Omega &\longrightarrow \{0, 1\} \\ \omega &\longmapsto \begin{cases} 1 & \text{if } \omega \in A \\ 0 & \text{if } \omega \notin A \end{cases} \end{aligned} \quad (1.6)$$

and it follows that

$$\mathbb{E}[\mathbb{1}_A] = \int_{\Omega} \mathbb{1}_A d\mathbb{P}(\omega) = \int_A d\mathbb{P}(\omega) = \mathbb{P}[A] \quad (1.7)$$

As we defined the notion of a r.v. in [Definition 1.2.5](#) as a measurable function from $\Omega \rightarrow S$, we can now focus on the measurable sets through X , by using in a sense the change of variable $x = X(\omega)$.

Definition 1.2.6 – Image (Pushforward) measure: Let $X : \Omega \rightarrow S$ be a random variable, and $A \subseteq S$. The image measure (also called pushforward measure) of $\mathbb{P}$ through X is denoted by $\mathbb{P}_X = \mathbb{P} \circ X^{-1}$. This notation can differ slightly depending on the community, so one can find also $\mathbb{P}_X = \mathbb{P} \circ X^{-1} = X_{\#}\mathbb{P}$, the latter notation being used in transport theory. The probability, for the r.v. X to be in A is equal to

$$\mathbb{P}[X \in A] = \mathbb{P}_X[A] = \int_A d\mathbb{P}_X(\omega) = \int_{X^{-1}(A)} d\mathbb{P}(\omega) = \mathbb{P}[X^{-1}(A)] = \mathbb{P}[\{\omega; X(\omega) \in A\}] \quad (1.8)$$

Similarly, for any measurable function h , the expectation taken with respect to a specific random variable X is

$$\mathbb{E}_X[h(X)] = \int_{\Omega} h(X(\omega)) d\mathbb{P}_X(\omega) \quad (1.9)$$

In most of this thesis, the sample space will be $S \subseteq \mathbb{R}^p$ for $p \geq 1$, so we are going to introduce useful tools and notations to characterize these particular real random variables.

1.2.4.b Real-valued random variables

We are now going to focus on real-valued random variables, so measurable function from Ω to the sample space $S = \mathbb{R}$.

Definition 1.2.7 – Distribution of a real-valued r.v.: The distribution of a r.v. can be characterized by a few functions:

- The *cumulative distribution function* (further abbreviated as cdf) of a real-valued r.v. X is defined as:

$$F_X(x) = \mathbb{P}[X \leq x] = \mathbb{P}_X[]-\infty, x] \quad (1.10)$$

and $\lim_{-\infty} F_X = 0$ and $\lim_{+\infty} F_X = 1$ If the cdf of a random variable is continuous, the r.v. is said to be *continuous* as well.

- The *quantile function* Q_X is the generalized inverse function of the cdf:

$$Q_X(p) = \inf\{q : F_X(q) \geq p\} \quad (1.11)$$

- If there exists a function $f : S \rightarrow \mathbb{R}^+$ such that for all measurable sets A

$$\mathbb{P}[X \in A] = \int_A d\mathbb{P}_X(\omega) = \int_A f(x) dx \quad (1.12)$$

then f is called the *probability density function* (abbreviated pdf), or *density* of X and is denoted p_X . As $\mathbb{P}[X \in S] = 1$, it follows trivially that $\int_S f(x) dx = 1$. One can verify that if F_X is derivable, then its derivative is the density of the r.v. :

$$\frac{dF_X}{dx}(x) = p_X(x) \quad (1.13)$$

Probability density functions are useful tools to characterize random variables, so the assumption of derivability of F_X is sometimes relaxed for ease of notation.

Remark 1.2.8: When restricting this search to “classical” functions, p_X may not exist. However, allowing generalized functions such as the *dirac delta function*, provides a way to consider simultaneously all types of real-valued random variables (continuous, discrete, and mixture of both). Dirac’s delta function can (in)formally be defined as

$$\delta_{x_0}(x) = \begin{cases} +\infty & \text{if } x = x_0 \\ 0 & \text{elsewhere} \end{cases} \quad \text{and} \quad \int_S \delta_{x_0}(x) dx = 1 \quad (1.14)$$

Example 1.2.9: Let us consider the random variable X that takes the value 1 with probability 0.5, and follows a uniform distribution with probability 0.5 over $[2, 4]$. Its cdf can be expressed as

$$F_X(x) = \begin{cases} 0 & \text{if } x < 1 \\ 0.5 & \text{if } 1 \leq x < 2 \\ 0.5 + \frac{x-2}{8} & \text{if } 2 \leq x < 4 \\ 1 & \text{if } 4 \leq x \end{cases} \quad (1.15)$$

and its pdf (as a generalized function)

$$p_X(x) = \frac{1}{2}\delta_1(x) + \frac{1}{4}\mathbb{1}_{\{2 \leq x < 4\}}(x) \quad (1.16)$$

The pdf and cdf are shown Fig. 1.2.

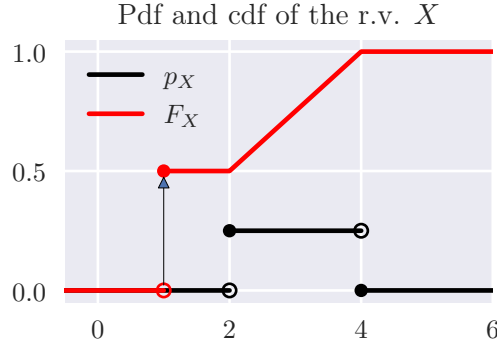


Figure 1.2 – Cdf and Pdf of X defined in Example 1.2.9. The arrow indicates Dirac's delta function

Definition 1.2.10 – Moments of a r.v. and L^s spaces: Let X be a random variable. The moment of order s is defined as $\mathbb{E}[X^s]$, and the centered moment of order s is defined as

$$\mathbb{E}[(X - \mathbb{E}[X])^s] = \int (X(\omega) - \mathbb{E}[X])^s d\mathbb{P}(\omega) = \int (x - \mathbb{E}[X])^s \cdot p_X(x) dx \quad (1.17)$$

To ensure that those moments exist, let us define $L^s(\mathbb{P})$ as the space of random variables X such that $\mathbb{E}[|X|^s] < +\infty$. If $X \in L^2(\mathbb{P})$, the centered moment of order 2 is called the variance:

$$\mathbb{E}[(X - \mathbb{E}[X])^2] = \text{Var}[X] \geq 0 \quad (1.18)$$

These definitions above hold for real-valued random variables, so 1D r.v., but can be extended for random vectors.

1.2.4.c Real-valued random vectors

Most of the definitions for a random variable extend component-wise to random vectors:

Definition 1.2.11 – Joint, marginal and conditional densities: Let $X = (X_1, \dots, X_p)$ be a random vector from $\Omega \rightarrow S \subseteq \mathbb{R}^p$. The expected value of a random vector is the expectation of the components

$$\mathbb{E}[X] = (\mathbb{E}[X_1], \dots, \mathbb{E}[X_p]) \quad (1.19)$$

The cdf of X at the point $x = (x_1, \dots, x_p)$ is

$$\begin{aligned} F_X(x) &= F_{X_1, \dots, X_p}(x_1, \dots, x_p) = \mathbb{P}[X_1 \leq x_1, \dots, X_p \leq x_p] \\ &= \mathbb{P}\left[\bigcap_{i=1}^p \{\omega; X_i(\omega) \leq x_i\}\right] \end{aligned} \quad (1.20)$$

Similarly as in the real-valued case, we can define the pdf of the random vector, or *joint pdf* by derivating with respect to the variables:

$$p_X(x) = p_{X_1, \dots, X_p}(x_1, \dots, x_p) = \frac{\partial^p F_X}{\partial x_1 \cdots \partial x_p}(x) \quad (1.21)$$

and it still integrates to 1: $\int_S p_{X_1, \dots, X_p}(x_1, \dots, x_p) d(x_1, \dots, x_p) = 1$

For two random vectors X and Y , the (cross-)covariance matrix of X and Y is defined as

$$\text{Cov}[X, Y] = \mathbb{E}[(X - \mathbb{E}[X])(Y - \mathbb{E}[Y])^T] = \mathbb{E}[XY^T] - \mathbb{E}[X]\mathbb{E}[Y]^T \quad (1.22)$$

and based on this definition, we can extend the notion of variance to vectors. The covariance matrix $\Sigma \in \mathbb{R}^{p \times p}$ of X , is defined as

$$\Sigma = \text{Cov}(X) = \text{Cov}[X, X] = \mathbb{E}[(X - \mathbb{E}[X])(X - \mathbb{E}[X])^T] = \mathbb{E}[XX^T] - \mathbb{E}[X]\mathbb{E}[X]^T \quad (1.23)$$

We can now define the *marginal densities*. For notation clarity, we are going to set $X = (Y, Z)$: the marginal densities of Y and Z are

$$p_Y(y) = \int_{\mathbb{R}} p_{Y,Z}(y, z) dz \quad \text{and} \quad p_Z(z) = \int_{\mathbb{R}} p_{Y,Z}(y, z) dy \quad (1.24)$$

The random variable Y given Z , denoted by $Y | Z$ has the conditional density

$$p_{Y|Z}(y | z) = \frac{p_{Y,Z}(y, z)}{p_Z(z)} \quad (1.25)$$

allowing us to rewrite the marginals as

$$p_Y(y) = \int_{\mathbb{R}} p_{Y|Z}(y | z) p_Z(z) dz = \mathbb{E}_Z[p_{Y|Z}(y | z)] \quad (1.26)$$

$$p_Z(z) = \int_{\mathbb{R}} p_{Z|Y}(z | y) p_Y(y) dy = \mathbb{E}_Y[p_{Z|Y}(z | y)] \quad (1.27)$$

1.2.4.d Bayes' Theorem

The classical Bayes' theorem is directly a consequence of the definition of the conditional probabilities in [Definition 1.2.4](#), and for random variables admitting a density in [Definition 1.2.11](#).

Theorem 1.2.12 – Bayes’ theorem: Let $A, B \in \mathcal{F}$. Bayes’ theorem states that

$$\begin{aligned}\mathbb{P}[A | B] \cdot \mathbb{P}[B] &= \mathbb{P}[B | A] \cdot \mathbb{P}[A] \\ \mathbb{P}[A | B] &= \frac{\mathbb{P}[B | A] \cdot \mathbb{P}[A]}{\mathbb{P}[B]} \text{ if } \mathbb{P}[B] \neq 0\end{aligned}$$

In terms of densities, the formulation is sensibly the same. Let Y and Z be two random variables. The conditional density of Y given Z can be expressed using the conditional density of Z given Y .

$$p_{Y|Z}(y | z) = \frac{p_{Z|Y}(z | y)p_Y(y)}{p_Z(z)} = \frac{p_{Z|Y}(z | y)p_Y(y)}{\int p_{Z,Y}(z, y) dy} \propto p_{Z|Y}(z | y)p_Y(y) \quad (1.28)$$

Bayes’ theorem is central as it links in a simple way conditional densities. In the inverse problem framework, if Y represents the state of information on the parameter space, while Z represents the information on the data space, $Z | Y$ can be seen as the forward problem. Bayes’ theorem allow us to “swap” the conditioning, and get information on $Y | Z$, that can be seen as the inverse problem.

The influence of one (or a set of) random variable(s) over another can be measured with the conditional probabilities. Indeed, if the state of information on a random variable does not change when observing another one, the observed one carries no information on the other. This notion of dependence (and independence) is first defined on events and then extended to random variables

Definition 1.2.13 – Independence: Let $A, B \in \mathcal{F}$. Those two events are said independent if $\mathbb{P}[A \cap B] = \mathbb{P}[A]\mathbb{P}[B]$. Quite similarly, two real-valued random variables Y and Z are said to be independent if $F_{Y,Z}(y, z) = F_Y(y)F_Z(z)$ or equivalently, $p_{Y,Z}(y, z) = p_Y(y)p_Z(z)$. Speaking in terms of conditional probabilities, this can be written as $p_{Y|Z}(y, z) = p_Y(y)$. If Y and Z are independent, $\text{Cov}[Y, Z] = 0$. The converse is false in general.

We discussed so far the different quantities that characterize random variables. Let us consider now two random variables which share the same sample space: $X, X' : \Omega \rightarrow S$. There exists various way to compare those two random variables, usually by quantifying some measure of distance between their pdf when they exist. One of the most used comparison tool for random variables is the Kullback-Leibler divergence.

Definition 1.2.14 – KL-divergence and entropy: The Kullback-Leibler divergence, introduced in [Kullback and Leibler \(1951\)](#) is a measure of dissimilarity between two distributions, based on information-theoretic considerations. Let X, X' be r.v. with the same sample space S , and p_X and $p_{X'}$ their densities, such that $\forall A \in \mathcal{F}$,

$\int_A p_X(x) dx = 0 \implies \int_A p_{X'}(x) dx = 0$. The KL-divergence is defined as

$$D_{\text{KL}}(p_X \| p_{X'}) = \int_S p_X(x) \log \frac{p_X(x)}{p_{X'}(x)} dx \quad (1.29)$$

$$= \mathbb{E}_X [-\log p_{X'}(X)] - \mathbb{E}_X [-\log p_X(X)] \quad (1.30)$$

$$= H[X', X] - H[X] \quad (1.31)$$

$H[X]$ is called the (differential) entropy of the random variable X , and $H[X', X]$ the cross-entropy of X' and X . Using Jensen's inequality, one can show that for all X and X' such that the KL-divergence exists, $D_{\text{KL}}(p_X \| p_{X'}) \geq 0$ with equality iff they have the same distribution, a desirable property when measuring dissimilarity. However, the KL-divergence is not a distance function, as it is not symmetric in general, and it does not verify the triangle inequality.

1.2.4.e Important examples of real random variables

One of the most important and simple distribution is the uniform distribution, translating the idea that the random variable takes a value on a given interval almost surely.

Example 1.2.15 – The uniform distribution: Let X be a r.v. from Ω to $\mathbb{R}$, and $a < b$. X is said to be uniformly distributed on $[a, b]$ if

$$p_X(x) = \mathbb{1}_{[a,b]}(x) \frac{1}{b-a} \quad (1.32)$$

One other well-known distribution is the normal distribution, also called Gaussian distribution, that appears in various situations, but most notably in the central limit theorem.

Example 1.2.16 – The Normal distribution: Let X be a r.v. from Ω to $\mathbb{R}$. If X follow the normal distribution of mean $\mu \in \mathbb{R}$ and variance $\sigma^2 > 0$, we write $X \sim \mathcal{N}(\mu, \sigma^2)$, and its pdf is

$$p_X(x) = \phi(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}\right) \quad (1.33)$$

For the multidimensional case, let X be a r.v. from Ω to $\mathbb{R}^p$, that follows a normal distribution of mean $\mu \in \mathbb{R}^p$ and covariance matrix $\Sigma \in \mathbb{R}^{p \times p}$, where Σ is semi-definite positive. In that case, $X \sim \mathcal{N}(\mu, \Sigma)$ the density of the random vector X can be written as

$$p_X(x) = (2\pi)^{-\frac{p}{2}} |\Sigma|^{-1} \exp\left(-\frac{1}{2} (x-\mu)^T \Sigma^{-1} (x-\mu)\right) \quad (1.34)$$

where $|\Sigma|$ is the determinant of the matrix Σ , and $(\cdot)^T$ is the transposition operator. As the covariance matrix appears through its inverse, another encountered parametrization

is to use the precision matrix Σ^{-1} . Examples of pdf of Gaussian normal distributions are displayed Fig. 1.3.

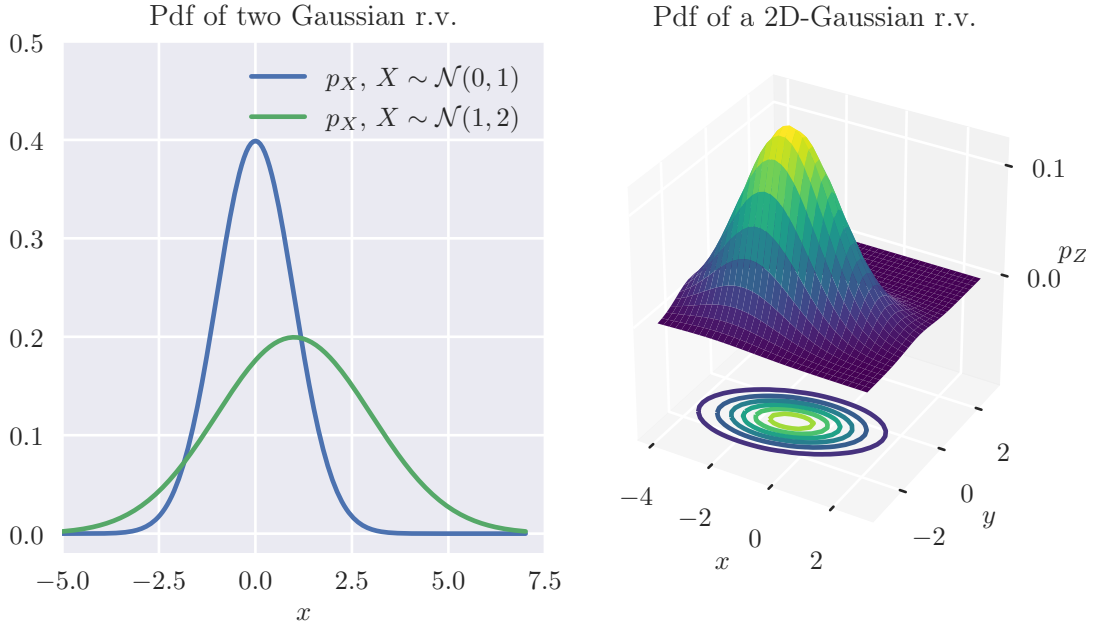


Figure 1.3 – Probability Density functions of 1D Gaussian distributed r.v. (left), and density of Z , a 2D Gaussian r.v. (right)

When adding independent squared samples of the normal distribution, the resulting random variable follows a χ^2 distribution.

Example 1.2.17 – The χ^2 distribution: Let $X_1, X_2, \dots, X_\nu$ be ν independent random variables, such that for $1 \leq i \leq \nu$, $X_i \sim \mathcal{N}(0, 1)$. We define the random variable X as

$$X = \sum_{i=1}^{\nu} X_i^2 \quad (1.35)$$

By definition, the random variable X follows a χ^2 distribution with ν degrees of freedom: $X \sim \chi_\nu^2$. The quantile of order β is written $\chi_\nu^2(\beta)$ and verifies

$$\mathbb{P}[X \leq \chi_\nu^2(\beta)] = \beta \quad (1.36)$$

The pdf of such a r.v. is displayed Fig. 1.4, for different degrees of freedom.

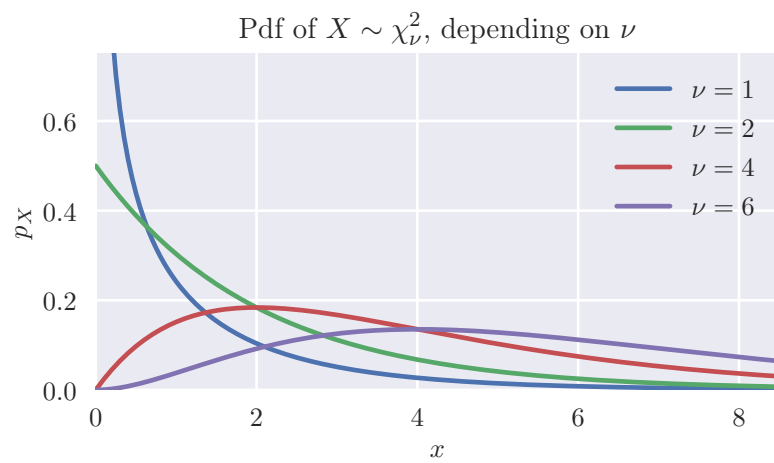


Figure 1.4 – Probability density functions of χ^2_ν random variables, for different degrees of freedom

1.3 Parameter inference

1.3.1 From the physical experiment to the model

We can represent both the reality and the computer simulation as models. The physical system (the reality) that is observed can be represented by a model $\mathfrak{M} = (\mathcal{M}, \Theta_{\text{real}})$, so by a forward operator $\mathcal{M}$, and a parameter space Θ_{real} . Observing the physical system means to get access to $y \in \mathbb{Y}$ that is the image of an *unknown* parameter value $\vartheta \in \Theta_{\text{real}}$ through the forward operator, so $y = \mathcal{M}(\vartheta) \in \mathbb{Y} \subseteq \mathbb{R}^n$.

On the other hand, let us assume that a numerical model of the reality has been constructed, by successive various assumptions, discretizations and simplifications giving $(\mathcal{M}, \Theta)$. The main objective of calibration is to find $\hat{\theta}$ such that the forward operator applied to $\hat{\theta}$: $\mathcal{M}(\hat{\theta})$ represents as accurately as possible the physical system, and thus matches as closely the data $\mathcal{M}(\vartheta) = y$. This is illustrated Fig. 1.5.

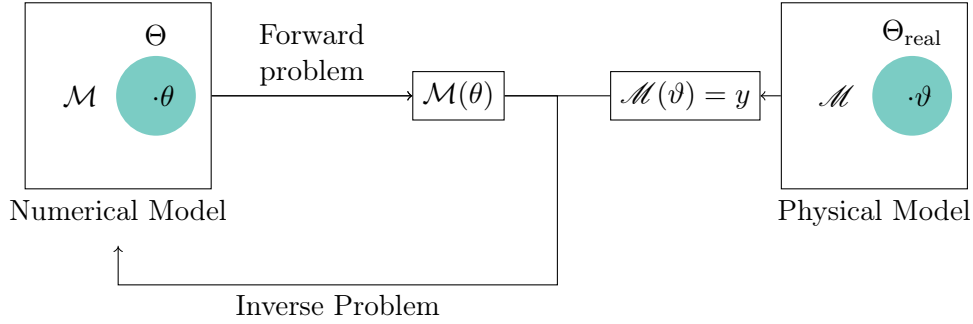


Figure 1.5 – Forward and inverse problem using models as defined Definition 1.2.1

For the sake of simplicity, let us assume that $\vartheta \in \Theta \subseteq \Theta_{\text{real}}$. In Kennedy and O’Hagan (2001); Higdon et al. (2004), the authors rewrite the link between the reality and the model at this value as

$$\mathcal{M}(\vartheta) = \mathcal{M}(\vartheta) + \epsilon(\vartheta) \in \mathbb{Y} \subseteq \mathbb{R}^n \quad (1.37)$$

The difference $\epsilon(\vartheta) = \mathcal{M}(\vartheta) - \mathcal{M}(\vartheta)$ is the error between the physical model and the model, called sometimes the misfit, or the residual error. This error is unknown and encompasses different sources of uncertainties, such as measurement errors, or model bias (with respect to the reality). To deal with this unknown, we are going to model it as a sample of a random variable, leading us to treat the obtained data as a random sample as well.

From the diverse assumptions we can make upon this sampled random variable, we can then treat the calibration procedure as a parameter estimation problem of a random variable. The estimated parameter will be written $\hat{\theta}$, and the subscript will denote additional information on the estimator. In this thesis, we focus on extremum estimators. Those estimators are defined as the optimiser of a given objective function J , $\hat{\theta} = \arg \min J$. In the next sections, we will see the probabilistic origins of a few classical objective functions.

1.3.2 Frequentist inference, MLE

As mentioned before, we can model the observations as a random variable, say Y (uppercase to highlight its random nature), and assume that this r.v. belongs to a family of parametric distributions, whose densities are

$$\{y \mapsto p_Y(y; \theta); \theta \in \Theta\} \quad (1.38)$$

This choice of notation has been made to keep explicit the dependency on θ . Assuming now that the residual are normally distributed with a given covariance matrix Σ , and that $\mathbb{Y} \subseteq \mathbb{R}^n$, Y is a random vector distributed as

$$Y \sim \mathcal{N}(\mathcal{M}(\theta), \Sigma) \quad (1.39)$$

whose one sample is $y = \mathcal{M}(\vartheta)$.

Now, instead of looking at the densities of Eq. (1.38) as functions taking as arguments the samples in $\mathbb{Y}$, we may look at it as a function of θ , as the observations $y \in \mathbb{Y}$ do not vary. We can then define the likelihood function and its associated extremum estimator.

Definition 1.3.1 – Likelihood function, MLE: The probability density function of the observations for a set of parameters is called the likelihood of those parameters given the observations, and is written $\mathcal{L}$. In the Gaussian case, this can be written as

$$\mathcal{L}(\cdot; y) : \theta \mapsto p_Y(y; \theta) = \mathcal{L}(\theta; y) \quad (1.40)$$

$$= (2\pi)^{-n/2} |\Sigma|^{-1/2} \exp \left(-\frac{1}{2} (\mathcal{M}(\theta) - y)^T \Sigma^{-1} (\mathcal{M}(\theta) - y) \right) \quad (1.41)$$

If $\Sigma = \text{diag}(\sigma_1^2, \dots, \sigma_n^2)$, the likelihood can be written as the product of 1D Gaussians:

$$\mathcal{L}(\theta; y) = \left(\prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma_i} \right) \exp \left(\sum_{i=1}^n -\frac{(\mathcal{M}(\theta)_i - y_i)^2}{2\sigma_i^2} \right) \quad (1.42)$$

$$= \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma_i} \exp \left(-\frac{(\mathcal{M}(\theta)_i - y_i)^2}{2\sigma_i^2} \right) \quad (1.43)$$

with $y = (y_1, \dots, y_n)$ and $\mathcal{M}(\theta) = (\mathcal{M}(\theta)_1, \dots, \mathcal{M}(\theta)_n)$. Based on the likelihood function, we can define the *Maximum Likelihood Estimator*, or *MLE*, that maximises the likelihood defined above:

$$\theta_{\text{MLE}} = \arg \max_{\theta \in \Theta} \mathcal{L}(\theta; y) \quad (1.44)$$

For practical and numerical reasons, the maximisation of the likelihood is often replaced by the minimisation of the negative log-likelihood:

$$\theta_{\text{MLE}} = \arg \min_{\theta \in \Theta} -\log \mathcal{L}(\theta; y) = \arg \min_{\theta \in \Theta} -\sum_{i=1}^n \log p_{Y_i|\theta}(y_i | \theta) \quad (1.45)$$

where

$$-\log \mathcal{L}(\theta; y) = \frac{1}{2}(\mathcal{M}(\theta) - y)^T \Sigma^{-1}(\mathcal{M}(\theta) - y) + \frac{n}{2} \log(2\pi) + \frac{1}{2} \log|\Sigma| \quad (1.46)$$

As the optimisation is performed on θ , we can remove the constant terms of the objective function, and rewrite the objective function as a L^2 norm in Eq. (1.47).

$$\begin{aligned} \theta_{\text{MLE}} &= \arg \min_{\theta \in \Theta} \frac{1}{2}(\mathcal{M}(\theta) - y)^T \Sigma^{-1}(\mathcal{M}(\theta) - y) \\ &= \arg \min_{\theta \in \Theta} \frac{1}{2} \|\mathcal{M}(\theta) - y\|_{\Sigma^{-1}}^2 \end{aligned} \quad (1.47)$$

Frequentist inference and Maximum Likelihood estimation boils down to Generalized non-linear least-square regression, that minimises the squared Mahalanobis distance between $\mathcal{M}(\theta)$ and y (Mahalanobis, 1936). This is only true as we assumed a Gaussian form of the errors in Eq. (1.39). Other choices of the sampling distribution will result in different objective functions. To reduce the sensitivity on outliers, some authors such as Rao et al. (2015) introduce Student or Laplace distributed errors, or specifically designed norm such as the Huber norm (Huber, 2011).

If the covariance matrix is diagonal, the residual errors are then uncorrelated, thus independent due to their Gaussian nature as defined in Eq. (1.39). The likelihood can be rewritten as the product of densities evaluated at the different samples y_i , obtained from their true distribution Y . A direct link can be written between the KL-divergence and the MLE. The KL-divergence between the true density p_Y and the parametric sampling distribution $p_Y(\cdot; \theta)$ is

$$D_{\text{KL}}(p_Y \| p_Y(\cdot; \theta)) = \mathbb{E}_Y [\log p_Y(Y)] - \mathbb{E}_Y [\log p_Y(Y; \theta)] \quad (1.48)$$

As the first term does not depend on θ , minimising this expression is equivalent to minimising the second part of the equation, so

$$\arg \min_{\theta \in \Theta} D_{\text{KL}}(p_Y \| p_Y(\cdot; \theta)) = \arg \min_{\theta \in \Theta} -\mathbb{E}_Y [\log p_Y(Y; \theta)] \quad (1.49)$$

The true distribution of the observation is unknown, but samples y_i are available. Using the empirical KL-divergence denoted $D_{\text{KL}}^{\text{empirical}}$, and replacing the theoretical expectation with the empirical one, the equation above becomes:

$$\arg \min_{\theta \in \Theta} D_{\text{KL}}^{\text{empirical}}(p_Y \| p_Y(\cdot; \theta)) = \arg \min_{\theta \in \Theta} \frac{1}{n} \sum_{i=1}^n -\log p_Y(y_i; \theta) = \theta_{\text{MLE}} \quad (1.50)$$

Thus, the MLE minimises the empirical KL-divergence between the true distribution of the observations and the sampling distribution of the observation (that depends on θ).

The MLE possesses desirable asymptotic properties, such as asymptotic normality when the number of observations grows large (Reid, 2013). Those properties permit the construction of asymptotic confidence intervals, and to test hypothesis for model selection. This aspect will be further developed in Section 1.5.

So far, the only information assumed on θ is its parameter space Θ . In the case where some belief on θ is present before the calibration, we can incorporate this information using the Bayesian framework.

1.3.3 Bayesian Inference

In Bayesian inference, the uncertainty present on θ is modelled by considering it as a random variable. Instead of having a precise value for θ , albeit unknown, we assume that we have a *prior distribution* on θ , denoted p_θ , that represents the initial state of belief upon the parameter, prior to any experiment and observations. The choice of this prior distribution will be discussed later. Using the experiment, whose sampling distribution is given by the likelihood, the prior distribution is updated to reflect the new state of belief upon the parameter. The Gaussian likelihood in Eq. (1.39) for the frequentist approach can be almost be rewritten as is in the Bayesian setting, just by conditioning Y with θ . Eq. (1.39) becomes

$$Y | \theta \sim \mathcal{N}(\mathcal{M}(\theta), \Sigma) \quad (1.51)$$

and the likelihood is the pdf $\mathcal{L}(\theta; y) = p_{Y|\theta}(y | \theta)$. Using Bayes' theorem, the *posterior distribution* of the parameters given the observed data is

$$p_{\theta|Y}(\theta | y) = \frac{p_{Y|\theta}(y | \theta)p_\theta(\theta)}{p_Y(y)} = \frac{\mathcal{L}(\theta; y)p_\theta(\theta)}{p_Y(y)} \quad (1.52)$$

The denominator can be seen as a normalizing constant, ensuring that $\int_{\Theta} p_{\theta|Y} = 1$. But it can also be seen as a measure of how well does the model explain the data obtained. This interpretation will be extended in Section 1.5

Definition 1.3.2 – Model Evidence: The model evidence, (or marginal likelihood, integrated likelihood) is defined as the distribution of the data marginalised over the parameters:

$$p_Y(y) = \int_{\Theta} p_{Y,\theta}(y, \theta) d\theta = \int_{\Theta} p_{Y|\theta}(y | \theta)p_\theta(\theta) d\theta \quad (1.53)$$

This quantity depends implicitly on the underlying mathematical model $\mathfrak{M} = (\mathcal{M}, \Theta)$. Comparing evidence of different models allows for the comparison of those different models. However, computing the model evidence requires the expensive evaluation of an integral over the whole parameter space, and no analytical form is available except for trivial cases. Specific techniques for this evaluation are reviewed in Friel and Wyse (2011).

When the model $(\mathcal{M}, \Theta)$ and the data y is fixed, the model evidence is constant with respect to the calibration parameter θ . The posterior distribution is thus often written and evaluated up to a multiplicative constant.

$$p_{\theta|Y}(\theta | y) \propto \mathcal{L}(\theta; y)p_\theta(\theta) \quad (1.54)$$

1.3.3.a Posterior inference

This posterior distribution is central in Bayesian analysis, as it gathers all the information we have on the parameter, given the observed data. Given Eq. (1.52), evaluating the posterior density at a point requires the evaluation of the model evidence, that is an expensive integral. To bypass this evaluation, several techniques have been developed to get samples from a unnormalized arbitrary function. One of the most

well-known method is based on the construction of a Markov-chain whose stationary state is the searched posterior. Classical MCMC algorithms such as Metropolis-Hastings require the use of a proposal density, and then to accept or reject the proposal based on the posterior distribution evaluated at the point.

A lot of refinement of these methods are available in the literature in order to better tackle the high-dimensionality of the parameter space, or to improve the mixing of the sampled MC chain. One important adaptation to mention is Hamiltonian Monte-Carlo ([Hanson, 2001](#); [Betancourt, 2017](#)), that improves the performance of the chain by using the value of the gradient of the log-posterior distribution. Obtaining this gradient (although for a different purpose) is discussed in [Section 1.4](#).

For time-dependent systems, Bayesian framework is particularly well-suited to treat observations sequentially, especially because Bayesian updating is done via multiplication. Bayes' theorem is the basis of many data assimilation methods, such as Kalman filter or various particle filters, that are often used for state estimation.

1.3.3.b Bayesian Point estimates

The whole posterior distribution aggregates a lot of information on the problem. However, as mentioned above, a certain work has to be done in order to get independent samples. Instead, one can try to find a point $\theta \in \Theta$ that summarizes this distribution. Consequently, the chosen estimate is often an indicator of the central tendency. In that sense, we wish to get a value that is quite close to all sampled values from the posterior ([Lehmann and Casella, 2006](#)).

Let us define a function L that measures a distance in the parameter space: $L : \Theta \times \Theta$. For a candidate θ' , the measured risk with respect to a sample from the posterior $\theta_{\text{sample}} \sim \theta \mid Y$ is $L(\theta', \theta_{\text{sample}})$. The *Bayesian risk* for θ' is then the expectation of this Bayesian loss functions L under the posterior distribution: $\mathbb{E}_{\theta \mid Y} [L(\theta', \theta) \mid y]$. A Bayesian point estimate is defined as a minimiser of the Bayesian risk:

$$\theta_L = \arg \min_{\theta' \in \Theta} \mathbb{E}_{\theta \mid Y} [L(\theta', \theta) \mid y] \quad (1.55)$$

Obviously, different loss functions will lead to different Bayesian point estimates, and we are going to evoke two of them.

Posterior mean

By defining L as the squared error $L(\theta', \theta) = (\theta' - \theta)^2$ (or $(\theta' - \theta)^T(\theta' - \theta)$ if $\dim \Theta > 1$), we can define the Mean Squared Error as $\text{MSE} : \theta' \mapsto \mathbb{E}_{\theta \mid Y} [(\theta' - \theta)^2 \mid y]$. Finally, the value corresponding to the Minimum Mean Squared Error is

$$\theta_{\text{MMSE}} = \arg \min_{\theta' \in \Theta} \mathbb{E}_{\theta \mid Y} [(\theta' - \theta)^2 \mid y] \quad (1.56)$$

Simple algebraic manipulations show that the minimiser is in fact the posterior mean:

$$\theta_{\text{MMSE}} = \mathbb{E}_{\theta \mid Y} [\theta \mid y] = \int_{\Theta} \theta \cdot p_{\theta \mid Y}(\theta \mid y) \, d\theta \quad (1.57)$$

In order to compute θ_{MMSE} , it is easier to compute directly the mean of the posterior samples obtained via posterior inference, than to solve the minimisation problem in Eq. (1.56).

Posterior Mode: the MAP

By reaching a bit on the notion of function for L and choosing $L(\theta', \theta) = -\delta_\theta(\theta')$, the dirac delta function defined in Eq. (1.14), one can show that the minimiser of $\mathbb{E}_{\theta|Y} [L(\theta', \theta) | y]$ is the mode of the posterior distribution, and is called the *Maximum A Posteriori* (MAP):

$$\begin{aligned}\theta_{\text{MAP}} &= \arg \min_{\theta' \in \Theta} \mathbb{E}_{\theta|Y} [\delta_\theta(\theta') | y] = \arg \min_{\theta' \in \Theta} -p_{\theta|Y}(\theta' | y) \\ &= \arg \max_{\theta' \in \Theta} p_{\theta|Y}(\theta' | y) = \arg \max_{\theta' \in \Theta} \mathcal{L}(\theta'; y) p_\theta(\theta')\end{aligned}\quad (1.58)$$

One interesting fact about the MAP, is that its evaluation does not require the full knowledge of the posterior distribution, nor samples to evaluate the integral of Eq. (1.57). We can resort to classical optimisation techniques for this evaluation. Similarly to the likelihood, taking the negative logarithm leads to the following minimisation problem.

$$\theta_{\text{MAP}} = \arg \min_{\theta' \in \Theta} -\log \mathcal{L}(\theta'; y) - \log p_\theta(\theta') \quad (1.59)$$

1.3.3.c Choice of a prior distribution

As seen in the application of Bayes' theorem in Eq. (1.52), the prior has a preponderant role in the formulation of the posterior distribution. Indeed, this prior distribution represents the current state of knowledge on the value of the parameter, before any experiment. This comes usually from an expert opinion, or some reasonable assumptions about the nature of θ .

Let us assume for instance that we have a Gaussian prior for θ : $\theta \sim \mathcal{N}(\theta_b, B)$ where B is called the background covariance error matrix and θ_b is called the *background value* that acts as a plausible reference value. Assuming a Gaussian form for the errors as well with covariance matrix Σ , the MAP can be written as

$$\theta_{\text{MAP}} = \arg \min_{\theta \in \Theta} \frac{1}{2} \|\mathcal{M}(\theta) - y\|_{\Sigma^{-1}}^2 + \frac{1}{2} \|\theta - \theta_b\|_{B^{-1}}^2 \quad (1.60)$$

Adding a Gaussian prior for the parameter comes down to adding a L^2 regularization term to the optimisation problem, also called Tikhonov regularization (Tikhonov and Arsenin, 1977). This expression is very analogous to the state estimation in the 3D-Var method in Data assimilation. Other choices of priors lead to other regularizations, such as the lasso regularization (Tibshirani, 2011) that is a consequence for choosing θ that follows a priori a Laplace distribution of mean 0.

The choice of a prior distribution has an influence on the inference of the parameter and its point estimation. Where there is no knowledge on the parameter beforehand, one can try to choose a non-informative prior in order to try to mitigate its effect. One can for instance choose a “flat” prior over the parameter space, but this can lead to

improper priors, in the sense that they do not integrate to 1. However, improper priors do not necessarily lead to improper posterior, allowing for the usual Bayesian analysis of the quantity. For instance, if $\Theta = \mathbb{R}^p$, the prior $p_\theta(\theta) \propto 1$ is improper, but the MAP estimation is equivalent to the MLE.

All in all, when looking for the MAP or the MLE, parameter estimation boils down to the minimisation of a well chosen objective function, that measures the misfit between the output of the numerical model and the observations. This objective function will be written J in the following, to match the notation of data assimilation. In this context of calibration, we can then summarize the estimation as a minimisation problem, where J represents some kind of distance between $\mathcal{M}(\theta)$ and the observations.

$$\theta = \arg \min_{\theta \in \Theta} J(\theta) \quad (1.61)$$

1.4 Calibration using adjoint-based optimisation

Point estimates in this context take the form of extremum estimators, that is an extremum of some given objective function J . This function is proportional to the log-likelihood for the MLE, or the log-posterior for the MAP, but other misfits can be considered, such as optimal transport based metrics. The formulation is then quite simple, but the problem of efficient optimisation remains. For differentiable problems, most of minimisation instances are solved using gradient-based methods, such as gradient descent or quasi-Newton methods.

This implies to be able to compute efficiently the gradient of the objective function J with respect to the parameter: $\nabla_\theta J$. The straightforward way is to compute the gradient using finite differences. Let us suppose that $\theta = (\theta_1, \dots, \theta_p)$, and e_i is 0 for all its component except the i th one which is 1. The gradient can be approximated by the usual 1st order forward finite-difference scheme:

$$\nabla_\theta J \approx \left[\frac{J(\theta + \epsilon e_1) - J(\theta)}{\epsilon}, \frac{J(\theta + \epsilon e_2) - J(\theta)}{\epsilon}, \dots, \frac{J(\theta + \epsilon e_p) - J(\theta)}{\epsilon} \right] \quad \text{for } \epsilon \ll 1 \quad (1.62)$$

In addition to the run of the model at θ , we have to evaluate the model p times, for each one of the coordinate of θ . If this is feasible in practice for low-dimensional problems, this is impossible for large problems that count more than hundreds of parameters. Nevertheless, different methods can be used to compute the gradient, at least approximately for optimisation purpose: for instance, [Boutet \(2015\)](#) uses Simultaneous Perturbation Stochastic Approximation to approximate the gradient using only one additional run, independently on the number of parameters.

In geophysical applications, parameter estimation and the subsequent optimisation is usually performed by deriving the adjoint equations in order to get the exact gradient for a relatively reasonable cost. This gradient is used afterward in optimisation methods such as conjugate gradient, or BFGS for instance. Adjoint methods are thus very popular in large-scale optimisation of Computational Fluid Dynamics codes, as the additional cost of implementation is often worth the gain in the short term. This situation is common in data assimilation, as shown in [Das and Lardner \(1991, 1992\)](#); [Honnorat et al. \(2010\)](#); [Couderc et al. \(2013\)](#), or in shape optimisation of airfoils in [Huyse and Bushnell \(2001\)](#).

To derive the adjoint equations, we will first rewrite the objective function as a function of the forward operator and the parameter: $J(\theta) = J(\mathcal{M}(\theta), \theta)$. The estimation of the parameter can be written as the following constrained optimisation problem:

$$\begin{aligned} \min_{\theta \in \Theta} J(\theta) &= J(y, \theta) \\ \text{such that } \mathcal{F}(y, \theta) &= 0 \end{aligned} \quad (1.63)$$

where the constraint on $\mathcal{F}$ signifies that the model is admissible, *i.e.* that $y = \mathcal{M}(\theta) \in \mathbb{Y}$.

Differentiating the Eq. (1.63) with respect to θ using the chain rule gives

$$\begin{aligned} \nabla_{\theta} J &= \frac{\partial J}{\partial y} \frac{\partial y}{\partial \theta} + \frac{\partial J}{\partial \theta} \\ \nabla_{\theta} \mathcal{F} &= \frac{\partial \mathcal{F}}{\partial y} \frac{\partial y}{\partial \theta} + \frac{\partial \mathcal{F}}{\partial \theta} \end{aligned} \quad (1.64)$$

In those equations, the partial derivatives with respect to θ are quite easily obtainable, while the real challenge is to obtain the derivative with respect to the state variable: $\frac{\partial}{\partial y}$.

To treat the constrained optimisation in Eq. (1.63), let us introduce the Lagrange multiplier $\lambda \in \mathbb{Y}$, so that we can write the Lagrangian $\mathcal{L}$:

$$\mathcal{L}(\theta, y, \lambda) = J(y, \theta) - \lambda^T \mathcal{F}(y, \theta) \quad (1.65)$$

and the unconstrained optimisation problem is then

$$\min_{\theta, y, \lambda} \mathcal{L}(\theta, y, \lambda) \quad (1.66)$$

The first-order condition of optimality for the Lagrangian: $\frac{\partial \mathcal{L}}{\partial \theta} = \frac{\partial \mathcal{L}}{\partial y} = \frac{\partial \mathcal{L}}{\partial \lambda} = 0$ translates into the optimality condition, adjoint equation and the state equation. Indeed, when differentiating with respect to the adjoint variable, we retrieve the state equation:

$$\frac{\partial \mathcal{L}}{\partial \lambda} = -\mathcal{F}(y, \theta) = 0 \quad (\text{State equation})$$

When differentiating with respect to the state variable, the equation that verifies the adjoint variable is called the adjoint equation

$$\frac{\partial \mathcal{L}}{\partial y} = \frac{\partial J}{\partial y} - \lambda^T \frac{\partial \mathcal{F}}{\partial y} = 0 \quad (\text{Adjoint equation})$$

Finally, when λ verifies the adjoint equation: $\left(\frac{\partial \mathcal{F}}{\partial y}\right)^T \lambda = \left(\frac{\partial J}{\partial y}\right)^T$, the gradient of the objective function can be expressed using the partial derivative *with respect to* θ of the objective function and of the forward model, and the adjoint variable:

$$\frac{\partial \mathcal{L}}{\partial \theta} = \nabla_{\theta} J = \frac{\partial J}{\partial \theta} - \lambda^T \frac{\partial \mathcal{F}}{\partial \theta} = 0 \quad (\text{Optimality condition})$$

So, to get $\nabla_{\theta} J$, as the partial derivatives with respect to the control variables are relatively easy to obtain, the challenge lies in solving the adjoint equation. Albeit

tedious, one can derive those equations by writing the tangent linear model of the original model, and implement a dedicated solver for the adjoint variables. A more common and way simpler approach is to derive the adjoint equations directly from the computer code implemented to solve the model, by using Automatic differentiation tools, such as TAPENADE (Hascoet and Pascual, 2013). Those programs directly translate the source code into a program that solves the original model equations and the adjoint equations, and outputs the gradient along with the objective function.

1.5 Model selection

So far, we have discussed the calibration of a specific model $(\mathcal{M}, \Theta)$ given some observations, thus solving an inverse problem and finding θ as an extremum of an objective function. But different models may be considered to explain the data. Those models may differ by their forward operator, by their parameter space, or by both at the same time.

But changing models also means changing the potential “best” fit attainable, in terms of minimum reached by the objective function. More complex models usually provide a better fit of the model but to the cost of a higher dimension in the parameter space. At the same time, more complex models may exhibit an overfitting behaviour.

Example 1.5.1: Figure 1.6 shows a curve-fitting problem using polynomial functions, where the y_i ’s are realisations of $Y_i \sim \mathcal{N}(i, 1)$ for $i = 0$ to 10. The objective function associated is $J(\theta) = \sum_{i=0}^{10} \|P(i; \theta) - y_i\|^2$, where P is a polynomial of degree $\dim \Theta$, and whose coefficients are given by $\theta = (\theta_1, \dots, \theta_n)$. For this problem of curve fitting, increasing the degree of the polynomial used (thus the dimensionality of the model) decreases the minimum value of J reached. However, the increase in the degree leads also to some oscillations between the sampled points, as the fitting procedure looks to account for the deviations due to the random origin of the y_i ’s

We can then look to reduce the complexity of the model, without decreasing significantly its performance.

We are first going to consider the case of nested models: models that share the same forward model, but whose parameter spaces are nested. Model selection in this case is a way to reduce the dimension of the model, by reducing the parameter space.

Finally, tools introduced in this section bring up model comparison. A calibrated model $(\mathcal{M}, \{\hat{\theta}\})$ is optimal given an objective function, but we can show that values close to the calibrated value $\hat{\theta}$ may also be of interest, as the decrease of performance may not be statistically significant. The “perturbated” model $(\mathcal{M}, \{\hat{\theta} + \varepsilon\})$ for small ε may accurately describe the data as well.

1.5.1 Likelihood ratio test and relative likelihood

Generally speaking, more complex models have a better ability to represent the observations, but all the parameters included in the model may not be relevant for the modelling. It can be interesting to test if a “simpler” model would give similar

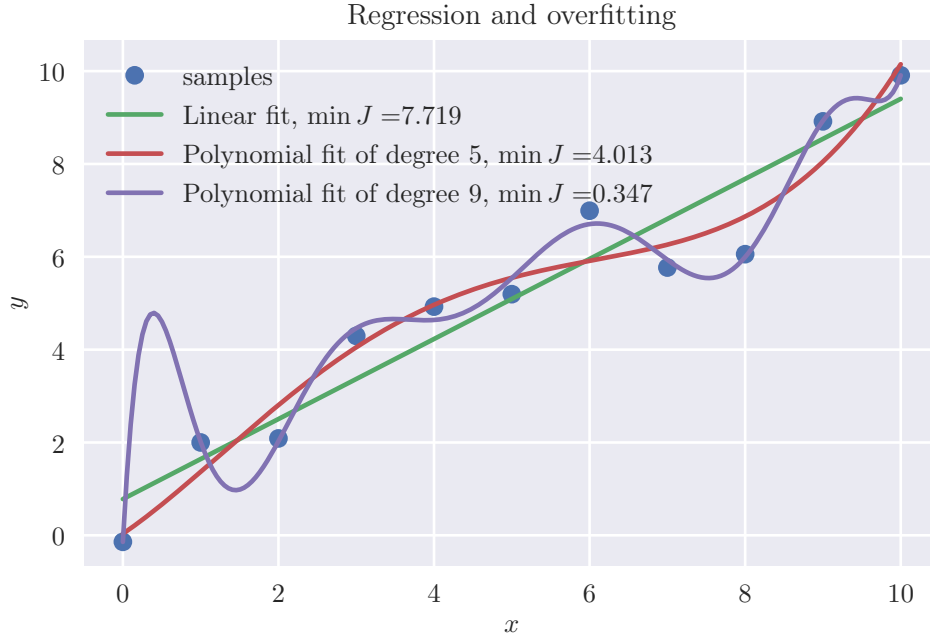


Figure 1.6 – Overfitting phenomenon, and reduction of the minimal value of the objective function

performances, or at least show a decrease in performances that is not statistically significant. One of the most well-known test is the likelihood-ratio test, that tests if two *nested models* are equivalent: Let us consider two nested models: $\mathfrak{M}_1 = (\mathcal{M}_1, \Theta_1)$, $\mathfrak{M}_2 = (\mathcal{M}_2, \Theta_2)$, such that $\mathcal{M}_1 = \mathcal{M}_2 = \mathcal{M}$ and $\Theta_2 \subsetneq \Theta_1$. In this case, $\mathfrak{M}_2$ represents the simpler model, with a reduced parameter space, while $\mathfrak{M}_1$ is the more general model. Recalling the notion of model dimension in [Remark 1.2.2](#), $\mathfrak{M}_1$ has dimension r , and $\mathfrak{M}_2$ has dimension d with $r > d$. As $\mathfrak{M}_1$ is more general, one can expect better performances.

The likelihood-ratio is defined as the ratio of the largest values taken by the likelihood on their respective parameter space, value that is assumed to be attained at $\hat{\theta}_1$ and $\hat{\theta}_2$.

$$\Lambda(y) = \frac{\sup_{\theta \in \Theta_2} \mathcal{L}(\theta; y)}{\sup_{\theta \in \Theta_1} \mathcal{L}(\theta; y)} = \frac{\mathcal{L}(\hat{\theta}_2; y)}{\mathcal{L}(\hat{\theta}_1; y)} \leq 1 \quad (1.67)$$

Based on this quantity, we can test whether the smaller model is sufficient to explain the observations as well as the larger model. The two hypothesis for this test are

- The null hypothesis $\mathcal{H}_0$: The two models are statistically equivalent: the difference between the maximal values of the likelihood is not statistically significant. This corresponds to Λ close to 1
- The alternative hypothesis $\mathcal{H}_1$: the two models are statistically different: the larger model performs better than the reduced one. This corresponds to Λ significantly smaller than 1.

Under the null hypothesis $\mathcal{H}_0$, $-2 \log \Lambda$ follows asymptotically (as the number of observations grows large) a χ^2 distribution introduced in [Example 1.2.17](#), according to Wilks' theorem ([Wilks, 1938](#)). The number of degrees of freedom of the χ^2 distribution is given by the difference of dimensionality between the two models:

$$-2 \log \Lambda(y) \xrightarrow{d} \chi_{r-d}^2 \quad (1.68)$$

By denoting $\chi_{r-d}^2(1 - \nu)$ the quantile of order $1 - \nu$ of the χ^2 distribution with $r - d$ degrees of freedom, the asymptotic rejection region of the null hypothesis at level ν is:

$$\text{RejReg}_\nu = \{y \mid -2 \log \Lambda(y) > \chi_{r-d}^2(1 - \nu)\} \quad (1.69)$$

So if the data $y \in \text{RejReg}_\nu$, we reject $\mathcal{H}_0$ at the ν -level, thus we accept $\mathcal{H}_1$. By reformulating using the log-likelihoods and objective functions $l(\theta; y) = \log \mathcal{L}(\theta; y) = -J(\theta)$

$$\text{RejReg}_\nu = \left\{ y \mid \left(\sup_{\theta \in \Theta_1} l(\theta; y) - \sup_{\theta \in \Theta_2} l(\theta; y) \right) > \frac{1}{2} \chi_{r-d}^2(1 - \nu) \right\} \quad (1.70)$$

$$= \left\{ y \mid J(\theta_2) - J(\theta_1) > \frac{1}{2} \chi_{r-d}^2(1 - \nu) \right\} \quad (1.71)$$

Example 1.5.2 – Likelihood-ratio test for the MLE: Let $\Theta_1 = \Theta \subset \mathbb{R}$ and $\Theta_2 = \{\theta\} \subset \Theta_1$. We have then $\sup_{\Theta_1} \mathcal{L} = \mathcal{L}(\theta_{\text{MLE}})$, and $\sup_{\Theta_2} \mathcal{L} = \mathcal{L}(\theta)$. The difference of dimensionality is then $r - d = 1$ and the .95 quantile of the χ^2 r.v. is $\chi_1^2(1 - 0.05) = 3.84$. When the data falls into the rejection region (*i.e.* $J(\theta_2) = J(\theta)$ significantly larger than $J(\theta_1) = J(\theta_{\text{MLE}})$), the null hypothesis is rejected, and the models can be asserted significantly different.

Conversely, the rejection region defined above allows us to define an asymptotic confidence interval for the parameter. Let us consider $\Theta_1 = \Theta$, and $\Theta_2 = \{\theta\}$ as in [Example 1.5.2](#), and let us introduce the *Relative Likelihood* ([Kalbfleisch, 1985](#)) which is the ratio of the likelihood evaluated at a point θ to the maximal value of the likelihood:

$$R(\theta) = \frac{\mathcal{L}(\theta; y)}{\mathcal{L}(\theta_{\text{MLE}}; y)} = \frac{\mathcal{L}(\theta; y)}{\sup_{\theta' \in \Theta} \mathcal{L}(\theta'; y)} \quad (1.72)$$

This ratio allows for comparing the plausibility of the value θ , compared to the MLE. The likelihood interval of level $\gamma \in]0, 1]$ is defined as

$$\mathcal{I}_{\text{Lik}}(\gamma) = \left\{ \theta \mid R(\theta) = \frac{\mathcal{L}(\theta; y)}{\mathcal{L}(\theta_{\text{MLE}}; y)} \geq \gamma \right\} \quad (1.73)$$

γ can be set to an arbitrary threshold, but it can also be chosen specifically so that $\mathcal{I}_{\text{Lik}}(\gamma)$ is the complement of the rejection region of a likelihood-ratio test with certain confidence.

Using the likelihood-ratio test, as in [Example 1.5.2](#), we can test how well a given value of the parameter θ performs compared to the MLE by comparing the models $(\mathcal{M}, \{\theta\})$

and $(\mathcal{M}, \Theta)$. Let $R(\theta)$ be the corresponding likelihood-ratio. The complement of the rejection region [Eq. \(1.71\)](#) written using [Eq. \(1.73\)](#) is

$$\mathcal{I}_{\text{Lik}} \left(\exp \left(-\frac{1}{2} \chi_{\dim(\Theta)}^2 (1 - \nu) \right) \right) = \left\{ \theta \mid R(\theta) \geq \exp \left(-\frac{1}{2} \chi_{\dim(\Theta)}^2 (1 - \nu) \right) \right\} \quad (1.74)$$

The values of the calibrated parameters in this set generate models that are statistically equivalent to the model comprising the MLE as its calibrated parameter.

For 1 dimensional models and the confidence level of 0.05, the threshold of [Eq. \(1.74\)](#) is $\exp \left(-\frac{1}{2} \chi_{\dim(\Theta)}^2 (1 - \nu) \right) = \exp \left(-\frac{1}{2} \chi_1^2 (.95) \right) \approx 0.15$, and at a level 0.10, $\exp \left(-\frac{1}{2} \chi_1^2 (.90) \right) \approx 0.26$.

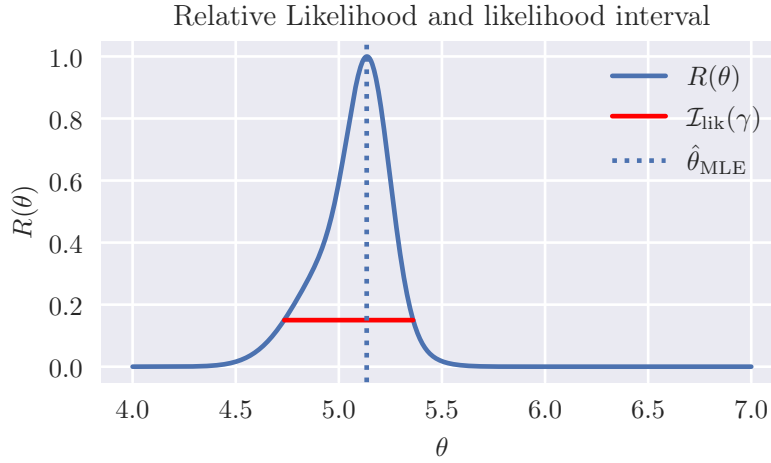


Figure 1.7 – Example of relative likelihood, and associated likelihood interval, for $\gamma = 0.15$

Due to the likelihood-ratio test and the relative likelihood, we can see that even though θ_{MLE} is the optimiser of the likelihood function, other values close to it may not be discarded, provided that the value of the log-likelihood does not drop off too much.

1.5.2 Criteria for non-nested model comparison

The likelihood-ratio test presented above is defined for nested models. In the more general case, we can also associate each model with a single numerical value that measures the balance between “fit” and complexity of the model. This usually takes the form of

$$\text{Crit}(\mathfrak{M}) = -2 \log \mathcal{L}(\theta_{\text{MLE}}) + \text{Complexity penalization} \quad (1.75)$$

where $\mathcal{L}$ is the likelihood function for the model $\mathfrak{M} = (\mathcal{M}, \Theta)$, and $\mathcal{L}(\theta_{\text{MLE}}) = \max \mathcal{L}$. The role of the complexity penalization is to avoid overfitting, and is often directly linked to the dimension of the parameter space Θ . Two quite popular examples of criteria are the AIC (Akaike Information Criterion) introduced in [Akaike \(1974\)](#) and the BIC (Bayesian Information Criterion) in [Schwarz \(1978\)](#).

Then, to compare two models $\mathfrak{M}_1$ and $\mathfrak{M}_2$, the magnitude of the difference $\text{Crit}(\mathfrak{M}_1) - \text{Crit}(\mathfrak{M}_2)$ is compared to some thresholds as shown in Burnham and Anderson (2004). The difference shows whether a model should be preferred, or if no substantial evidence exists for either model. These criteria, as well as the likelihood-ratio test, are based on the evaluation of the likelihood at its maximal value.

Another approach is to marginalize the likelihood with respect to the prior distribution of the calibration parameter, giving Bayesian model comparison. The criterion for a model $\mathfrak{M}_1$ is the evidence of the model, *i.e.* the pdf of the data given the model $\mathfrak{M}_1$, and $\text{Crit}(\mathfrak{M}_1) = \log \int_{\Theta} \mathcal{L}(\theta) p_{\theta}(\theta) d\theta$. Again, the logarithm of Bayes' factor is compared with specified thresholds (Kass and Raftery, 1995; Burnham and Anderson, 2004).

1.6 Parametric model misspecification

We introduced earlier the mathematical model $(\mathcal{M}, \Theta)$, and based our analysis on the fact that the “target model”, *i.e.* the reality is $(\mathcal{M}, \Theta_{\text{real}} \supseteq \Theta)$, so the parameter spaces are the same for the two models. However, between the reality and the numerical model, various simplifications are introduced, thus the reality is not often completely *representable* by the numerical model: we have then misspecified models.

Definition 1.6.1 – Misspecified model: Let Y be the random variable associated with the observations, and p_Y its pdf. Let $\{p_{Y|\theta}; \theta \in \Theta\}$ the parametric family of densities, among which we are looking to find p_Y . The model is said to be misspecified, if $p_Y \notin \{p_{Y|\theta}; \theta \in \Theta\}$.

We can also define this misspecification in terms of numerical models defined in this chapter: let $(\mathcal{M}, \Theta_{\text{real}})$ be the physical system under study and $(\mathcal{M}, \Theta)$ the numerical model that is to be calibrated with respect to the observations $y = \mathcal{M}(\vartheta)$. $(\mathcal{M}, \Theta)$ is said to be misspecified, if $\vartheta \notin \Theta$.

Finally, when we have a family of models $\{(\mathcal{M}(\cdot, u), \Theta) \mid u \in \mathbb{U}\}$, where each element depends on some additional parameter, and for all $u \in \mathbb{U}$, $\mathcal{M}(\cdot, u) \neq \mathcal{M}$, we talk about *parametric misspecification*.

In practice, in addition to the simplifications, the parameter space Θ does not contain necessarily all the parameters needed to run the forward model, but represents the space of the parameters of interest, or calibration parameters. In addition to them, some other parameters are at play, that we are going to call the *environmental parameters*, or *uncertain parameters* written $u \in \mathbb{U}$. These parameters come from instance from the external forcings.

Bayesian framework and more specifically Bayesian update of the prior by the likelihood puts the emphasis on the update of the information on the *parameter of interest*. However the environmental parameters are assumed to have an inherent variability. In that sense, it may not be worth spending time and resources to infer these parameter values, as they are bound to change. Moreover, we can only get information on the environmental conditions used to generate the observations.

In terms of models, each choice of $u \in \mathbb{U}$ gives a different model $\mathfrak{M}(u) = \{\mathcal{M}(\cdot, u), \Theta\}$. Let us assume that we can model the uncertain parameters as a random variable U . Let us consider that we chose a specific $u_0 \in \mathbb{U}$, and that we are given some observation $y = \mathcal{M}(\vartheta)$. We can formulate an inverse problem, and an objective function $J: \theta \mapsto J(\theta, u_0)$, that we wish to minimise with respect to θ . Some estimators still carry nice properties. The MLE for instance, defined [Section 1.3.2](#) can still be written as the minimiser of the empirical KL-divergence. We assume that we can write the sampling distribution as $p_{Y|\theta, U}$, and

$$\theta_{\text{MLE}}(u_0) = \arg \min_{\theta \in \Theta} D_{\text{KL}}^{\text{empirical}}(p_Y \| p_{Y|\theta, U}(\cdot | \theta, U = u_0)) \quad (1.76)$$

and can be seen as the “best” value given $U = u_0$. However, the asymptotic properties of the MLE are slightly different as described in [White \(1982\)](#). So when the model is misspecified, minimising the same objective function still makes sense.

However, the calibration will depend on the chosen u_0 : $\theta(u_0) = \arg \min_{\theta \in \Theta} J(\theta, u_0)$, and there is no guarantee that $\theta(u_0)$ will minimise $J(\cdot, u_1)$ for $u_0 \neq u_1$, as illustrated [Fig. 1.8](#).

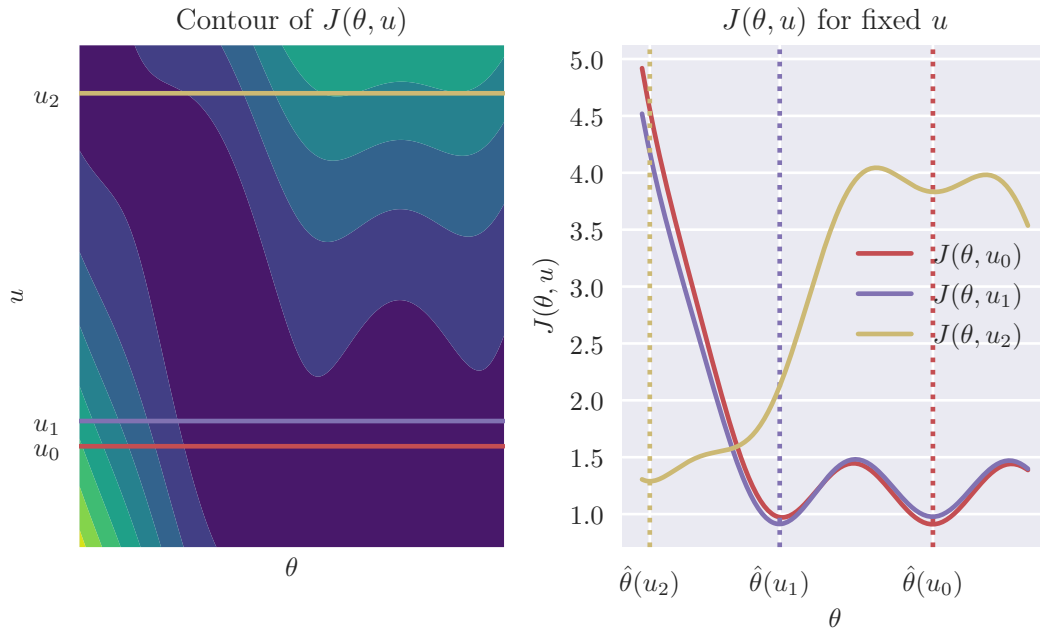


Figure 1.8 – Effect of the misspecification on the minimiser.

u_0 and u_1 are close to each other, but $\theta(u_0)$ and $\theta(u_1)$ are not. However, as the objective function shows similar values at those points, choosing either one would lead to a satisfactory outcome given $U = u_0$ or $U = u_1$. If $U = u_2$ is considered as well, the modeller may have a preference and choose $\theta(u_1)$ as the final estimator.

In terms of model selection, the asymptotic distribution of the likelihood ratio statistic defined [Section 1.5.1](#) is also slightly different. Instead of following a χ_r^2 distribution,

where r is the number of dimensions for the test, $-2 \log \Lambda$ will asymptotically have the same distribution as a weighted sum of r random variables, where each one have a χ_1^2 distribution and whose weights are the eigenvalues of a matrix involving the Jacobian and the Hessian of the log-likelihood (Kent, 1982).

This random misspecification leads to some issues in the calibration of the model, and it asks for a notion of robustness with respect to the environmental parameters.

1.7 Partial conclusion

In this chapter, starting from a probabilistic point of view, we established the usual tools encountered in model calibration: the misfit between the data and the numerical model is measured by an objective function J , that can be minimised using for instance gradient descent. From this optimisation, we can define a “acceptable” region for the estimate. In other words, values in this set yield an misfit that is not different enough to be completely discarded.

Adding an environmental variable as a random parameter introduces a random parametric misspecification to the model: each realization of this underlying random variable will yield a different estimation. In [Chapter 2](#), we will discuss the notion of robustness under this random misspecification, and introduce a family of robust estimators, inspired by model selection.

* * *

CHAPTER 2

ROBUST ESTIMATORS IN THE PRESENCE OF UNCERTAINTIES

Contents

2.1	Defining robustness	35
2.1.1	Classifying the uncertainties	35
2.1.2	Robustness and/or reliability	35
2.1.3	Robustness under parametric misspecification	36
2.2	Probabilistic inference	37
2.2.1	Frequentist approach	37
2.2.2	Bayesian approach	38
2.3	Variational approach	41
2.3.1	Decision under deterministic uncertainty set	41
2.3.1.a	Global optimisation	41
2.3.1.b	Worst-case optimisation	42
2.3.1.c	Regret maximin	42
2.3.2	Robustness based on the moments of an objective function	43
2.3.2.a	Expected loss minimisation, central tendency	44
2.3.2.b	Variance optimisation	45
2.3.2.c	Multiobjective optimisation	46
2.3.2.d	Higher moments in optimisation	47
2.4	Regret-based families of estimators	48
2.4.1	Conditional minimum and minimiser	49
2.4.2	Regret and model selection	51
2.4.2.a	Objective as the negative log-likelihood	51
2.4.2.b	Interval and probability of acceptability	52

2.4.3	Relative-regret	54
2.4.3.a	Absolute and relative error	54
2.4.3.b	Relative-regret estimators family	55
2.4.4	The choice of the threshold	57
2.5	Partial Conclusion	58

In the previous chapter, we introduced the problem of calibration of a numerical model with respect to a *calibration parameter* θ . This takes the form of the optimisation of an objective function. We also raised the problem of parametric misspecification of the numerical model with respect to the reality: $u \in \mathbb{U}$. Moreover, this misspecification is modelled by a random variable U with known distribution. One desirable property is that the calibrated model shows relatively good performances when the environmental variables vary, or in other words, we want the calibrated model to be *robust* with respect to the varying environmental parameters. In this chapter, we are going to introduce some criteria that aim at solving this *robust optimisation problem*. The actual computation of those estimates will be discussed in the next chapter.

2.1 Defining robustness

2.1.1 Classifying the uncertainties

In the Bayesian formulation of the problem, the uncertainty on the calibration parameter is modelled through the prior distribution, while the uncertain parameter, u has its own distribution. While mathematically similar, those two representations actually encompasses a significant difference: we are actively trying to reduce the uncertainty of the calibration parameter by Bayesian update, while the uncertainty on the environmental parameter is seen as a nuisance.

In that context, the very notion of uncertainty can be roughly split in two, as described in Walker et al. (2003):

- Aleatoric uncertainties, coming from the inherent variability of a phenomenon, *e.g.* intrinsic randomness of some environmental variables
- Epistemic uncertainties coming from a lack of knowledge about the properties and conditions of the phenomenon underlying the behaviour of the system under study

According to this distinction, the epistemic uncertainty can be reduced by investigating the effect of the calibration parameter θ upon the physical system, and choose it accordingly to an objective function. The uncertain variable u on the other hand is uncertain in the aleatoric sense, and cannot be controlled directly, as its value is destined to change. This is why we model it using a random variable U . This distinction, illustrated in Fig. 2.1, is a bit simplistic, as Kiureghian and Ditlevsen (2009) points out that deciding the type of uncertainties is up to the modeller, who decides on which parameters inference is worth doing.

2.1.2 Robustness and/or reliability

The notion of *robustness* is dependent on the context in which it is used. In this work, the term “robust” qualifies a model that behaves still nicely under uncertainties, or to put it in an other way, that is insensitive up to a certain extent to some perturbations. Moreover, robustness is often linked and sometimes confused to the semantically close notion of *reliability*. In Lelièvre et al. (2016) we can find summarized in Table 2.1 the difference between these notions, by defining optimality as the deterministic counterpart of robustness, and admissibility as the counterpart of reliability.

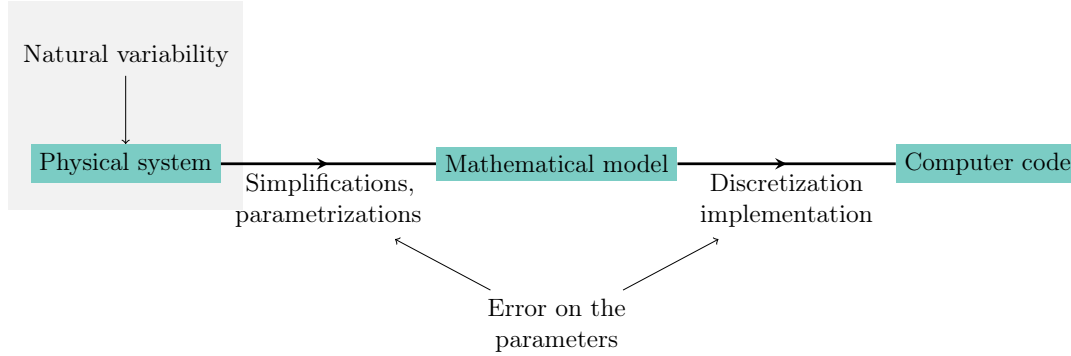


Figure 2.1 – Sources of uncertainties and errors in the modelling. The natural variability of the physical system can be seen as aleatoric uncertainties, and the errors on the parameters as epistemic uncertainties

	No objective	Objective with deterministic inputs	Objective with uncertain inputs
Unconstrained		Optimal	Robust
Deterministic constraints	Admissible	Optimal and admissible	Robust and admissible
Uncertain constraints	Reliable	Optimal and reliable	Robust and reliable

Table 2.1 – Types of problems, depending on their deterministic nature for the constraints or the objective. Shaded cells correspond to problems comprising an uncertain part. Reproduced from [Lelièvre et al. \(2016\)](#)

Other definitions of robustness can be encountered in the literature, and will not be treated in this work: Bayesian approaches are sometimes criticized for their use of subjective probabilities that represent the state of beliefs, especially on the choice of prior distributions. In that sense, robust Bayesian analysis aims at quantifying the sensitivity of the choice of the prior distribution on the resulting inference and relative Bayesian quantities derived. In the statistical community, robustness is often implied as the non-sensitivity on the outliers in the sample set.

2.1.3 Robustness under parameteric misspecification

Given a family of models $\{(\mathcal{M}(\cdot, u), \Theta), u \in \mathbb{U}\}$ and some observations $y \in \mathbb{Y}$ sampled from a random variable Y , we can derive a problem of parameter estimation for each $u \in \mathbb{U}$. As detailed in [Chapter 1](#), we can formulate the likelihood $\mathcal{L}$ and the posterior distribution, and then compute the MLE and the MAP.

Not taking into account the uncertainty on u may be an issue in the modelling, especially if the influence of this variable is non-negligible. Choosing a specific $u \in \mathbb{U}$ leads to *localized optimisation* ([Huyse and Bushnell, 2001](#)) and *overcalibration*, that is choosing a value $\hat{\theta}$ that is optimal for the given situation (which is induced by u). This value does not carry the optimality to other situations, or in Layman’s term according to [Andréassian et al. \(2012\)](#), being lured by “fool’s gold”. In geophysics and especially in hydrological models, this overcalibration may lead to the appearance of aberrations in the predictions as those uncertainties become prevalent sources of errors. In hydrology,

uncertainties are the principal culprit of the existence of “Hydrological monsters” (Kuczera et al., 2010), that are calibrated models that perform really badly.

There are two main ways to tackle this problem. Since the environmental parameter is random by nature with a known distribution, we can introduce it directly in the probabilistic inference framework, by appending u to the calibration parameter and to consider (θ, u) for the inference. This will be treated in Section 2.2. In this context, the additional environmental variables are usually called *nuisance parameters*.

Another way, that we are calling the *variational* approach, is to consider instead the objective function $(\theta, u) \mapsto J(\theta, u)$ that we want to minimise, as introduced in the previous chapter. Due to the uncertainty on u , we can then study the family of random variables indexed by $\theta \in \Theta : \{J(\theta, U); \theta \in \Theta\}$. This will be addressed in Section 2.3.

2.2 Probabilistic inference

In probabilistic inference, the environmental parameters are sometimes called *nuisance parameters*, and different ways have been studied to remove their influence. We will first detail likelihood-based methods and then the extension to Bayesian framework.

2.2.1 Frequentist approach

From a frequentist approach, we define the joint likelihood $\mathcal{L}(\theta, u; y) = p_{Y|\theta, U}(y | \theta, u)$. Under a Gaussian assumption, the sampling distribution, $Y | \theta, U$ is

$$Y | \theta, U \sim \mathcal{N}(\mathcal{M}(\theta, U), \Sigma) \quad (2.1)$$

where Σ is a covariance matrix.

There are two common ways to get rid of the nuisance parameters: one by *profiling*, one by *marginalization*. Profiling implies to perform first a maximisation of the likelihood with respect to the nuisance parameters:

$$\mathcal{L}_{\text{profile}}(\theta; y) = \max_{u \in \mathbb{U}} \mathcal{L}(\theta, u; y) \quad (2.2)$$

and

$$\theta_{\text{prMLE}} = \arg \max_{\theta \in \Theta} \mathcal{L}_{\text{profile}}(\theta; y) \quad (2.3)$$

In other words, considering the most favorable case of the likelihood given the nuisance parameters. Comparing the MLE over $\Theta \times \mathbb{U}$ for the original joint likelihood and the profile MLE on Θ for the profile likelihood, it is straightforward to verify that their components on Θ coincide as

$$\max_{(\theta, u) \in \Theta \times \mathbb{U}} \mathcal{L}(\theta, u; y) = \max_{\theta \in \Theta} \mathcal{L}_{\text{profile}}(\theta; y) \quad (2.4)$$

The resulting estimator does not take into account the uncertainty upon u , and can perform quite badly when the likelihood presents sharp ridges (Berger et al., 1999).

Another alternative is to define the *integrated*, or *marginalized* likelihood as

$$\mathcal{L}_{\text{integrated}}(\theta; y) = \int_{\mathbb{U}} \mathcal{L}(\theta, u; y) p_U(u) \, du \quad (2.5)$$

$$= \int_{\mathbb{U}} p_{Y|\theta, U}(y | \theta, u) p_U(u) \, du \quad (2.6)$$

$$= \int_{\mathbb{U}} p_{Y, U|\theta}(y, u | \theta) \, du \quad (2.7)$$

$$= p_{Y|\theta}(y | \theta) \quad (2.8)$$

and by maximising this function,

$$\theta_{\text{intMLE}} = \arg \max_{\theta \in \Theta} \mathcal{L}_{\text{integrated}}(\theta; y) \quad (2.9)$$

Example 2.2.1: In order to illustrate the difference between those two methods, the profile and integrated likelihood have been computed for the following likelihood:

$$Y | \theta, U \sim \mathcal{N}(\theta + u^2, 2^2) \quad (2.10)$$

and the observations $y = (y_1, \dots, y_{10})$ have been generated using $\theta + u^2 = 1$. We set $\Theta = [-5, 5]$ and $\mathbb{U} = [-2, 2]$. The likelihood evaluated on $\Theta \times \mathbb{U}$ is displayed Fig. 2.2, with the integrated and profile likelihood. We can see that there is not unicity of the maximiser for the profile likelihood: $\mathcal{L}_{\text{profile}}(\theta; y)$ is constant for $\theta \in [-3, 1]$. This is due to the fact that the observations can have been generated with any θ and u verifying $\theta + u^2 = 1$. For the integrated likelihood however, there is a unique maximum, attained for $\theta_{\text{intMLE}} \approx 0.8$.

2.2.2 Bayesian approach

Similarly as in Section 1.3.3, we can incorporate information on θ by introducing a prior distribution p_θ , and we can derive the posterior distribution using Bayes' theorem. We assume that U and θ are independent: $p_{\theta, U} = p_\theta \cdot p_U$. The likelihood of the data given θ and u is

$$\mathcal{L}(\theta, u; y) = p_{Y|\theta, U}(y | \theta, u) \quad (2.11)$$

The joint posterior distribution can be written as:

$$p_{\theta, U|Y}(\theta, u | y) = \mathcal{L}(\theta, u; y) p_\theta(\theta) p_U(u) \frac{1}{p_Y(y)} \quad (2.12)$$

$$\propto \mathcal{L}(\theta, u; y) p_\theta(\theta) p_U(u) \quad (2.13)$$

Here, the posterior is used to do inference on θ and u jointly. In order to suppress the dependency in u , we integrate with respect to U and get the marginalized posterior $p_{\theta|Y}$:

$$p_{\theta|Y}(\theta | y) = \int_{\mathbb{U}} p_{\theta, U|Y}(\theta, u | y) \, du \quad (2.14)$$

$$= \int_{\mathbb{U}} p_{\theta|Y, U}(\theta | y, u) p_{U|Y}(u | y) \, du \quad (2.15)$$

We can then define the *marginalized maximum a posteriori* (MMAP) (Doucet et al., 2002) as the maximiser of this marginalized posterior:

$$\theta_{\text{MMAP}} = \arg \max_{\theta \in \Theta} p_{\theta|Y}(\theta | y) \quad (2.16)$$

or, by taking the negative logarithm to get a minimisation problem, another writing is

$$\theta_{\text{MMAP}} = \arg \min_{\theta \in \Theta} -\log p_{\theta|Y}(\theta | y) \quad (2.17)$$

Unfortunately, neither the integration with respect to the nuisance parameter in Eq. (2.14) nor the subsequent optimisation is analytically easy. Assuming that we are able to get i.i.d. samples $\{(\theta_i, u_i)\}_{1 \leq i \leq n_{\text{samples}}}$ from the posterior distribution using MCMC methods for instance, by discarding the u components, the samples $\{\theta_i\}_{1 \leq i \leq n_{\text{samples}}}$ are distributed according to the marginal posterior $p_{\theta|Y}$, and thus can be used to get the MMAP. More direct techniques, such as Doucet et al. (2002), introduce methods in order to estimate iteratively the MMAP, through sampling of the joint posterior.

Example 2.2.2: Using the same data as in Example 2.2.1, we add a prior distribution of θ as a centered normal distribution, truncated on Θ . On Fig. 2.2 we can see the influence of the prior distribution, as it nudges the MMAP θ_{MMAP} toward 0, compared to the integrated likelihood.

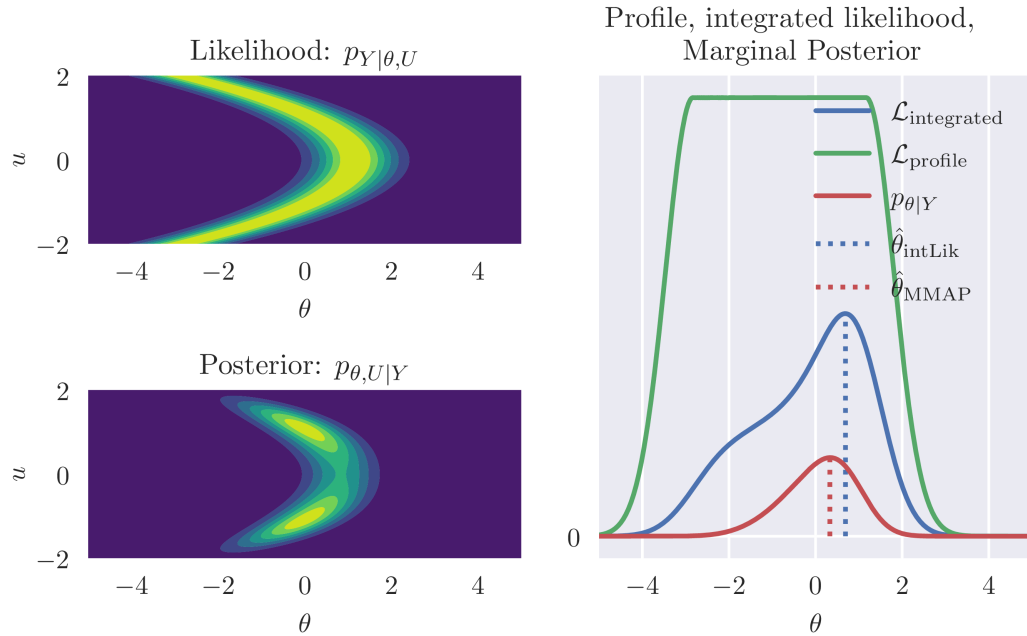


Figure 2.2 – Joint likelihood and posterior (left). Profile and integrated likelihood for an uniform nuisance parameter and marginal posterior distribution (right)

Those three estimators appear quite naturally in the probabilistic formulations. The main difference between the MMAP and the integrated likelihood is the presence of the prior distribution of θ in the formulation, thus similarly as in the inference problem of the previous chapter (without nuisance parameter), we can incorporate information on the calibration parameter. Regarding the profile likelihood however, this estimation relies on the optimisation with respect to u , and thus does not really take into account its random nature.

2.3 Variational approach

We discussed so far the calibration problem with nuisance parameters in the formulation of the likelihood or the posterior distribution. However, in data assimilation for instance, problems of parameter estimation are often formulated directly by introducing a cost function:

$$\begin{aligned} J : \Theta \times \mathbb{U} &\longrightarrow \mathbb{R}^+ \\ (\theta, u) &\longmapsto J(\theta, u) \end{aligned} \quad (2.18)$$

This function in a calibration context is measuring the misfit between the data y and the forward operator, and can be written as the negative log-likelihood, or the negative log-posterior distribution. Still, the objectives and criteria introduced in the following are not specific to this context, and J can represent other properties that ought to be reduced such as a loss or some unwanted physical properties such as the drag in airfoil design optimisation. This general problem is sometimes quoted as *Optimisation under uncertainties* (OOU) (Cook, 2018; Seshadri et al., 2014)

All in all, $J(\theta, u)$ represents the cost of taking the decision $\theta \in \Theta$ when the environmental variable is equal to u . We are going to make several assumptions:

- Θ is convex and bounded
- For all $\theta \in \Theta$ and $u \in \mathbb{U}$, $J(\theta, u) > 0$
- For all $\theta \in \Theta$, $J(\theta, \cdot)$ is measurable
- For all $\theta \in \Theta$, $J(\theta, U) \in L^s(\mathbb{P}_U)$ and $s \geq 2$. So for each θ , mean and variance exist and are finite.

As the function represents a cost, *i.e.* an undesirable property, we are interested in minimising in some sense this random variable, which depends on θ . Most of existing methods rely on the definition of a *robust* counterpart of the minimisation problem, which implies to remove the dependence on the uncertain variable U . This counterpart being a deterministic optimisation problem, we can solve it using classical methods of minimisation.

2.3.1 Decision under deterministic uncertainty set

We will first introduce some estimators that can be argued robust, even though the random nature of U is not directly taken into account. The uncertainty is modelled here by assuming that no information is available on u , except that $u \in \mathbb{U}$. In this paradigm, $\mathbb{U}$ is called the uncertainty set (Bertsimas et al., 2010).

2.3.1.a Global optimisation

A global optimisation criterion, as its name suggests, advocates for minimising the cost function over the whole space $\Theta \times \mathbb{U}$, giving this optimisation problem:

$$\min_{(\theta, u) \in \Theta \times \mathbb{U}} J(\theta, u) \quad (2.19)$$

Rearranging slightly this problem, the θ -component of the minimiser can be written as

$$\theta_{\text{global}} = \arg \min_{\theta \in \Theta} \min_{u \in \mathbb{U}} J(\theta, u) \quad (2.20)$$

The global minimum is the equivalent of profile likelihood maximisation defined [Eq. \(2.2\)](#), when J is the negative log-likelihood. This method exhibits some flaws: we are optimising the cost function only over the most favourable cases of the environmental parameter, thus there is no guarantee on the behaviour of J outside of those optimistic situations. It then makes sense to “separate” θ and u in the optimisation.

2.3.1.b Worst-case optimisation

As global optimisation is inherently optimistic, we can easily derive a criterion which is pessimistic in the sense that we want to minimise over the *least favourable* cases, thus minimising the objective in the worst-case scenarios. The optimisation problem in this case becomes

$$\min_{\theta \in \Theta} \max_{u \in \mathbb{U}} J(\theta, u) \quad (2.21)$$

This criterion is sometimes called Wald’s Minimax criterion ([Wald, 1945](#)), and the associated estimator is

$$\theta_{\text{WC}} = \arg \min_{\theta \in \Theta} \max_{u \in \mathbb{U}} J(\theta, u) \quad (2.22)$$

Minimising in the worst-case sense also possesses some flaws, especially from a computational point of view. First, the maximum on $\mathbb{U}$ may not exist, especially if $\mathbb{U}$ is unbounded: we could make the model perform as badly as possible by taking extreme values of u . Additionally, if it exists, the resulting estimator is most likely very conservative as only the worse cases are considered.

2.3.1.c Regret maximin

One other approach, called Savage’s maximin regret ([Savage, 1951](#)) is to compare the current objective to the best performance given the uncertain variable u . The translated objective is called the *regret* and is defined as

$$r(\theta, u) = J(\theta, u) - \min_{\theta \in \Theta} J(\theta, u) \quad (2.23)$$

Using the regret as the new objective function, we can optimise it in the worst-case sense, as introduced in [Section 2.3.1.b](#), and the minimum is attained at θ_{rWC} :

$$\theta_{\text{rWC}} = \arg \min_{\theta \in \Theta} \max_{u \in \mathbb{U}} r(\theta, u) \quad (2.24)$$

Example 2.3.1: [Figure 2.3](#) shows global, worst-case and regret optimisation for the analytical cost function

$$J(\theta, u) = (1 + u(\theta + 0.1))^2 (1 + (\theta - u)^2) \quad (2.25)$$

We can see how the worst-case minimisation (in blue) and Savage's maximin regret (in green) compare in this example. Maximin regret will favour values of θ giving an objective that is never too far from the optimal value available, in contrast to the worst-case that focuses on the absolute objective.

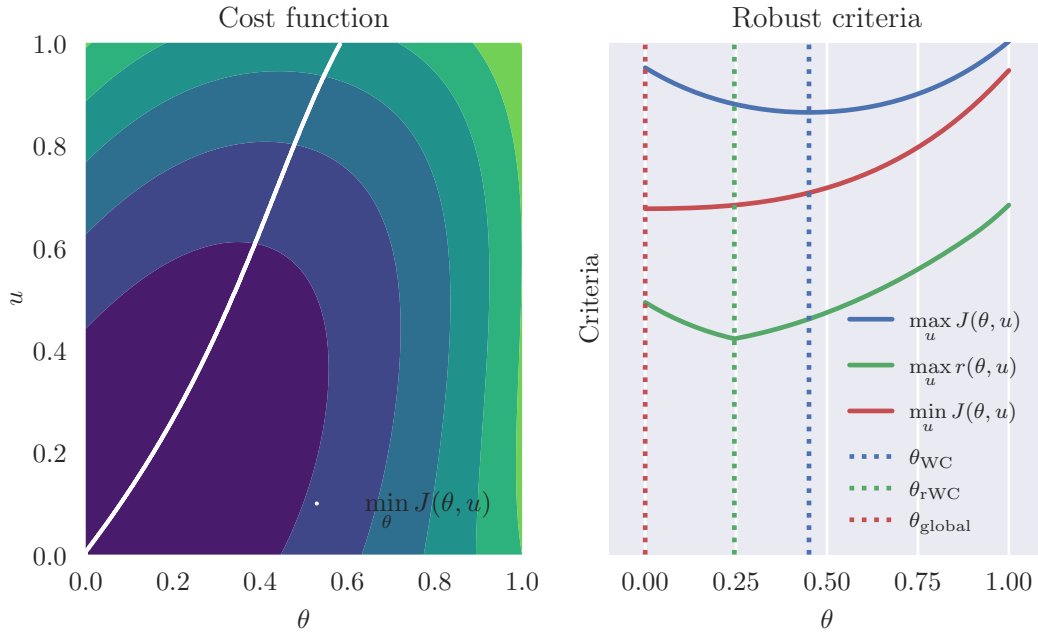


Figure 2.3 – Illustration of global optimisation, worst-case, and regret worst-case. The red points on the contour of the cost function correspond to the minimisers of J at each u fixed

So far, we did not use the fact that u was a realisation of a random variable, and did not take advantage of the knowledge we have upon it. In the next sections, we will see how to incorporate the knowledge of the distribution of U in the estimations.

2.3.2 Robustness based on the moments of an objective function

In the presence of uncertainties, choosing a parameter value θ can also be seen as making a choice under risk. Let $J : \Theta \times \mathcal{U} \rightarrow \mathbb{R}^+$ be an objective function, and assume that for all $\theta \in \Theta$, $J(\theta, \cdot)$ is a measurable function. J can be seen as the opposite of the *utility* function, often encountered in game theory or econometrics. Because of the random nature of U , we can define a family of real random variables $\{J(\theta, U) \mid \theta \in \Theta\}$, indexed by $\theta \in \Theta$. In [Beyer and Sendhoff \(2007\)](#), the authors define an *aggregation approach*, based on the integration with respect to the uncertain variable, in order to get an *aggregated objective*, which is a deterministic function that depends only on θ . An example of this aggregation is the integration of the successive powers of the cost function, in order to get the moments of the associated random variable, that we will detail in [Sections 2.3.2.a, 2.3.2.c and 2.3.2.d](#). The aggregated objective is then minimised with respect to the control variable.

2.3.2.a Expected loss minimisation, central tendency

One of the simplest approach when facing such a problem is to look to optimise a central tendency of those random variables. The mean value being an obvious candidate, we define the expected objective (or expected loss) as

$$\mu(\theta) = \mathbb{E}_U [J(\theta, U)] = \int_{\mathbb{U}} J(\theta, u) p_U(u) du \quad (2.26)$$

The expected objective $\mu(\theta)$ is sometimes called the conditional mean given θ . Taking the expectation of the objective function is very common in many problems of classification and regression (Bishop, 2006).

The conditional mean is minimised, giving θ_{mean} . Assuming that $J(\theta, u) \propto -\log \mathcal{L}(\theta, u; y)$, we have

$$\theta_{\text{mean}} = \arg \min_{\theta \in \Theta} \mu(\theta) = \arg \min_{\theta \in \Theta} \int_{\mathbb{U}} J(\theta, u) p_U(u) du \quad (2.27)$$

$$= \arg \min_{\theta \in \Theta} - \int_{\mathbb{U}} \log \mathcal{L}(\theta, u; y) p_U(u) du \quad (2.28)$$

$$= \arg \min_{\theta \in \Theta} - \int_{\mathbb{U}} \log (p_{Y|\theta, U}(y | \theta, u)) p_U(u) du \quad (2.29)$$

Taking the average of an objective function is the basis of *stochastic programming*. However, the integral of Eq. (2.26) is often intractable analytically, so instead of computing it exactly, one usually resorts to minimising the empirical mean risk. For $1 \leq i \leq n_U$, let u_i be i.i.d. samples from U . We can then use those samples to approximate μ : the empirical mean is

$$\mu^{\text{emp}}(\theta) = \frac{1}{n_U} \sum_{i=1}^{n_U} J(\theta, u_i) \quad (2.30)$$

and the minimisation problem defined as

$$\min_{\theta \in \Theta} \frac{1}{n_U} \sum_{i=1}^{n_U} J(\theta, u_i) \quad (2.31)$$

is called the *sample average problem* (Juditsky et al., 2009), or *empirical risk minimisation* problem in Machine Learning (see e.g. Vapnik (1992)). Other indicators of central tendency can be considered for optimisation, such as the mode or the median of the cost function.

Despite some similarities with the integrated likelihood introduced Eq. (2.5), θ_{mean} and θ_{intMLE} are not equal in general, as shown Fig. 2.4 for the likelihood introduced in Example 2.2.1 and Fig. 2.2.

A low expected value is to be taken with caution, as it refers to a behaviour *in the long run*. Indeed, the mean value is equivalent to averaging over all the outcomes, but there can be a compensation effect, where “good surprises” balance the “bad surprises”. An example is the following problem:

$$J(\theta_1, U) \sim \mathcal{N}(2, 2^2) \quad (2.32)$$

$$J(\theta_2, U) \sim \mathcal{N}(3, 1^2) \quad (2.33)$$

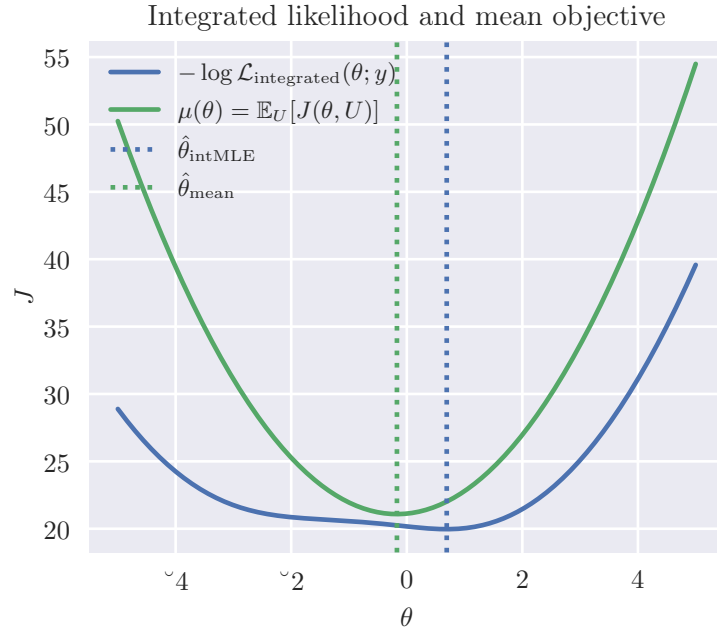


Figure 2.4 – Difference between the negative logarithm of the integrated likelihood defined in Eq. (2.5), and the mean loss of $J = -\log \mathcal{L}$ defined in Eq. (2.26) and the subsequent difference in estimators

and we have to choose either $\theta = \theta_1$ or $\theta = \theta_2$. It is clear that $\mathbb{E}_U[J(\theta_1, U)] < \mathbb{E}_U[J(\theta_2, U)]$. However, making the decision $\theta = \theta_2$ leads to less extreme values:

$$\mathbb{P}_U[J(\theta_1, U) > 5] = 0.06681 > \mathbb{P}_U[J(\theta_2, U) > 5] = 0.02275 \quad (2.34)$$

Depending on the application, such a behaviour could be prohibitive. The difference in these probabilities is explained by the difference in the variance of the random variable $J(\theta, U)$. Accounting for the variance in the objective function is discussed in Section 2.3.2.c.

2.3.2.b Variance optimisation

In Section 2.3.2.a, we used the mean as a measure of the central tendency that we want to minimise. Jointly with the central tendency, information about the dispersion of the random variable may also be relevant, in order to predict how much deviation should be expected around the mean. Let us define the variance of the objective function:

$$\sigma^2(\theta) = \mathbb{V}\text{ar}[J(\theta, U)] \quad (2.35)$$

and minimising this variance yields

$$\theta_{\text{var}} = \min_{\theta \in \Theta} \sigma^2(\theta) = \min_{\theta \in \Theta} \sigma(\theta) \quad (2.36)$$

Depending on the application, the equivalent formulation using the standard deviation may be used instead of the variance. In both formulations, such a computation requires the

evaluation of an expensive integral, this problem can be tackled using sample averaging, and the minimisation problem becomes

$$\min_{\theta \in \Theta} \frac{1}{n_U - 1} \sum_{i=1}^{n_U} (J(\theta, u_i) - \mu^{\text{emp}}(\theta))^2 \quad (2.37)$$

Figure 2.5 shows the conditional mean and conditional standard deviation for the objective function J defined Eq. (2.25).

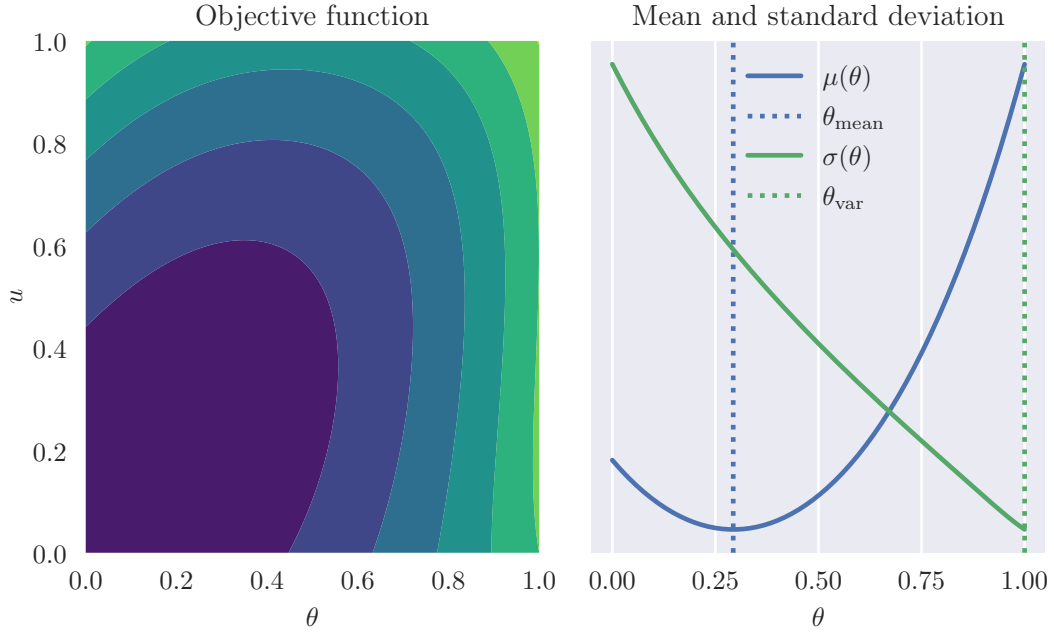


Figure 2.5 – Illustration of conditional mean and conditional standard deviation, as a function of θ . Those quantities have been rescaled to share the same range on the right plot.

2.3.2.c Multiobjective optimisation

Minimising the variance is often irrelevant without additional constraints, as it could just point toward really high values of the objective function, but steady with respect to θ . Taking both objectives: low mean value and low variance together to the following multiobjective optimisation problem:

$$\min_{\theta \in \Theta} (\mu(\theta), \sigma(\theta)) \quad (2.38)$$

This problem can be tackled in different ways using multiobjective optimisation. To compare θ_1 and θ_2 , we can compare component-wise the objective vectors $(\mu(\theta_i), \sigma(\theta_i))$ for $i = 1, 2$. If $\mu(\theta_1) \leq \mu(\theta_2)$ and $\sigma(\theta_1) \leq \sigma(\theta_2)$, θ_2 is said to be *dominated* by θ_1 . The Pareto frontier is defined as the set of points in Θ that cannot be dominated by any other

points. For points on this front, one cannot decrease further one of the objective without increasing the other. On Fig. 2.6 is illustrated the Pareto frontier for the multiobjective problem of Eq. (2.38). The red point corresponding to θ_1 is dominated by the green point θ_0 on the frontier, but not by the green point of θ_2 . A solution of the multiobjective problem can then be chosen within the Pareto frontier.

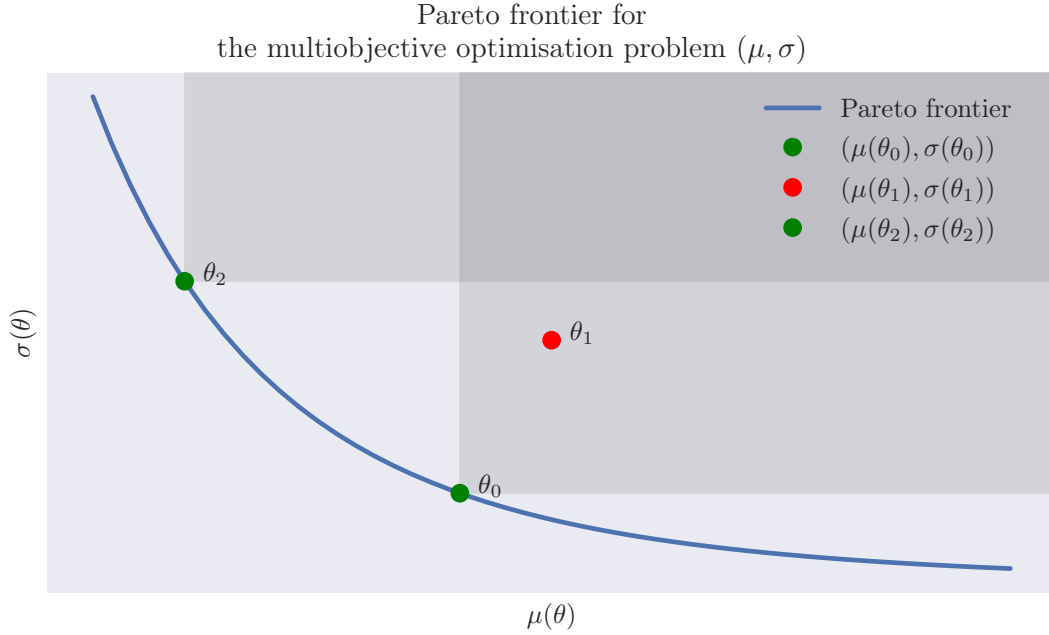


Figure 2.6 – Illustration of the Pareto frontier for the multiobjective problem of Eq. (2.38). The shaded regions corresponds to the domain dominated by each points

Instead of finding the Pareto frontier, the multiobjective problem is often “scalarized” by summing the weighted objectives (Marler and Arora, 2010), provided that such an operation makes sense with regards to the units of the quantities, justifying the use of the standard deviation instead of the variance.

$$\min_{\theta \in \Theta} \lambda \mu(\theta) + (1 - \lambda) \sigma(\theta) = \min_{\theta \in \Theta} \lambda \mathbb{E}_U[J(\theta, U)] + (1 - \lambda) \sqrt{\text{Var}[J(\theta, U)]} \quad (2.39)$$

where $\lambda \in [0, 1]$ is chosen to reflect the preference toward one or another objective.

2.3.2.d Higher moments in optimisation

Higher moments can also be considered as additional criteria, especially in Portfolio optimisation (Lai et al., 2006; Bricc et al., 2007). For a real random variable $X \in L^3$ (with respect to Lebesgue’s measure on $\mathbb{R}$), the skewness coefficient measures the asymmetry in the distribution, and is the (normalized) centered moment of order 3:

$$\text{sk}[X] = \mathbb{E} \left[\left(\frac{X - \mu}{\sigma} \right)^3 \right] \quad (2.40)$$

where $\mu = \mathbb{E}[X]$ and $\sigma = \sqrt{\text{Var}[X]}$.

Adding the skewness in the optimisation translates to a preference toward a risk-averse or a risk-seeking approach. Indeed, as the main goal is the optimisation of an objective function, deviations of the value of the random variable toward lower values is more desirable than deviations toward larger values.

This is illustrated Fig. 2.7: all three of the random variables displayed have the same mean and variance. If the skewness coefficient is negative, the distribution presents a heavier left tail than right. In other words, a sample taken from this distribution has a higher probability of being a “good surprise”. On the other hand, if a big deviation occurs for a sample from a right-skewed distribution, it is more probable to be a large deviation toward large values of the sample space, hence the term “bad surprise”.

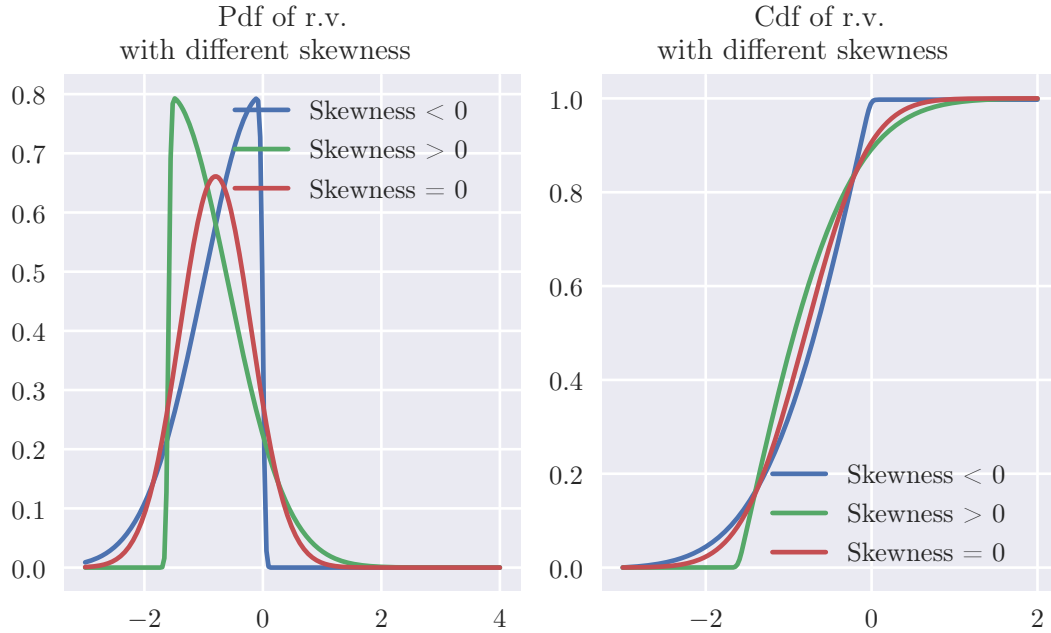


Figure 2.7 – Pdf and cdf of random variables with same mean, variance but different skewness

Other extensions have been developed around the cdf of the random variable, especially in portfolio optimisation. Indeed, integrating and comparing the cdf allows to introduce a *domination order* between random variables. These concepts of Stochastic Dominance (Ogryczak and Ruszczyński, 1997) are then used to take decisions under uncertainties.

2.4 Regret-based families of estimators

All the methods introduced above required first to eliminate in some sense the dependency on the environmental parameter, in order to transform the random variable $J(\theta, U)$ into an objective that depends solely on θ , and to optimise this deterministic

counterpart. For a given $\theta \in \Theta$, this elimination is done by aggregating all the possible outcomes $J(\theta, u)$ when u is a sample of U .

We propose now to reverse these steps, by first optimising the objective function with respect to θ , and, from the set of minimisers that depend on u obtained, derive an estimator. The rationale behind this permutation is that every situation induced by a realisation u is to be taken separately, quite similarly as Savage's regret introduced [Section 2.3.1.c](#). In turn, this avoids aggregation (and in a sense compensation) between the different u .

The work detailed in this section is largely based on [Trappler et al. \(2020\)](#).

2.4.1 Conditional minimum and minimiser

We assume that U is a continuous random variable, with a compact support.

Definition 2.4.1 – Conditional minimum, minimiser: Let $J : \Theta \times \mathbb{U}$ be an objective function, and let us assume that for each $u \in \mathbb{U}$, $\min_{\theta \in \Theta} J(\theta, u)$ exists and is attained at a unique point. We denote

$$J^*(u) = \min_{\theta \in \Theta} J(\theta, u) \quad (2.41)$$

the *conditional minimum* of J given u , and

$$\theta^*(u) = \arg \min_{\theta \in \Theta} J(\theta, u) \quad (2.42)$$

is defined as the *conditional minimiser*

As u is thought to be a realization of a random variable U , we can consider the two random variables $\theta^*(U)$ and $J^*(U)$. The conditional minimum $J^*(U)$ is then a random variable describing the best performances of the calibration, if we could optimise the objective function for each realization of U .

Similarly, let us assume that the conditional minimiser is well defined for all $u \in \mathbb{U}$. We can study the image of the random variable through this mapping, that we will denote $\theta^*(U)$. This random variable in itself gives already information on the “identifiability” of a robust estimate, depending on the information carried by its distribution.

Example 2.4.2: Let $\Theta = \mathbb{U} = [0, 1]$, and $U \sim \text{Unif}(\mathbb{U})$, and the following objective functions:

$$J_1(\theta, u) = (1 + u) + (\theta - 0.5)^2 \quad (2.43)$$

$$J_2(\theta, u) = (\theta - u)^2 + 1 \quad (2.44)$$

We have

$$\theta_1^*(u) = \arg \min_{\theta \in \Theta} J_1(\theta, u) = 0.5 \quad (2.45)$$

$$\theta_2^*(u) = \arg \min_{\theta \in \Theta} J_2(\theta, u) = u \quad (2.46)$$

In the first case, $\theta_0 = \arg \min_{\theta} J(\theta, U)$, so $\theta_1^*(U)$ is a degenerate random variable almost surely equals to θ_0 . In other words, the minimiser is not dependent on the value taken by the environmental parameter. The minimal value attained J^* might be dependent though. On the other hand, for J_2 , as $\Theta = \mathbb{U}$ and $U \sim \text{Unif}(\mathbb{U})$, $\theta_2^*(U)$ is uniformly distributed on Θ , no value shows a better affinity of being a minimiser than the other.

In general, this random variable cannot be classified as continuous or discrete without further study. However, in the following, we are going to assume that it is a *continuous random variable*. The entropy of the random variable $\theta^*(U)$ (see [Definition 1.2.14](#)) can be seen as a measure of the sensitivity of the calibration when the environmental variable varies. If the support of the random variable $\theta^*(U)$ is bounded, the distribution with the highest entropy on this support is a uniform distribution. Per the continuity assumption, this entropy can be estimated by various methods (see for instance [Beirlant et al. \(1997\)](#)).

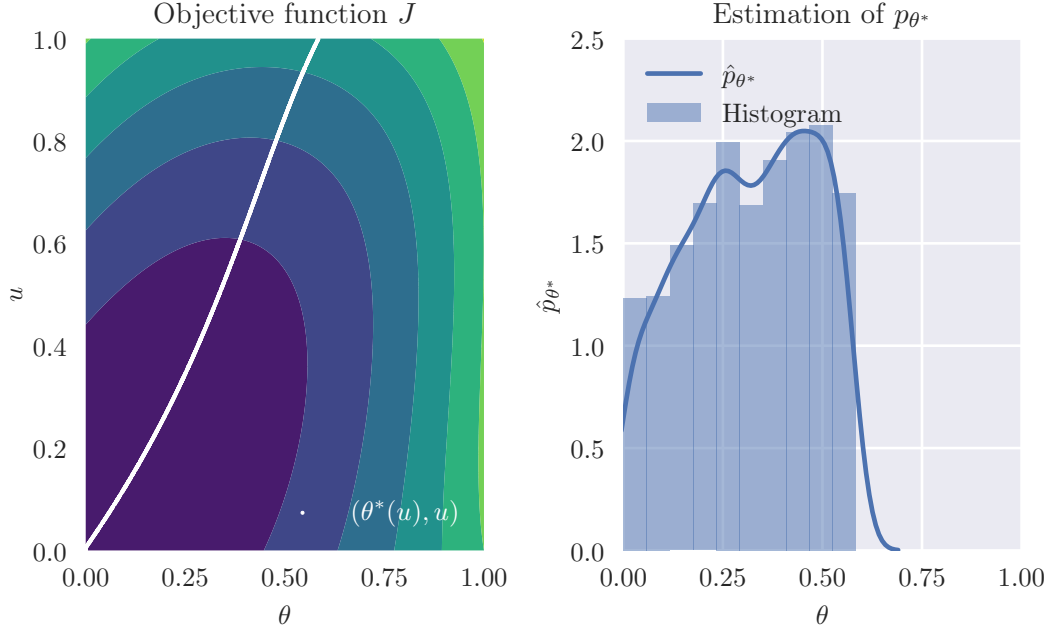
This distribution of the minimisers and its entropy can be used for global optimisation, as outlined in [Hennig and Schuler \(2011\)](#). Furthermore, the authors provide an analytical expression of the pdf of the minimisers, and the nature of the infinite product is discussed:

$$p_{\theta^*}(\theta) = \int_{\mathbb{U}} p_U(u) \prod_{\substack{\tilde{\theta} \in \Theta \\ \tilde{\theta} \neq \theta}} \mathbb{1}_{\{J(\tilde{\theta}, u) > J(\theta, u)\}} du \quad (2.47)$$

However, except for simple analytical problems, the pdf p_{θ^*} cannot be obtained analytically, and needs to be estimated. This can be done by different methods, depending on the assumptions we can make upon $\theta^*(U)$. Let $\{u_i\}_{1 \leq i \leq n_U}$ be n_U i.i.d. samples of U , and $\{\theta^*(u_i)\}_{1 \leq i \leq n_U}$ the corresponding minimisers, as defined [Eq. \(2.42\)](#). Among the methods of density estimation, one of the easiest to implement and widespread method is *Kernel Density Estimation* (KDE). Given the samples u_i and the minimisers $\theta^*(u_i)$ for $1 \leq i \leq n_U$, the isotropic KDE is given by

$$\hat{p}_{\theta^*}(\theta^*) = \frac{1}{n_U h^{\dim \Theta}} \sum_{i=1}^{n_U} \mathcal{K} \left(\frac{\theta^* - \theta^*(u_i)}{h} \right) \quad (2.48)$$

where $h > 0$ is the bandwidth (that measures the influence of each sample), and $\mathcal{K}$ is a kernel of dimension $p = \dim \Theta$, usually defined as the product of one-dimensional kernels $\mathcal{K}_{1D}$: $\mathcal{K}(\theta) = \prod_{j=1}^{\dim \Theta} \mathcal{K}_{1D}(\theta_j)$. Several choices of 1D kernels are available, and one of the most common one is the Gaussian Kernel: $\mathcal{K}_{1D}(\theta_j) = (2\pi)^{-1/2} \exp(-\theta_j^2/2)$. [Figure 2.8](#) shows the estimated density $\hat{p}_{\theta^*}$ using KDE and Scott's rule for the bandwidth ([Scott, 1979](#)), along with the histogram of the minimisers.

Figure 2.8 – Density estimation of the minimisers of J

Finally, when we have an estimation of the density of θ^* , and if it exists, we can compute its mode, that we are going to call the *Most Probable estimator*:

$$\theta_{\text{MPE}} = \arg \max_{\theta \in \Theta} p_{\theta^*}(\theta) \quad (2.49)$$

This mode can be sought directly using appropriate algorithms, such as the Mean-shift algorithm ([Yizong Cheng, 1995](#)), based on the KDE, or clustering methods, such as the Expectation-Maximisation algorithm introduced in [Dempster et al. \(1977\)](#).

Choosing θ_{MPE} means to select the value that is “most often” the minimiser of the objective function. However, we have no indication on its performances when it is *not* optimal, and how often this non-optimality happens. In [Section 2.4.2](#), we are going to introduce the notions of regret (additive and relative), in order to try to be “almost optimal” with high probability.

2.4.2 Regret and model selection

In this section, we will first focus on objective functions defined as the negative log-likelihood in order to link the additive regret and the likelihood ratio test.

2.4.2.a Objective as the negative log-likelihood

Let the objective function be the negative log-likelihood:

$$J(\theta, u) = -\log p_{Y|\theta, U}(y | \theta, u) = -\log \mathcal{L}(\theta, u) \quad (2.50)$$

We can then link the notion of Wald's regret introduced earlier in Eq. (2.23) to the likelihood-ratio test:

$$r(\theta, u_0) = J(\theta, u_0) - \min_{\theta' \in \Theta} J(\theta', u_0) = J(\theta, u_0) - J^*(u_0) \quad (2.51)$$

$$= \max_{\theta' \in \Theta} \log \mathcal{L}(\theta', u_0) - \log \mathcal{L}(\theta, u_0) \quad (2.52)$$

$$= -\log \frac{\mathcal{L}(\theta, u_0)}{\max_{\theta' \in \Theta} \mathcal{L}(\theta', u_0)} = -\log \frac{\max_{\theta' \in \{\theta\}} \mathcal{L}(\theta', u_0)}{\max_{\theta' \in \Theta} \mathcal{L}(\theta', u_0)} \quad (2.53)$$

As the model $(\mathcal{M}(\cdot, u_0), \Theta)$ is misspecified, we can apply the misspecified likelihood-ratio test, and for a candidate $\theta \in \Theta$ we can test for the following hypotheses:

- $\mathcal{H}_0$: the model $(\mathcal{M}(\cdot, u_0), \{\theta\})$ is statistically equivalent to $\{\mathcal{M}(\cdot, u_0), \Theta\}$
- $\mathcal{H}_1$: the models are statistically different

The statistic of the test, *i.e.* the regret r is to be compared with half the quantile of the r.v. $X(u_0)$ defined as

$$X(u_0) = \sum_{i=1}^{\dim \Theta} c_i(u_0) \Xi_i \quad \text{with } \Xi_i \sim \chi_1^2 \text{ i.i.d.} \quad (2.54)$$

where $\{c_i(u_0)\}_{1 \leq i \leq \dim \Theta}$ are coefficients linked to the eigenvalues of the Fisher information matrix as evoked in Section 1.6.

The null hypothesis $\mathcal{H}_0$ is rejected at a level $\eta \in]0; 1[$ if

$$r(\theta, u_0) = J(\theta, u_0) - J^*(u_0) > \beta \quad (2.55)$$

Where β is half the $1 - \eta$ quantile of the random variable $X(u_0)$ defined Eq. (2.54).

Using this rejection region, we can construct a likelihood interval (as defined Eq. (1.73)), which depends on u_0 :

$$\mathcal{I}_{\text{Lik}}(u_0; \beta) = \{\theta \in \Theta \mid J(\theta, u_0) - J^*(u_0) \leq \beta\} \quad (2.56)$$

So, for $\theta \in \mathcal{I}_{\text{Lik}}(u_0; \beta)$, the model $(\mathcal{M}(\cdot, u_0), \{\theta\})$ is acceptable at the η -level, per the likelihood-ratio test.

From a computational point of view, the coefficients $\{c_i(u_0)\}$ are hard to obtain, and depend on u_0 . A first approximation would be to suppose that $X(u_0) \sim \chi_{\dim \Theta}^2$, *i.e.* to apply the “well-specified” likelihood-ratio test. In the more general case, we can choose $\beta > 0$ in a more arbitrary way in order to avoid the computations of the coefficients $\{c_i(u_0)\}_i$, as we are going to see Section 2.4.2.b.

2.4.2.b Interval and probability of acceptability

We assumed before that J was the negative log-likelihood. In the more general case, J represents a loss function, that we want to minimise. The generalization of Eq. (2.56) is what we are calling the *interval of acceptability*.

Definition 2.4.3 – Interval of acceptability: By analogy with the likelihood interval defined Eq. (1.73) and Eq. (2.56), we can construct a set for an arbitrary threshold $\beta \geq 0$ such that

$$\mathcal{I}_\beta(u) = \{\theta \in \Theta \mid J(\theta, u) \leq J^*(u) + \beta\} \quad (2.57)$$

As J may not stem from a likelihood, we call $\mathcal{I}_\beta(u)$ the *interval of acceptability*. In other words, we say that $\theta \in \Theta$ is β -acceptable for $U = u$ if $\theta \in \mathcal{I}_\beta(u)$.

Figure 2.9 shows an objective function evaluated for different fixed u_i . The β -acceptable intervals for those environmental variables are plotted below the curves.

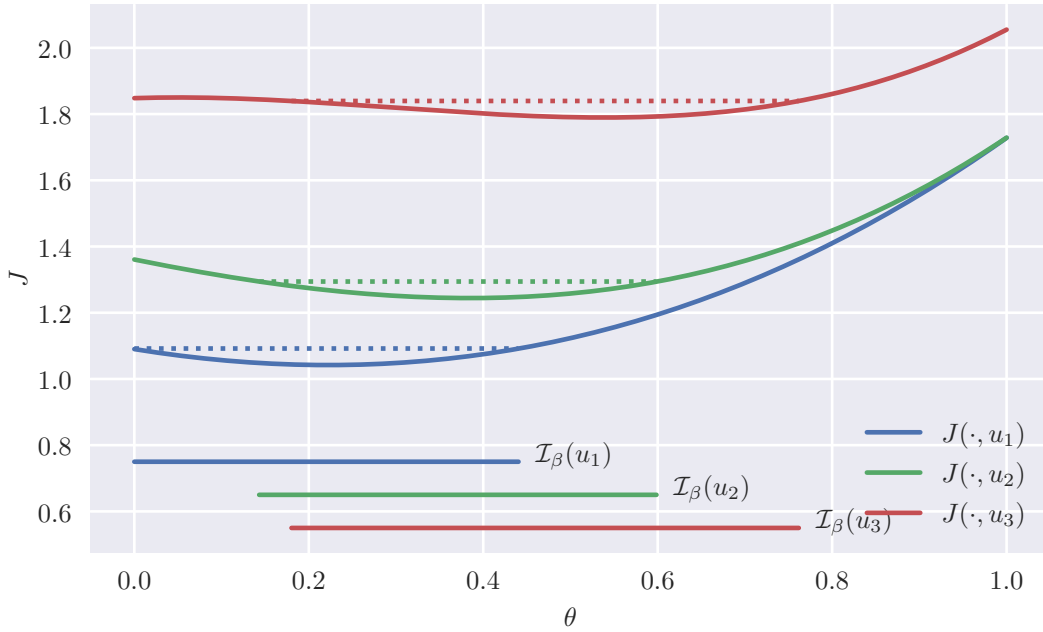


Figure 2.9 – Different acceptable regions corresponding to different $u \in \mathbb{U}$

Now, for a given $\theta \in \Theta$, we can define the set of $u \in \mathbb{U}$ such that $\theta \in \mathcal{I}_\beta(u)$, i.e.

$$R_\beta(\theta) = \{u \in \mathbb{U} \mid \theta \in \mathcal{I}_\beta(u)\} \quad (2.58)$$

$$= \{u \in \mathbb{U} \mid J(\theta, u) \leq J^*(u) + \beta\} \quad (2.59)$$

This set is measurable, and by measuring this subset of $\mathbb{U}$ with respect to the distribution of U , we get

$$\Gamma_\beta(\theta) = \mathbb{P}_U[R_\beta(\theta)] \quad (2.60)$$

Loosely speaking, for a given θ , $\Gamma_\beta(\theta)$ is the probability that the model $(\mathcal{M}(\cdot, U), \{\theta\})$ is “statistically equivalent” to the “full model” $(\mathcal{M}(\cdot, U), \Theta)$, at a certain level linked to

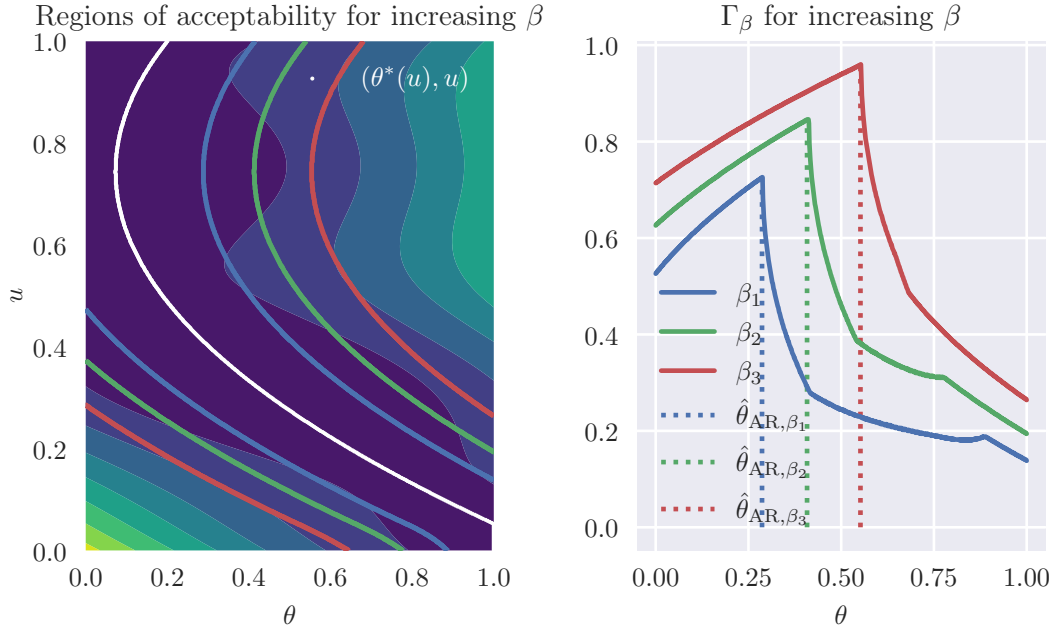


Figure 2.10 – Boundaries of regions of acceptability for increasing β , and Γ_β . The coloured lines on the left plot are the boundaries of those regions

the value of β . Figure 2.10 shows the regions of acceptability for different β , and the associated Γ_β .

Γ_β is then maximised with respect to its argument, in order to get the value θ which has the highest probability of being acceptable at the level β . For different $\beta \geq 0$, we can define the family of additive-regret estimators.

Definition 2.4.4 – Additive-regret family of estimators: For $\beta \geq 0$, we define the family of robust estimators as the maximisers of Eq. (2.60):

$$\left\{ \theta_{\text{AR},\beta} = \arg \max_{\theta \in \Theta} \Gamma_\beta(\theta) \mid \beta > 0 \right\} \quad (2.61)$$

Among this family of estimators, we can then choose a particular value, either by setting a threshold β arbitrarily, or by choosing it so that the probability of being acceptable $\max \Gamma_\beta$ reaches a particular value. This will be discussed later in Section 2.4.4.

2.4.3 Relative-regret

2.4.3.a Absolute and relative error

We examined before regret that can be qualified as *additive* as this is the difference between J and J^* that is compared to fixed thresholds. However, we can argue that the relative magnitude of the objective function has an importance in the comparison. For

illustration purposes, let us consider the situation described Table 2.2, with $\Theta = \{\theta_1, \theta_2\}$ and $\mathbb{U} = \{u_1, u_2\}$, and $\mathbb{P}[U = u_1] = \mathbb{P}[U = u_2] = 1/2$.

J	u_1	u_2	$\mathbb{E}[J(\cdot, U)]$	$J - J^*$	u_1	u_2	$\frac{J - J^*}{J^*}$	u_1	u_2
θ_1	10000	110	5055	θ_1	0	100	θ_1	0	10
θ_2	10100	10	5055	θ_2	100	0	θ_2	0.01	0

Table 2.2 – Illustration of an objective function, expected loss, additive regret and relative error

In this situation, for both u_1 and u_2 , the maximal additive regret is $\max_{\theta} J(\theta, u) - J^*(u) = 100$, so no clear preference could be inferred toward one or another value. However, choosing θ_1 over θ_2 means to choose to improve the performance of an already pretty bad situation (10000 instead of 10100), while increasing tenfold the loss for the situation $U = u_2$.

From the example developed Table 2.2, an alternative to the additive regret may be considered, as the difference in magnitude of the objective function between $U = u_1$ and $U = u_2$ is probably due to the effect of a misspecification. So to take into account this difference, we are now going to consider the *relative regret* J/J^* , instead of the *absolute regret* $J - J^*$, and derive a family of estimators in a similar fashion.

2.4.3.b Relative-regret estimators family

Analogously as the additive regret defined before, we are going first to define the notions of acceptability in the case of the relative regret.

Definition 2.4.5 – α -acceptability: Let $\alpha \geq 1$. A point (θ, u) is said to be α -acceptable if $J(\theta, u) \leq \alpha J^*(u)$. We define the α -acceptable interval as

$$\mathcal{I}_{\alpha}(u) = \{\theta \in \Theta \mid J(\theta, u) \leq \alpha J^*(u)\} \quad (2.62)$$

Then, for a given α and θ , we can define the set of $u \in \mathbb{U}$ such that θ is α -acceptable:

$$R_{\alpha}(\theta) = \{u \in \mathbb{U} \mid \theta \in \mathcal{I}_{\alpha}(u)\} = \{u \in \mathbb{U} \mid J(\theta, u) \leq \alpha J^*(u)\} \quad (2.63)$$

$R_{\alpha}(\theta)$ is a measurable subset of $\mathbb{U}$, and by integrating this set with respect to $\mathbb{P}_U$, we get

$$\Gamma_{\alpha}(\theta) = \mathbb{P}_U[R_{\alpha}(\theta)] \quad (2.64)$$

the probability of being α -acceptable. Using this function, we can define an estimator as the value which maximises this probability. And by varying the threshold α , we get the family of relative-regret estimators.

Definition 2.4.6 – Relative-regret family of estimators: Given α , the value of θ that maximises the probability of being α -acceptable is called the relative-regret (RR) estimator $\theta_{\text{RR}, \alpha}$.

We define the family of relative regret estimators as the set of those estimators:

$$\left\{ \theta_{\text{RR},\alpha} = \arg \max_{\theta \in \Theta} \Gamma_{\alpha}(\theta) \mid \alpha > 1 \right\} \quad (2.65)$$

For different increasing α , the corresponding regions of acceptability are represented in Fig. 2.11, along with the functions Γ_{α} .

Among those estimators and the associated quantities, two limiting cases appear. One particular choice is to set α to 1. In this case, we have

$$\mathcal{I}_1(u) = \{\theta^*(u)\} \quad (2.66)$$

$$R_1(\theta) = \{u \in \mathbb{U} \mid J(\theta, u) = J^*(u)\} = \{u \in \mathbb{U} \mid \theta = \theta^*(u)\} \quad (2.67)$$

We have then that $\Gamma_1(\theta)$ is non-zero if the set $R_1(\theta)$ has non-zero measure with respect to $\mathbb{P}_U$. In other words, $\Gamma_1(\theta)$ is non-zero if θ is the minimiser of $J(\cdot, u)$ for a non-negligible subset of $\mathbb{U}$. If we consider that Θ is a discrete space (due to a discretization for instance), $\theta^*(U)$ is a discrete random variable. $\Gamma_1(\theta)$ is then the probability mass function (discrete parallel of the pdf) of the discrete r.v. $\theta^*(U)$. In this case, $\theta_{\text{RR},\alpha=1} = \theta_{\text{MPE}}$. A similar argument can be made for the additive regret and $\beta = 0$.

We can see that the thresholds act like a relaxation of the optimality condition, as for $\alpha = 1$ and $\beta = 0$, we are measuring the probability of being optimal, while increasing those values means to measure the probability of being nearly optimal.

Another choice is to set the threshold large enough so that the probability of being acceptable reaches a unique maximum, being 1. In this situation, the regret is bounded almost surely. Let β_{inf} and α_{inf} , which verify

$$\beta_{\text{inf}} = \inf \{\beta \geq 0 \mid \max \Gamma_{\beta} = 1\} \text{ and } \alpha_{\text{inf}} = \inf \{\alpha \geq 1 \mid \max \Gamma_{\alpha} = 1\} \quad (2.68)$$

To put it differently, α_{inf} and β_{inf} are the smallest thresholds where there exists a value in Θ acceptable almost surely. This value shares similarities with the minimiser of Savage's regret: θ_{WC} introduced Section 2.3.1.c, as it minimises almost surely (*i.e.* for all u in a non-negligible set) the regret, either additive or relative.

For both regret-based approaches, the interval of acceptability at a fixed $u \in \mathbb{U}$ grows with the threshold, and the sharper the minimum is (in the sense of a large curvature), the faster it grows. A more telling illustration of the differences between relative and additive regret is shown Fig. 2.12. The regions of acceptability of an objective function J have been plotted, along with an interval $\mathcal{I}(u)$ and the region R for a given θ for both regrets. For u around 0, J is quite flat, but also has very low values. In this case, the interval $\mathcal{I}_{\beta}(u)$ is also large. But for u around 1, the objective function presents higher values, and a sharper minimum (*i.e.* a higher curvature). Additive regret in this case puts stronger confidence on the value of the parameter as indicated by the smaller interval of acceptability.

For the relative regret, the situation is reversed. Although sharp, the large value attained by the minimum $J^*(u)$, for u around 1 leads to a large interval of acceptability,

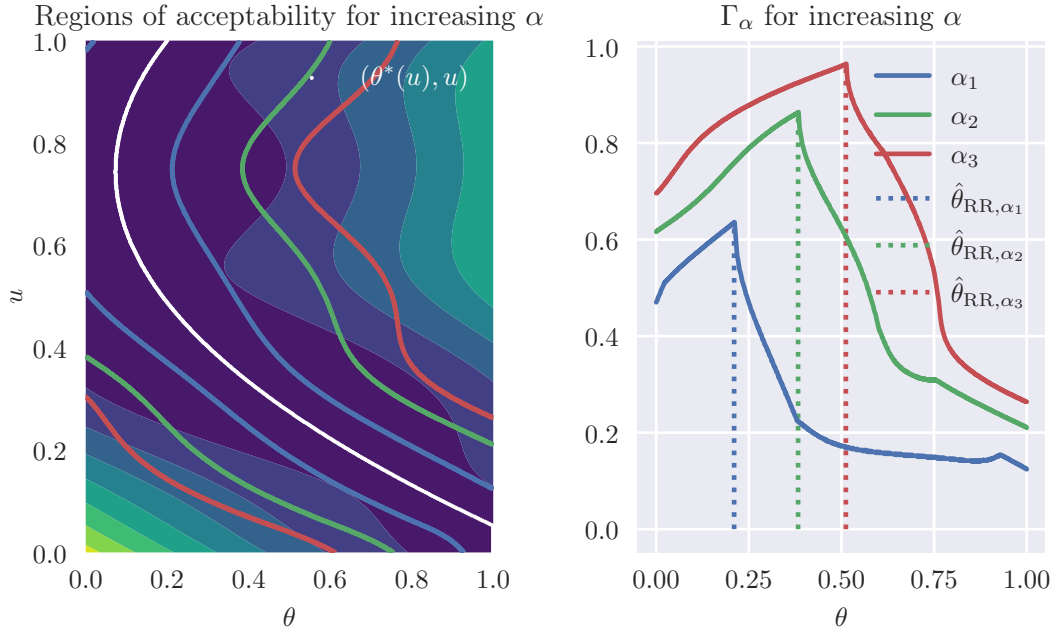


Figure 2.11 – Regions of acceptability for relative regret and increasing α . The coloured lines on the left plot are the boundaries of those regions

meaning that we can deviate a bit more from this minimum, as the situation $u \approx 1$ is already pretty bad. For $u \approx 0$, which is close to the global minimum of J , the interval is smaller, as this criterion does not favour a large deviation from this global minimum.

2.4.4 The choice of the threshold

We are now going to focus exclusively on the relative-regret, and the associated threshold α , but similar arguments can be made for the absolute regret and the threshold β .

In order to have an insight on the potential robustness of an estimator $\theta_{RR, \alpha}$ both values can be studied together: the threshold α and the maximal probability reached $\max \Gamma_\alpha$. Finding a relevant threshold can be a thorny issue, especially with no further information on J . Setting it too large will lead to large α -acceptable intervals, and Γ_α may reach 1 for several different values. On the other hand, choosing a threshold too small may give a maximal acceptability probability too low to assess the robustness of the chosen solution. We can contemplate three starting points:

- Set a relaxation parameter $\alpha > 1$. From this, the maximal probability is $p_\alpha = \max_{\theta \in \Theta} \Gamma_\alpha = \max_{\theta \in \Theta} \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)]$. This task can be seen as the estimation and optimisation of a specific probability.
- Set a probability $0 < p \leq 1$ we want to reach, and define $\alpha_p = \inf_{\alpha \geq 1} \{\max_{\theta \in \Theta} \Gamma_\alpha(\theta) \geq p\}$, so the problem can be thought as the estimation and the optimisation of quantile of a particular random variable: we can define α_p as the solution of the following

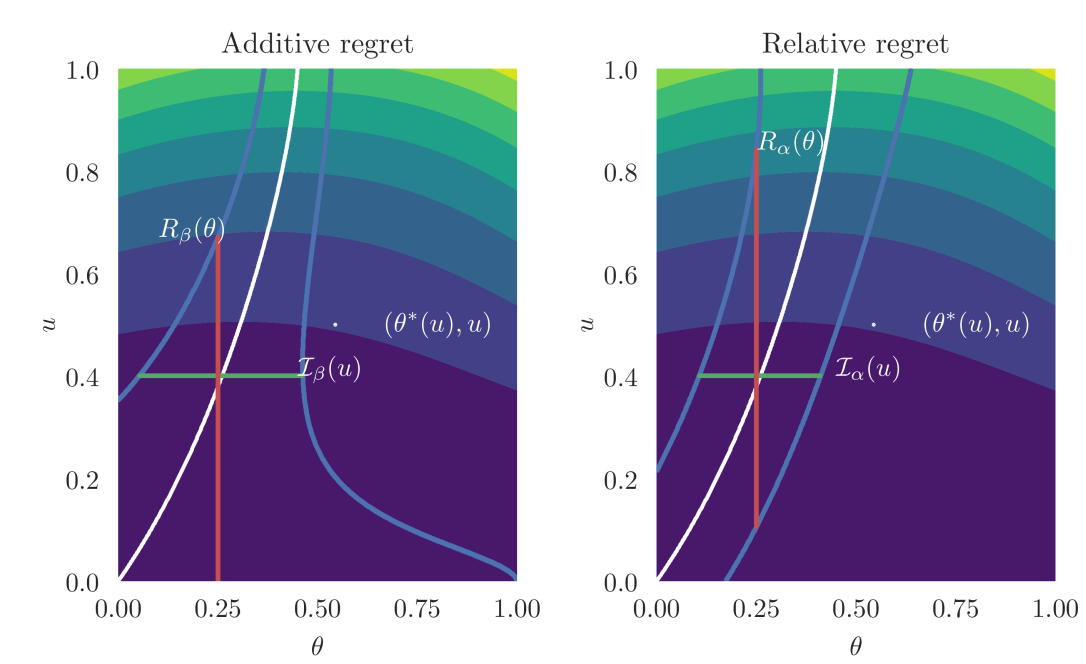


Figure 2.12 – Comparison of the regions of acceptability for additive and relative regret

chance constrained problem:

$$\begin{cases} \min_{\theta \in \Theta} & q(\theta) \\ \text{such that} & \mathbb{P}_U[r(\theta, U) \leq q(\theta)] \geq p \end{cases} \quad (2.69)$$

where $r(\theta, u) = J(\theta, u)/J^*(u)$ or $r(\theta, u) = J(\theta, u) - J^*(u)$ depending on the context, and $q(\theta)$ is the quantile of order p of the r.v. $r(\theta, U)$.

- Study the evolution of one quantity with respect to the other, in order to find a balance between the probability of being acceptable, and the relaxation needed to reach it.

Specific techniques may be applied in order to perform the first two approaches efficiently, which will be introduced in the next chapter. The third one however requires the knowledge of the objective function on the whole joint space $\Theta \times \mathbb{U}$, in order to compute $\max \Gamma_\alpha$ for a various number of thresholds, thus may not be adapted for costly computer simulations.

2.5 Partial Conclusion

As shown throughout this chapter, when optimising under uncertainties, a lot of criteria can be defined in order to satisfy the idea of *robustness*, depending on the interpretation of this term. A summary of those introduced here can be found [Table 2.3](#). Some criteria are commonly encountered in optimisation under uncertainty, such as

expected loss minimisation. We introduced also new families of estimators, which aim at maximising a probability based on the regret, either additive or relative.

Objective name	Objective to <i>minimise</i> with respect to θ
Profile Likelihood	$-\log (\max_{u \in \mathbb{U}} p_{Y \theta, U}(y \theta, u))$
Integrated Likelihood	$-\log \int_{\mathbb{U}} p_{Y \theta, U}(y \theta, u) p_U(u) du = -\log p_{Y \theta}(y \theta)$
Marginal maximum a posteriori	$-\log p_{\theta Y}(\theta y)$
Global Optimum	$\min_{u \in \mathbb{U}} J(\theta, u)$
Worst-case	$\max_{u \in \mathbb{U}} J(\theta, u)$
Regret worst-case	$\max_{u \in \mathbb{U}} \{J(\theta, u) - \min_{\theta' \in \Theta} J(\theta', u)\}$
Mean	$\mathbb{E}_U[J(\theta, U)]$
Mean and variance	$\lambda \mathbb{E}_U[J(\theta, U)] + (1 - \lambda) \sqrt{\text{Var}_U[J(\theta, U)]}$
Most Probable Estimate	$-\log p_{\theta^*}(\theta)$
Additive-regret	$-\mathbb{P}_U [J(\theta, U) \leq J^*(U) + \beta] = -\Gamma_{\beta}(\theta)$
Relative-regret	$-\mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)] = -\Gamma_{\alpha}(\theta)$

Table 2.3 – Summary of single objective robust estimators

Obviously, other criteria can be defined, that satisfy other robustness requirements. Furthermore, we did not treat the possibility of combining some of those objectives by using them to set constraints. An example is the minimisation of the variance, under the constraint that the mean value does not exceed a certain threshold T as in [Lehman et al. \(2004\)](#):

$$\begin{aligned} \min \text{Var}_U [J(\theta, U)] \\ \text{s.t. } \mathbb{E}_U [J(\theta, U)] \leq T \end{aligned}$$

All of the criteria introduced above require costly numerical procedures, such as integration and optimisation. Solving these robust estimation problems is then expensive in term of computer resources, as one would need to run the forward model a very large number of times in order to get accurate numerical integration or optimisation. In the next chapter we will discuss methods based on surrogate modelling, that can be used to solve efficiently such problems, in order to make the best of the evaluations of the numerical model $\mathcal{M}$ on the space $\Theta \times \mathbb{U}$.

* * *

CHAPTER 3

ADAPTIVE STRATEGIES FOR CALIBRATION USING GAUSSIAN PROCESSES

Contents

3.1	Introduction	63
3.2	Gaussian process regression	63
3.2.1	Random processes	64
3.2.2	Kriging equations	65
3.2.3	Covariance functions	67
3.2.4	Initial design and validation	67
3.3	General enrichment strategies for Gaussian Processes	68
3.3.1	1-step lookahead strategies	69
3.3.2	Batch selection of points: sampling-based methods	70
3.4	Criteria of enrichment	70
3.4.1	Criteria for exploration of the input space	71
3.4.2	Criteria for optimisation of the objective function	72
3.4.3	Contour and volume estimation	75
3.4.3.a	Reduction of the augmented IVPC	75
3.4.3.b	Sampling in the Margin of uncertainty	77
3.5	Adaptive strategies for robust optimisation using GP	78
3.5.1	PEI for the conditional minimisers	79
3.5.2	Gaussian formulations for the relative and additive regret families of estimators	80
3.5.3	GP-based methods for the estimation of Γ_α	83

3.5.3.a	Estimation of Γ_α	83
3.5.3.b	Reduction of the augmented IMSE	83
3.5.3.c	Reduction of the augmented IVPC	84
3.5.3.d	Sampling in the margin of uncertainty	86
3.5.4	Optimisation of the quantile of the relative regret	92
3.5.4.a	Lognormal approximation of the ratio	93
3.5.4.b	Reduction of the IMSE	94
3.5.4.c	Sampling-based method, adaptation of the QeAK-MCS	95
3.6	Partial conclusion	96

3.1 Introduction

In the previous chapter, we introduced different notions of robustness for the estimation of a calibration parameter. Those objectives usually require the evaluation of various integrals, and to perform several optimisations. A large number of model evaluations is then needed, but it may bring some practical issues, as numerical models are usually very expensive to run in terms of computational resources. Indeed, for most realistic physical simulations, systems of PDEs over large discretized domains have to be solved. Even though the programs are optimised and parallelized to take best advantage of high-performance computers, the time required to compute the quantities of interest may range from a few seconds to days. Because of that, methods requiring a large number of runs of the model for exhaustivity should be avoided.

In this chapter, we will focus on the use of *surrogate models* to solve robust optimisation problems, according to some of the criteria introduced in the previous chapter.

Definition 3.1.1 – Surrogate function: Let $f : \mathbb{X} \rightarrow \mathbb{R}$ be a function representing the computation of a quantity of interest. A *surrogate*, or *metamodel*, or *emulator* of f , say g , is a function from $\mathbb{X}$ to $\mathbb{R}$ which possesses two main properties:

- g is an approximation of f
- g is cheaper to evaluate than f

We will focus exclusively on kriging (Krige, 1951; Matheron, 1962), also called Gaussian Process regression, but other methods can be used to solve such problems. Polynomial Chaos regression for instance Wiener (1938); Xiu and Karniadakis (2002); Sudret (2015); Miranda et al. (2016).

Those surrogates then be used directly instead of f in a *plug-in* approach.

Definition 3.1.2 – Plug-in method: In this work, we will use the term *plug-in* as the fact of using directly an estimated quantity, such as a surrogate g instead of the expensive-to-evaluate original function f in computations.

In this chapter, after defining the usual kriging equations in Section 3.2.2, we are going to introduce a few classic and useful criteria in Section 3.3 for global optimisation and/or exploration of an unknown function f . Finally, in Section 3.5, we are going to focus on robust optimisation, by splitting the input space in Θ and $\mathbb{U}$, and introduce other strategies to take advantage of the nature of the surrogate in order to efficiently compute the robust estimates introduced before, especially regret-based estimators.

3.2 Gaussian process regression

In the following, we will consider a generic function f , that maps a space $\mathbb{X}$ to $\mathbb{R}$. Depending on the application, $\mathbb{X} = \Theta$ or $\mathbb{X} = \Theta \times \mathbb{U}$. This function is unknown, and expensive to evaluate.

3.2.1 Random processes

Let us assume that we have a map f from a p -dimensional space to $\mathbb{R}$:

$$\begin{aligned} f : \mathbb{X} \subset \mathbb{R}^p &\longrightarrow \mathbb{R} \\ x &\longmapsto f(x) \end{aligned} \quad (3.1)$$

This function is assumed to have been evaluated on a design of n points, $\mathcal{X} = \{(x_i, f(x_i))\}_{1 \leq i \leq n}$, called the *initial design*. For notational simplicity, we write $x \in \mathcal{X}$ if $(x, f(x)) \in \mathcal{X}$, *i.e.* if the point x has been evaluated by f . As this function is unknown, there is uncertainty on the values outside the initial design, which can be classified as *epistemic* since it can be reduced by directly evaluating the function (see [Section 2.1](#)). This uncertainty on the value taken by the function leads us to the definition of random processes:

Definition 3.2.1 – Random process: Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a probability space, and $\mathbb{X} \subset \mathbb{R}^p$. A random process Z is a collection of random variables indexed on $\mathbb{X}$, so for each $x \in \mathbb{X}$, $Z(x)$ is a real random variable (*i.e.* $Z(x) : \Omega \rightarrow \mathbb{R}$):

$$\begin{aligned} Z : \mathbb{X} &\longrightarrow (\Omega \rightarrow \mathbb{R}) \\ x &\longmapsto Z(x) \end{aligned} \quad (3.2)$$

A sample from this random process, that is $Z(\cdot)(\omega)$ for $\omega \in \Omega$ will be shortened as $Z(\cdot, \omega)$ for notational purpose, and is called a *sample trajectory*, or a *sample path*. When ω is omitted, $Z(x)$ represents the random process at the point $x \in \mathbb{X}$.

From a Bayesian point of view, such a random process can act as a prior on the function f , or in other words, f can be thought as a particular sample path of Z . Evaluating the function at an additional point $x \notin \mathcal{X}$ provides new information on the random process, and we can update our belief on f .

In this work, we are going to focus exclusively on a specific type of random process, namely the Gaussian process (abbreviated as GP). Other types of random process can be encountered in the literature: Student t-processes in [Shah et al. \(2014\)](#) are introduced as alternatives to GP to account for larger tails, or various graphical models such as Gaussian and Markov Random Fields in [Bishop \(2006\)](#); [Li \(2009\)](#) are used to model images for instance.

Definition 3.2.2 – Gaussian process: Let Z be a random process on $\mathbb{X}$, *i.e.* a collection of random variables indexed by $\mathbb{X}$. Z is a Gaussian process (GP) if any finite number of those random variables have a multivariate joint Gaussian distribution. In that case, Z is uniquely defined by its mean function $m_Z : \mathbb{X} \rightarrow \mathbb{R}$ and its covariance function $C_Z : \mathbb{X} \times \mathbb{X} \rightarrow \mathbb{R}$:

$$m_Z(x) = \mathbb{E}[Z(x)] \quad (3.3)$$

$$C_Z(x, x') = \text{Cov}[Z(x), Z(x')] \quad (3.4)$$

and we write $Z \sim \text{GP}(m_Z, C_Z)$. Due to the definition of a GP, we also have for all $x \in \mathbb{X}$

$$Z(x) \sim \mathcal{N}(m_Z(x), \sigma_Z^2(x)) \quad (3.5)$$

with $\sigma_Z^2(x) = C_Z(x, x)$

A more thorough description of Gaussian processes and their applications can be found in [Rasmussen and Williams \(2006\)](#).

Let Z be a GP with a known covariance function. The construction of a covariance function will be discussed [Section 3.2.3](#). Based on the initial design $\mathcal{X}$, we can construct a surrogate for the unknown function f by conditioning the GP on the design which comprises some evaluations of f .

3.2.2 Kriging equations

Given a Gaussian process Z as a prior on f , and by conditioning it by the initial design $\mathcal{X} = \{(x_i, f(x_i))\}_{1 \leq i \leq n}$, the conditioned random process is still a GP:

$$Z \mid \mathcal{X} \sim \text{GP}(m_{Z|\mathcal{X}}, C_{Z|\mathcal{X}}) \quad (3.6)$$

Given the nature of Z and the initial design, we can derive the joint distribution of the GP at the points of the design and the GP at an unobserved point $x \in \mathbb{X}$:

$$\begin{pmatrix} Z(\mathbf{x}) \\ Z(x) \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} \mu_Z \\ m_Z(x) \end{pmatrix}; \begin{pmatrix} \mathbf{K}_{\mathcal{X}} & K_{\mathcal{X}}(x) \\ K_{\mathcal{X}}(x)^T & C_Z(x, x) \end{pmatrix} \right) \quad (3.7)$$

where $\mathbf{x} = (x_1, \dots, x_n)$, $x_i \in \mathcal{X}$, for $1 \leq i \leq n$ are the points of the design and

$$Z(\mathbf{x}) = (Z(x_1), \dots, Z(x_n)) \quad (3.8)$$

$$K_{\mathcal{X}}(x) = (C_Z(x, x_1), C_Z(x, x_2), \dots, C_Z(x, x_n))^T \quad (3.9)$$

$$\mathbf{K}_{\mathcal{X}} = (C_Z(x_i, x_j))_{1 \leq i, j \leq n} \quad (3.10)$$

and $\mu_Z = \mathbb{E}[Z(\mathbf{x})]$ is the mean of the unconditioned GP. When this mean is assumed to be known, the kriging procedure is qualified of *Simple Kriging*, otherwise, we often talk about *Ordinary Kriging*. Finally, when this mean is a deterministic function (that may be estimated), we talk about *Universal Kriging* ([Le Riche, 2014](#)). In the following, we consider Simple Kriging.

Using the properties of the multivariate distribution, the GP conditioned on the observations of [Eq. \(3.6\)](#) has mean and covariance function defined by

$$m_{Z|\mathcal{X}}(x) = m_Z(x) + K_{\mathcal{X}}(x)^T \mathbf{K}_{\mathcal{X}}^{-1} (f(\mathbf{x}) - \mu_Z) \quad (3.11)$$

$$C_{Z|\mathcal{X}}(x, x') = C_Z(x, x') - K_{\mathcal{X}}(x)^T \mathbf{K}_{\mathcal{X}}^{-1} K_{\mathcal{X}}(x') \quad (3.12)$$

These are called the *kriging equations*.

Given the fact that a conditioned GP is still a GP, at a point $x \in \mathbb{X}$, we have

$$Z(x) \mid \mathcal{X} \sim \mathcal{N} \left(m_{Z|\mathcal{X}}(x), \sigma_{Z|\mathcal{X}}^2(x) \right) \quad \text{with} \quad \sigma_{Z|\mathcal{X}}^2(x) = C_{Z|\mathcal{X}}(x, x) \quad (3.13)$$

The function $m_{Z|\mathcal{X}} : \mathbb{X} \rightarrow \mathbb{R}$ can then be defined as a surrogate for f . This surrogate provides an interpolation of the unknown function f , since for $x \in \mathcal{X}$, $f(x) = m_{Z|\mathcal{X}}(x)$. In addition to that interpolation property, for points not in the design, we have a measure of the epistemic uncertainty modelled by the normal r.v. in Eq. (3.13). As an illustration, for a few training points and their respective evaluations, we represented the GP regression of an unknown function f Fig. 3.1.

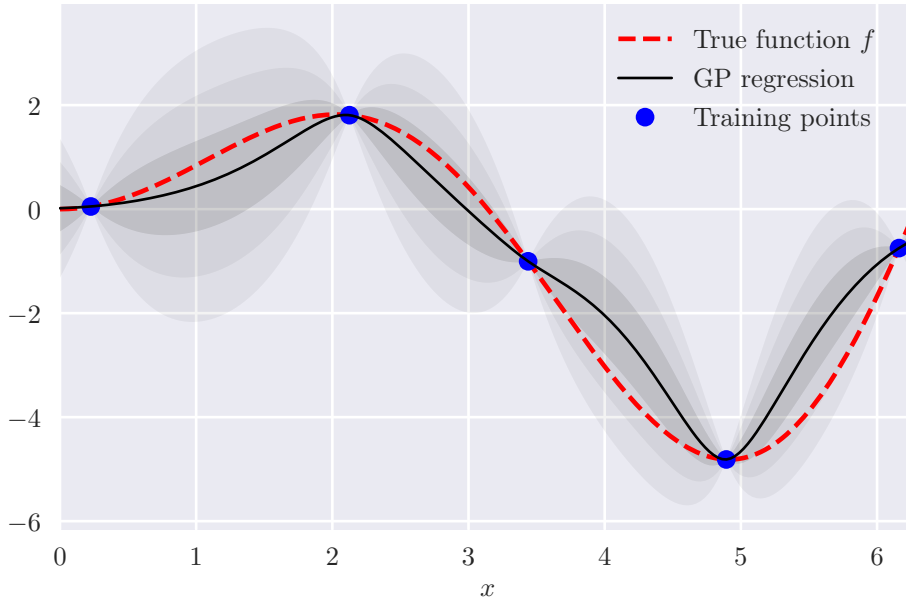


Figure 3.1 – Example of a Gaussian Process, as a surrogate for a function f evaluated at 4 inputs. The shaded regions correspond to the regions $m_Z \pm i \cdot \sigma_Z$ for $i = 1, 2, 3$.

From a computational point of view, GP are relatively cheap to handle:

- If the covariance function C_Z is known, the most expensive step in the estimation of $m_{Z|\mathcal{X}}$ and $C_{Z|\mathcal{X}}$ is the inversion of the matrix $\mathbf{K}_{\mathcal{X}}$. This matrix does not depend on the points x where we wish to evaluate $m_{Z|\mathcal{X}}$, thus its inversion needs only to be performed once. However, if there is a very large number of points in the design, and thus a very large matrix $\mathbf{K}_{\mathcal{X}}$, this inversion can reveal itself numerically complicated in practice. This renders GP suited for problems involving only a moderate number of dimensions
- In order to get a sample path from this GP, one needs to be able to sample from a multivariate normal distribution, and thus only need to compute the Cholesky decomposition of the conditioned covariance matrix.

However, if some points of the design are very close to each other, the covariance matrix can be ill-conditioned. In this case, for numerical stability, we can add a *nugget* effect to

$\mathbf{K}_{\mathcal{X}}$, and Eq. (3.14) becomes:

$$\begin{pmatrix} Z(\mathbf{x}) \\ Z(x) \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} \mu_Z \\ m_Z(x) \end{pmatrix}; \begin{pmatrix} \mathbf{K}_{\mathcal{X}} + \varsigma^2 \text{Id} & K_{\mathcal{X}}(x) \\ K_{\mathcal{X}}(x)^T & C_Z(x, x) \end{pmatrix} \right) \quad (3.14)$$

From a probabilistic point of view, we can see the nugget effect as an added Gaussian noise when evaluating the function, *i.e.* that we are observing $f(x_i) + \epsilon$ with $\epsilon \sim \mathcal{N}(0, \varsigma^2)$. When $\varsigma^2 \neq 0$, the GP is no longer interpolant, as $\text{Var}[Z(x) \mid \mathcal{X}] \neq 0$ for $x \in \mathcal{X}$.

3.2.3 Covariance functions

In the previous section, we described the equations to solve to get the surrogate $m_{Z|\mathcal{X}}$ of f based on GP. The kriging equations are based on the covariance function C_Z . A covariance is said to be stationary, if for all $x, x' \in \mathbb{X}$, the covariance of the GP between those two points depends only on the difference $h = x - x' = (h_1, \dots, h_p)$. In that case, we will write $C_Z(x, x') = C_Z(h)$. For multidimensional problems, covariance functions are usually chosen as the product of one dimensional covariance functions:

$$C_Z(h) = s^2 \prod_{i=1}^p C_i(h_i; l_i) \quad (3.15)$$

These covariance functions introduce an additional parameter $l = (l_1, \dots, l_p)$ of dimension $p = \dim \mathbb{X}$, and a variance parameter s^2 . l is called the *length scale* and quantify the radius of influence of an evaluation along the different input variables. If the length scales are all equals, the covariance kernel is said *isotropic*. Otherwise, the kernel is *anisotropic*. A few common stationary 1D-covariance functions are introduced Table 3.1.

Name	$C(h; l)$	Regularity of sample paths
Gaussian	$\exp\left(-\frac{h^2}{2l^2}\right)$	C^∞
Exponential	$\exp\left(-\frac{ h }{l}\right)$	C^0
Matérn 3/2	$(1 + \sqrt{3}\frac{h}{l}) \exp\left(-\sqrt{3}\frac{h}{l}\right)$	C^1
Matérn 5/2	$\left(1 + \sqrt{5}\frac{h}{l} + \frac{5}{3}\frac{h^2}{l^2}\right) \exp\left(-\sqrt{5}\frac{h}{l}\right)$	C^2

Table 3.1 – Common stationary covariance functions

The choice of a particular covariance function is usually motivated by the wanted regularity of the sample paths. For example, if the unknown function f is assumed to be infinitely differentiable, a Gaussian kernel is suited for the modelling. One common choice is the Matérn kernel of order 5/2, so that the sample paths are twice-differentiable. Figure 3.2 shows the shape of the different covariance functions introduced Table 3.1, and also a few sample paths of an unconditioned Gaussian Process $Z \sim \text{GP}(0, C(\cdot, \cdot))$ for each of the covariance functions.

3.2.4 Initial design and validation

The $(p + 1)$ hyperparameters of the covariance, *i.e.* the length scales $(l_1, \dots, l_p)$ and the variance parameter s have to be estimated based on the training set $\mathcal{X}$. This is

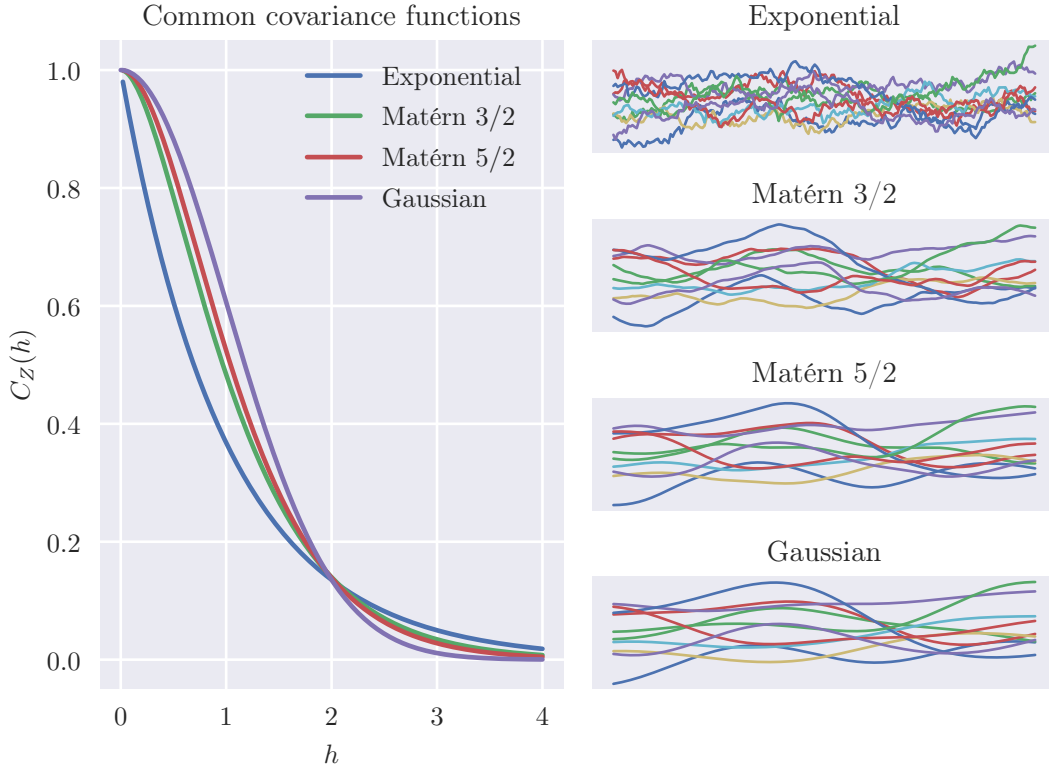


Figure 3.2 – Common covariance functions for GP regression. The right plots show (unconditioned) sample paths for those different covariance functions with same length scale.

usually done by MLE (see for instance [Ribaud et al. \(2019\)](#)), or by cross-validation (see [Ginsbourger et al. \(2009\)](#)).

In order to construct a Gaussian process, or a general metamodel for that matter, the model must first be evaluated on the initial design $\mathcal{X}$ which should present good space-filling properties. An usual choice is to use Latin Hypercube Sampling (LHS), and a widespread rule of thumb for the initial number of points of $\mathbb{X}$ to be evaluated is about $10p$. In order to validate the GP, *i.e.* to verify that the GP does not overfit the data, one can use cross-validation as described in [Dubrule \(1983\)](#).

3.3 General enrichment strategies for Gaussian Processes

We detailed so far the construction of a GP on a unknown function f , based on an initial design. However the surrogate is usually not an end in itself, as it may be used afterward to replace the original function for expensive procedures in a plug-in approach (see [Definition 3.1.2](#)). Obviously, if the surrogate is not accurate enough, the subsequent computations can become meaningless, that is why we are going to see how one can improve such surrogate.

The approximation given by the metamodel can be improved by adding points to the design, *i.e.* by evaluating some points in the input space. In order to incorporate as most information on f as possible, choosing those points can be done sequentially, by optimising a specified criterion, as described [Section 3.3.1](#), or we can look to add a batch of K points each iteration, as described in [Section 3.3.2](#). Those procedures can be performed until some specific stopping criterion is met. Alternatively, as the unknown function f is expensive to evaluate, we can define a maximal number of calls allowed to the function, and the procedures stop when the budget of evaluations is spent.

3.3.1 1-step lookahead strategies

For an unknown function f , a GP is initially constructed based on a design $\mathcal{X}_0 = \{(x_1, f(x_1)), \dots, (x_{n_0}, f(x_{n_0}))\}$, that consists of n_0 points of $\mathbb{X}$, and their corresponding evaluations. This GP is denoted $Z \mid \mathcal{X}_0$ and is defined as

$$Z \mid \mathcal{X}_0 \sim \mathcal{N}(m_{Z \mid \mathcal{X}_0}(x), \sigma_{Z \mid \mathcal{X}_0}^2(x)) \quad (3.16)$$

For notational convenience, the conditioning with respect to $\mathcal{X}_0$ will be omitted if the experimental design is clear from the context.

In order to enrich the design of experiments sequentially, we can adopt a *Stepwise Uncertainty Reduction* strategy ([Villemonteix et al., 2006](#)), which is based on the definition of a *criterion* (also called *learning function*), say κ , that measures the uncertainty upon a certain objective associated with the GP and with f . This criterion implicitly depends on the GP and thus the design used to construct it.

To select the next point to evaluate with f , this criterion is maximised on $\mathbb{X}$. This procedure is then repeated at each iteration: at the iteration n , with the set of evaluated points $\mathcal{X}_n$, the next point is then chosen as

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} \kappa(x) = \arg \max_{x \in \mathbb{X}} \kappa(x; Z \mid \mathcal{X}_n) \quad (3.17)$$

The GP is updated according to the new evaluation which is added to the design. The general principle of this strategy is summarized in [Algorithm 1](#)

Algorithm 1 SUR strategy: adaptive enrichment using a 1-step criterion

Require: Initial design $\mathcal{X}_0$, criterion function κ

Fit Z , a GP using the design $\mathcal{X}_0$

$n \leftarrow 0$

while (stopping criterion not met) **or** (evaluation budget not reached) **do**

$x_{n+1} \leftarrow \arg \max_{x \in \mathbb{X}} \kappa(x; Z \mid \mathcal{X}_n)$

 Evaluate $f(x_{n+1})$

$\mathcal{X}_{n+1} \leftarrow \mathcal{X}_n \cup \{(x_{n+1}, f(x_{n+1}))\}$

 Update the GP using $\mathcal{X}_{n+1}$

$n \leftarrow n + 1$

end while

3.3.2 Batch selection of points: sampling-based methods

Due to the structure of many computer codes, it may sometimes be beneficial to evaluate a batch of K points instead of selecting a unique point at each step of the procedure. SUR strategies as introduced before require the optimisation of a criterion function to select a single point, but in [Ginsbourger et al. \(2010\)](#), the authors introduce a few ways to transform the criterion in order to be able to select K points that will be evaluated by the true function.

Instead of adapting a 1-step criterion, we can rely on sampling in order to get K points of $\mathbb{X}$ to evaluate. This technique, named AK-MCS in [Echard et al. \(2011\)](#), which stands for *Adaptive-Kriging Monte-Carlo Sampling* is described in [Dubourg et al. \(2011\)](#), and has been applied and refined more recently in [Schöbi et al. \(2017\)](#); [Razaaly \(2019\)](#); [Razaaly et al. \(2020\)](#) for the estimation of extreme quantiles and small probabilities. Those are qualified as Monte-Carlo sampling methods, as they rely on obtaining samples drawn from a given distribution which represent points whose evaluation could be interesting with regards to the problem at stake.

Let κ be a 1-step criterion, and let $\bar{\kappa}(x) = \frac{\kappa(x)}{\int_{\mathbb{X}} \kappa(u) du}$ be the normalized criterion, so that $\int_{\mathbb{X}} \bar{\kappa} = 1$. In this case, $\bar{\kappa}$ can be seen as a density. Using an appropriate sampler, we can generate N i.i.d. samples $S = \{x_j^s\}_{1 \leq j \leq N}$ from this criterion.

However, as N should be large, there is no point in evaluating all the samples in S : using statistical reduction of the samples, we can find a reduced number of points, which are the most representative of the N samples. This is usually done by means of clustering algorithms, especially those who can take a fixed number of clusters such as the KMeans algorithm ([MacQueen, 1967](#)), in order to select precisely K points, number chosen in compliance with the wanted batch size. For each one of those centroids, we then look for the closest sampled point as the centroid can be located in a region of low or even zero $\bar{\kappa}$. Finally, the closest sampled points can then be considered to be evaluated. Additional filtering can then be performed in order to avoid numerical issues, if the chosen points are too close to each other or to points of the design. The principle of AK-MCS (using KMeans) is described [Algorithm 2](#).

3.4 Criteria of enrichment

We are now going to introduce a few criteria which are used for three different problems involving an unknown function f :

- The exploration of the input space will be addressed [Section 3.4.1](#), in order to target regions where the kriging prediction may differ significantly from the true function;
- the global optimisation of the function f in [Section 3.4.2](#), to focus the exploration of the input space $\mathbb{X}$ in the regions where f is minimal;
- the estimation of different level sets $f^{-1}(\{T\})$ and/or excursion sets in the form $f^{-1}(]-\infty, T])$ for $T \in \mathbb{R}$ in [Section 3.4.3](#).

Algorithm 2 AK-MCS: enrichment of the design using sampling**Require:** Initial design $\mathcal{X}_0$, Sampler according to $\bar{\kappa}$ **Require:** Batch size K , number of samples N Fit Z , a GP using the design $\mathcal{X}_0$ $n \leftarrow 0$ **while** (stopping criterion not met) **or** (evaluation budget not reached) **do**Sample N points $S = \{x_j^s\}_{1 \leq j \leq N}$ according to the distribution with pdf $\bar{\kappa}$ Get the K centroids $\{x_j^{\text{center}}\}_{1 \leq j \leq K}$ using KMeans on S **for** $j = 1$ to K **do** $x_{n+j} = \arg \min_{x \in S \setminus \{x_{n+1}, \dots, x_{n+j-1}\}} \|x - x_j^{\text{center}}\|$ **end for**Evaluate $f(x_{n+1}), \dots, f(x_{n+K})$ $\mathcal{X}_{n+K} \leftarrow \mathcal{X}_n \cup \{(x_{n+1}, f(x_{n+1})), \dots, (x_{n+K}, f(x_{n+K}))\}$ Condition the GP according to $\mathcal{X}_{n+K}$ $n \leftarrow n + K$ **end while****3.4.1 Criteria for exploration of the input space**

We are going to introduce first some common criteria of enrichment, that aim at exploring the input space $\mathbb{X}$. We suppose that the current GP has been conditioned with the design $\mathcal{X}_n$, composed of the points x_i and their evaluations $f(x_i)$ for $1 \leq i \leq n$.

Maximum of variance

A measure of uncertainty on the GP is $\max_{x \in \mathbb{X}} \sigma_{Z|\mathcal{X}_n}^2(x)$, the maximum value of the prediction variance on the space. A simple criterion is to select and evaluate the point corresponding to this maximum of variance:

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} \sigma_{Z|\mathcal{X}_n}^2(x) \quad (3.18)$$

This criterion by its simplicity is easy to implement, as the prediction variance is cheap to compute given a GP. It does not depend directly on the evaluations of the function at x_i for $1 \leq i \leq n$, uniquely on the distance between the candidate point x and the points of the design $\mathcal{X}_n$ and the covariance parameters.

Integrated Mean Square Error

The prediction variance is directly given by $\sigma_{Z|\mathcal{X}_n}^2$ and represents the uncertainty on the Gaussian regression. To summarize this uncertainty on the whole space $\mathbb{X}$, we define the Integrated Mean Square Error (IMSE) (Sacks et al., 1989) as

$$\text{IMSE}(Z | \mathcal{X}_n) = \int_{\mathbb{X}} \sigma_{Z|\mathcal{X}_n}^2(x) \, dx \quad (3.19)$$

For practical reasons, we can consider to integrate the MSE only on a subset $\mathfrak{X} \subset \mathcal{X}_n$ that yields

$$\text{IMSE}_{\mathfrak{X}}(Z | \mathcal{X}_n) = \int_{\mathbb{X}} \sigma_{Z|\mathcal{X}_n}^2(x) \mathbb{1}_{\mathfrak{X}}(x) \, dx = \int_{\mathfrak{X}} \sigma_{Z|\mathcal{X}_n}^2(x) \, dx \quad (3.20)$$

For notational convenience, when the GP considered is clear from the context, only the design used to fit it will be an argument: $\text{IMSE}(Z \mid \mathcal{X}_n) = \text{IMSE}(\mathcal{X}_n)$. Unfortunately, exact evaluation of this integral is impossible, so it needs to be approximated using numerical integration, such as Monte-carlo or quadrature rules such as Gaussian, or Hermite rule to cite a few.

We can then study how adding a certain point would affect the IMSE. For a given $x \in \mathbb{X}$ and an outcome $z = f(x) \in \mathbb{R}$, the *augmented design* is defined as $\mathcal{X}_n \cup \{(x, z)\}$, and the IMSE of the augmented design is $\text{IMSE}(\mathcal{X}_n \cup \{(x, z)\})$. Before the actual experiment though, z is unknown, but we can model it by its distribution given by the GP (per Eq. (3.16)). So for a given candidate x , the IMSE when evaluating the point x will be on average

$$\mathbb{E}_{Z(x)} \left[\text{IMSE}(\mathcal{X}_n \cup \{(x, Z(x))\}) \right] \quad (3.21)$$

where the expectation is to be taken with respect to the random variable $Z(x)$. The expression of Eq. (3.21) requires then the evaluation of an integral in dimension $1 + \dim \mathbb{X}$. We then use the term *augmented* to signal that a quantity has been evaluated with an augmented design, and the output have been averaged in this fashion.

As each outcome of $Z(x)$ requires to fit a GP and to compute the IMSE, a precise evaluation is quite expensive. A strategy found for instance in Villemonteix et al. (2006) is to take M possible outcomes for $Z(x)$, corresponding to evenly spaced quantiles of its distribution, or using Gauss-Hermite quadratures (Bernard, 2019) in order to take advantage of the Gaussian nature of $Z(x)$. The hyperparameters of the GP should not be reevaluated when augmenting the design, in order to get comparable values for the IMSE.

Finally, we can maximise a criterion which is the opposite of the expression in Eq. (3.21) to enrich the design:

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} -\mathbb{E}_{Z(x)} \left[\text{IMSE}(\mathcal{X}_n \cup \{(x, Z(x))\}) \right] \quad (3.22)$$

3.4.2 Criteria for optimisation of the objective function

The criteria detailed above aim at reducing the epistemic uncertainty modelled through the Gaussian Process. In other words, we try to improve our knowledge on the unknown function globally. We are now going to evoke a few criteria which are driven by the global optimisation of the function.

Those methods usually aim at striking a balance between the *exploration* of the whole space $\mathbb{X}$ and *intensification*, *i.e.* evaluating the function near its optimum. Let f be the unknown function, and Z be a GP constructed based on an initial design $\mathcal{X}_0 = \{(x_i, f(x_i))\}$.

Probability of improvement

We are first going to introduce the probability of improvement PI, which is the probability that the GP is smaller than a threshold $f_{\min}$. Due to the Gaussian nature

of $Z(x)$, this probability can be written in closed form using $\Phi = F_{\mathcal{N}(0,1)}$ the cdf of a centered standard Gaussian r.v.

$$\text{PI}(x) = \mathbb{P}[Z(x) < f_{\min}] \quad (3.23)$$

$$= \Phi\left(\frac{m_Z(x) - f_{\min}}{\sigma_Z(x)}\right) \quad (3.24)$$

This threshold can have different forms

- $f_{\min} = \min_i f(x_i)$: the GP is compared with the current minimal value reached by the function
- $f_{\min} = \min_i f(x_i) + \epsilon$: by introducing a small tolerance ϵ , we encourage exploration instead of intensification.

Using the probability of improvement tends to select points quite close to the point evaluated so far, thus does favor intensification at the expense of exploration.

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} \text{PI}(x) \quad (3.25)$$

Expected improvement and EGO

One of the most common criteria for global optimisation is the *Expected Improvement* (EI) (Moćkus, 1974), and the SUR strategy using it as an enrichment criterion is called *Efficient Global Optimisation* (EGO) (Jones et al., 1998). Analogously to the probability of improvement, we define the improvement $I(x)$ as the random variable defined as

$$I(x) = [f_{\min} - Z(x)]_+ \quad (3.26)$$

where $[y]_+ = \max(y, 0)$. The Expected Improvement EI is

$$\text{EI}(x) = \mathbb{E}_{Z(x)}[I(x)] = \mathbb{E}_{Z(x)}[[f_{\min} - Z(x)]_+] \quad (3.27)$$

Again, a closed form is available to compute the expected improvement, that does not require the direct evaluation of the expectation Eq. (3.27).

$$\text{EI}(x) = (f_{\min} - m_Z(x)) \Phi\left(\frac{f_{\min} - m_Z(x)}{\sigma_Z(x)}\right) + \sigma_Z(x) \phi\left(\frac{f_{\min} - m_Z(x)}{\sigma_Z(x)}\right) \quad (3.28)$$

The EI is then quite easy to evaluate and furthermore it is possible to compute the gradient of the EI and use it for the optimisation as done in Marmin et al. (2015).

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} \text{EI}(x) \quad (3.29)$$

IAGO

Another criterion worth mentioning is a criterion based on the distribution of the minimisers (Villemonteix et al., 2006; Hennig and Schuler, 2011), called IAGO (*Informational Approach to Global Optimisation*). Let z_i be a sample path of Z , and let x_i^* the global minimiser of z_i . We denote then X^* the random variable corresponding to the global

minimiser of Z . We consider the differential entropy of X^* introduced [Definition 1.2.14](#), given the augmented design $\mathcal{X}_n \cup \{(x, Z(x))\}$: $H[X^* | \mathcal{X}_n \cup \{(x, Z(x))\}]$. So at each step, we choose the point that gives the smallest expected uncertainty on the location of the global minimisers of the sample paths. The criterion can then be written as

$$x_{n+1} = \arg \max_{x \in \mathbb{X}} -\mathbb{E}_{Z(x)} \left[H[X^* | \mathcal{X}_n \cup \{(x, Z(x))\}] \right] \quad (3.30)$$

[Figure 3.3](#) shows different criteria introduced before for an unknown function.

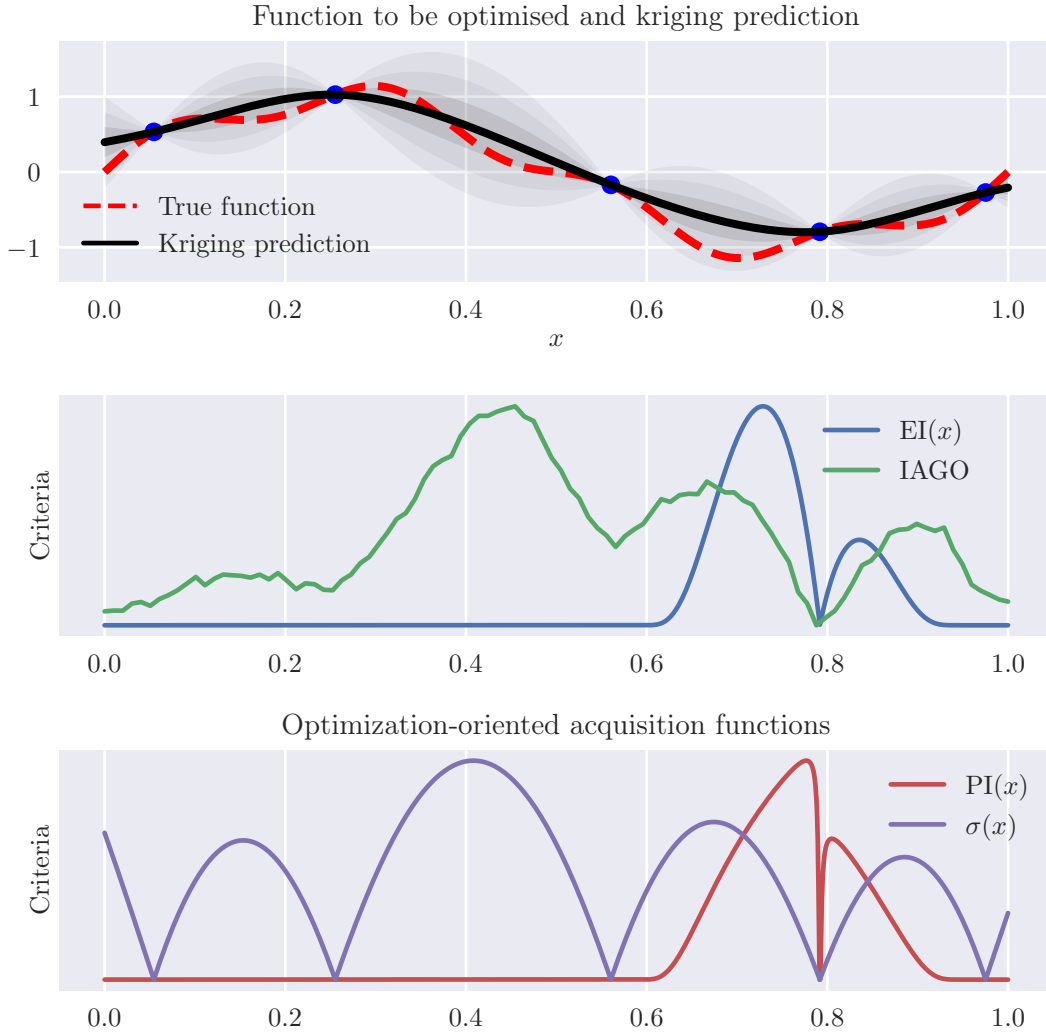


Figure 3.3 – Example of optimisation criteria. At this iteration, EI and PI aim toward intensification, while IAGO and the maximum of variance favorize exploration

3.4.3 Contour and volume estimation

In this section, we are interested in the estimation of the inverse image of sets: $\{x \in \mathbb{X} \mid f(x) = T\} = f^{-1}(\{T\})$ for level sets, or $\{x \in \mathbb{X} \mid f(x) \leq T\} = f^{-1}(]-\infty, T])$ called excursion sets. An example of this estimation is illustrated Fig. 3.4, where the focus is on the estimation of the set $f^{-1}(]-\infty, T])$. For that, we are going to introduce enrichment strategies which rely on the evaluation of the function around the level set $\{f = T\}$.

3.4.3.a Reduction of the augmented IVPC

For a function f , let us assume that we are interested in the estimation of the set $f^{-1}(B) = \{x \in \mathbb{X} \mid f(x) \leq T\}$, with $B =]-\infty, T]$. Given $\mathbb{P}_X$ a probability measure on $\mathbb{X}$, and $B \subset \mathbb{R}$, we can compute $V = \mathbb{P}_X[f^{-1}(B)]$. For $B =]-\infty, T]$, it is equivalent to compute the volume of the *excursion set* of f below T : $V = \mathbb{P}_X[f(X) \leq T]$. Such an example is illustrated Fig. 3.4.

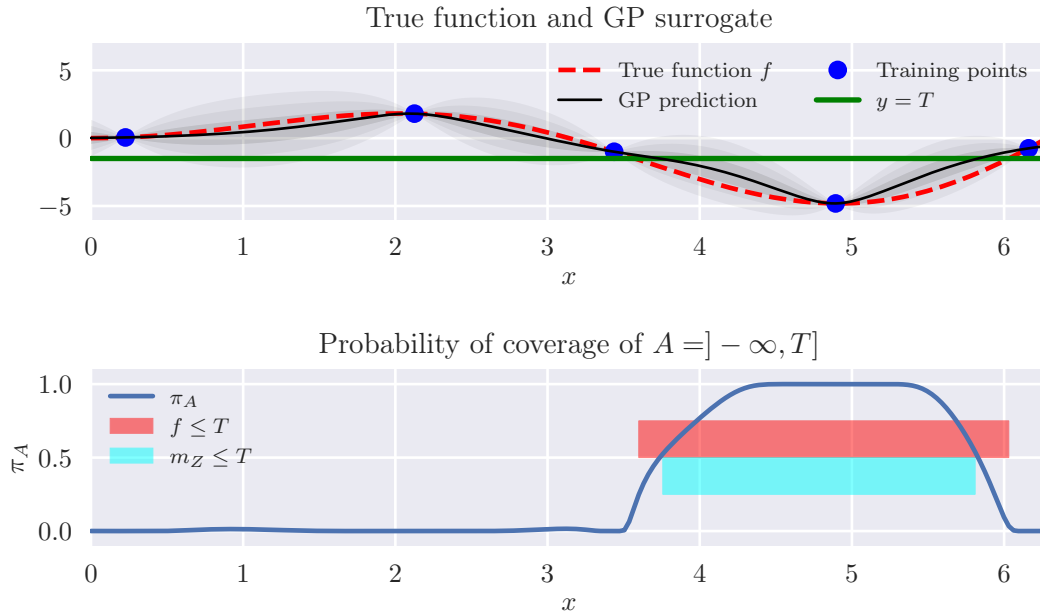


Figure 3.4 – Estimation of $f^{-1}(]-\infty, T])$ using GP and probability of coverage associated

Let Z be a GP, indexed by $\mathbb{X}$ with continuous sample paths. For a measurable set $B \subset \mathbb{R}$, $A = Z^{-1}(B)$ is a random closed set. We define π_A its probability of coverage:

$$\pi_A(x) = \mathbb{P}_{Z(x)}[x \in A] = \mathbb{P}_{Z(x)}[Z(x) \in B] \quad (3.31)$$

Figure 3.4 shows the probability of coverage of the set A , the true set $f^{-1}(B)$, and the *plug-in* estimation of this set: $m_Z^{-1}(B) = \{x \in \mathbb{X} \mid m_Z(x) \leq T\}$.

For a given $x \in \mathbb{X}$, the event “ x belongs to A ” is Bernoulli random variable, with probability $\pi_A(x)$, thus has variance

$$\mathcal{V}_A(x) = \pi_A(x)(1 - \pi_A(x)) \quad (3.32)$$

We can also define the probability of missclassification. If $\pi_A(x) > 0.5$, x is classified in A so the probability of misclassification is $1 - \pi_A(x)$, and on the other hand, if $\pi_A(x) < 0.5$, x is classified in A^C , and the probability of missclassification is $\pi_A(x)$. This can be condensed as

$$\mathcal{P}_{\text{mis}}(x) = \min(\pi_A(x), 1 - \pi_A(x)) \quad (3.33)$$

For both $\mathcal{P}_{\text{mis}}$ and $\mathcal{V}_A$, the maximum is reached for $\pi_A(x) = 1/2$.

As we want to classify each point either *in* A or *out of* A , we can look for the different level-sets of the probability of coverage: for $\eta \in [0, 1]$, we define the η -level set of π_A , also called *Vorob'ev quantiles* (see [Vorobyev and Vorobyev \(2003\)](#))

$$Q_\eta = \{x \in \mathbb{X} \mid \pi_A(x) \geq \eta\} \quad (3.34)$$

Those sets are decreasing (with respect to the inclusion) when η increases:

$$0 \leq \eta \leq \xi \leq 1 \implies Q_\xi \subseteq Q_\eta \quad (3.35)$$

Using the properties of the GP stating that $x \in \mathbb{X}$, $Z(x) \sim \mathcal{N}(m_Z(x), \sigma_Z^2(x))$, we can express the probability of coverage of the random set A using $F_{Z(x)}$, the cdf of the r.v. $Z(x)$, or Φ , the cdf of the centered standard Gaussian distribution:

$$\pi_A(x) = \mathbb{P}_{Z(x)}[Z(x) \leq T] = F_{Z(x)}(T) = \Phi\left(-\frac{m_Z(x) - T}{\sigma_Z(x)}\right) \quad (3.36)$$

$$\mathcal{V}_A(x) = \pi_A(x)(1 - \pi_A(x)) \quad (3.37)$$

$$\mathcal{P}_{\text{mis}}(x) = \min(F_{Z(x)}(T), 1 - F_{Z(x)}(T)) = \min(F_{Z(x)}(T), F_{-Z(x)}(-T)) \quad (3.38)$$

$$= \Phi\left(-\frac{|m_Z(x) - T|}{\sigma_Z(x)}\right) \quad (3.39)$$

[Echard et al. \(2011\)](#) defines the argument of the cdf [Eq. \(3.39\)](#) as the *reliability index* ρ (denoted U in its original definition [Echard et al. \(2011\)](#)):

$$\rho(x) = \frac{|m_Z(x) - T|}{\sigma_Z(x)} \quad (3.40)$$

Based on those quantities, we will introduce two criteria which aim at improving the classification problem of $x \in A = f^{-1}(B)$.

First, we can choose to evaluate the point which has the maximal probability of misclassification, or equivalently since Φ is monotonously increasing, maximise the criterion given by

$$x_{n+1} = \arg \max_{\theta \in \Theta} -\frac{|m_Z(x) - T|}{\sigma_Z(x)} \quad (3.41)$$

This criteria will be maximal when either the kriging prediction is close to the level set T (intensification aspect), or if there is a large prediction variance (exploration aspect).

Alternatively, we can also use the variance of the probability of coverage $\mathcal{V}_A$ as an indication of the uncertainty of the classification problem. As the maximum of variance will be obtained for points having a probability of coverage $\pi_A(x) = 0.5$, looking and evaluating the maximiser of this criterion is not really relevant: any point verifying $m_Z(x) = T$ would be suited, without taking into account its prediction error $\sigma_Z^2(x)$, and a point already well predicted (*i.e.* close to other points already in the design) could be evaluated, incorporating little new knowledge.

Instead, by integrating $\mathcal{V}_A$ over the whole space $\mathbb{X}$, we get a measure of the uncertainty in the classification problem associated with the design $\mathcal{X}_n$, which we will denote $\text{IVPC}(\mathcal{X}_n)$ (Integrated Variance of the Probability of Coverage).

$$\text{IVPC}(\mathcal{X}_n) = \int_{\mathbb{X}} \mathcal{V}_A(x) dx \quad (3.42)$$

A similar function has been introduced in [Beet et al. \(2012\)](#). We can finally introduce a criterion based on the augmented design, as done for the IMSE.

$$x_{n+1} = \arg \min_{x \in \mathbb{X}} \mathbb{E}_{Z(x)} [\text{IVPC}(\mathcal{X}_n \cup \{(x, Z(x))\})] \quad (3.43)$$

Other criteria have also been developed, such as the use of the theory of random sets ([El Amri et al., 2019](#)), in order to define a measure of uncertainty based on Vorob'ev mean and deviation ([Vorobyev and Vorobyev, 2003](#)).

3.4.3.b Sampling in the Margin of uncertainty

Using the level sets, we can construct the η -margin of uncertainty, as introduced in [Dubourg et al. \(2011\)](#), that is the set of points $x \in \mathbb{X}$ that we cannot classify in or out of A with high enough probability. Setting the classical level $\eta = 0.05$ for instance, $Q_{1-\frac{\eta}{2}} = Q_{0.975}$ is the set of points whose probability of coverage is higher than 0.975, while $Q_{\frac{\eta}{2}} = Q_{0.025}$ is the set of points whose probability of coverage is higher than 0.025, thus its complement in $\mathbb{X}$, denoted by $Q_{\frac{\eta}{2}}^C$ is the set of points whose probability of coverage is lower than 0.025. Obviously, $Q_{1-\frac{\eta}{2}} \subset Q_{\frac{\eta}{2}}$.

The η -margin of uncertainty $\mathbb{M}_\eta$ for $\eta = 0.05$ is then defined as the sets of points whose coverage probability is between 0.025 and 0.975:

$$\mathbb{M}_\eta = \left(Q_{1-\frac{\eta}{2}} \cup Q_{\frac{\eta}{2}}^C \right)^C = Q_{1-\frac{\eta}{2}}^C \cap Q_{\frac{\eta}{2}} = Q_{\frac{\eta}{2}} \setminus Q_{1-\frac{\eta}{2}} \quad (3.44)$$

$$= \left\{ x \in \mathbb{X} \mid \frac{\eta}{2} \leq \pi_A(x) \leq 1 - \frac{\eta}{2} \right\} \quad (3.45)$$

Based on this set, we can construct easily a sampling criterion by using the indicator function of the margin of uncertainty: $\mathbb{1}_{\mathbb{M}_\eta}(x) = \kappa_{\mathbb{M}}(x) = \bar{\kappa}_{\mathbb{M}}(x) \cdot \text{Vol}(\mathbb{M}_\eta)$. Evaluating points in this region allows to reduce the uncertainty at x (completely if the GP is

interpolant, as $\sigma_Z(x)$ becomes 0 after evaluating $f(x)$ on the classification of the point x .

Finally, using this indicator function, [Algorithm 2](#) can be applied to enrich the design, by adding a batch of K points at each iterations. Using an acceptance-rejection method (since the pdf of the sampling distribution is an indicator function), we can easily obtain samples within the margin of uncertainty: they are shown on the right plot of [Fig. 3.5](#), with the frontiers of the margin. Using the KMeans algorithm, we can compute $K = 10$ centroids, which represent statistically the samples (red stars on the figure). We can then see that they are located close to the true level-set of the function in this case.

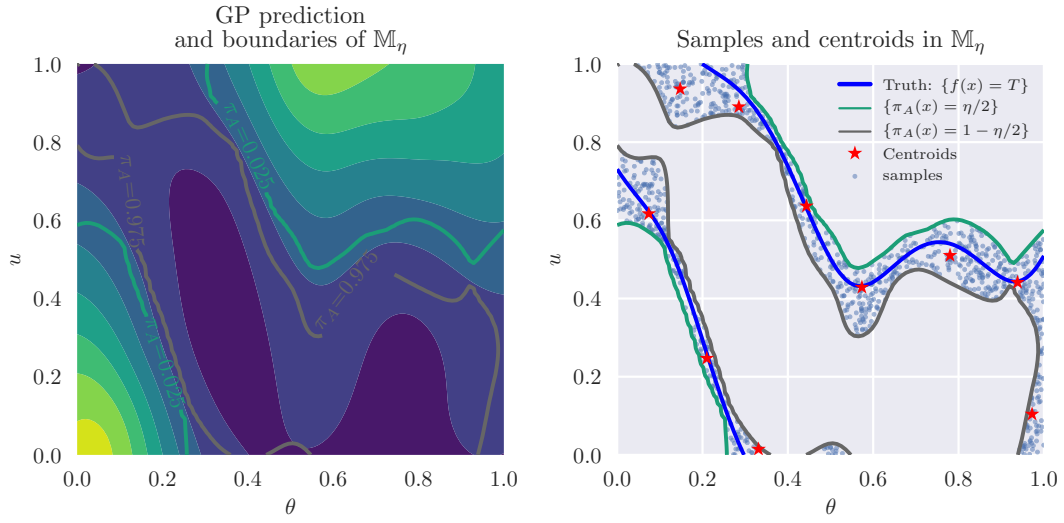


Figure 3.5 – Margin of uncertainty, and construction of the K points of the batch to be evaluated, using samples and centroids

3.5 Adaptive strategies for robust optimisation using GP

In the previous section, we introduced enrichment strategies to solve different problems for a generic function $f : \mathbb{X} \rightarrow \mathbb{R}$. Before introducing specific strategies for robust calibration (as defined in [Chapter 2](#)), we are first going to recall briefly the setting and corresponding problems of robust calibration. Let J be an objective function taking two inputs:

$$\begin{aligned} J : \Theta \times \mathbb{U} &\longrightarrow \mathbb{R} \\ (\theta, u) &\longmapsto J(\theta, u) \end{aligned} \quad (3.46)$$

In a context of calibration, this function measures the misfit between the output of the numerical model and the observations.

The first input of the function, $\theta \in \Theta$, is the calibration parameter which we wish to control and $u \in \mathbb{U}$ is the environmental parameter. We model this intrinsic variability of the environmental parameter with a random variable U . In this context, different objectives can be derived in order to get a satisfying robust estimate $\hat{\theta}$. Robustness here

is to be understood as the ability to get satisfying performances of the objective function $J(\hat{\theta}, u)$ for any $u \in \mathbb{U}$, realisation of U .

Based on an initial design $\mathcal{X}_0 = \{((\theta_i, u_i), J(\theta_i, u_i))\}_{1 \leq i \leq n_0}$ of n_0 points in $\Theta \times \mathbb{U}$, we assume that we constructed a GP, say Z , on $\Theta \times \mathbb{U}$. In the following, if the design used to construct the GP is clear from the context, it will be omitted in the notation.

As a GP, Z is described by its mean function $m_Z : \Theta \times \mathbb{U} \rightarrow \mathbb{R}$ and a covariance function $C : (\Theta \times \mathbb{U})^2 \rightarrow \mathbb{R}$, and the prediction variance is $\sigma_Z^2(\theta, u) = C((\theta, u), (\theta, u))$.

As a GP, for any $(\theta, u) \in \Theta \times \mathbb{U}$, $Z(\theta, u)$ is normally distributed:

$$Z(\theta, u) \sim \mathcal{N}(m_Z(\theta, u), \sigma_Z^2(\theta, u)) \quad (3.47)$$

A surrogate of J , constructed using the GP Z is then m_Z , which will be used to compute the different estimates. However, the initial design $\mathcal{X}_0$ is probably too sparse to get m_Z accurately enough to compute the wanted estimates. For that purpose, we propose enrichment strategies in order to add *relevant* points to the design $\mathcal{X}_0$.

One main challenge for this problem lies in the decomposition of the input space, ($\mathbb{X}$ previously) into $\Theta \times \mathbb{U}$. The objective on each of these spaces is different: one needs to explore sufficiently the space $\mathbb{U}$ according to the distribution of U , in order to be able to compute $\mathbb{P}_U$ and $\mathbb{E}_U$, while Θ does not need to be explored as much because the objective function is supposed to be *optimised* with respect to θ .

In what follows, we are going to focus on the estimation of the conditional minimisers $\theta^* : u \mapsto \theta^*(u)$, which can be used to estimate the mode of the distribution of the conditional minimisers θ_{MPE} , and on the estimators which satisfy properties based on regret. For that purpose, the methods introduced here will usually involve the computation and the optimisation of a criterion on the joint space directly, following the method of [Algorithm 1](#) on [Page 69](#), and sampling in regions of large uncertainties, as detailed in [Algorithm 2](#) [Page 71](#).

For other robust estimators, such as the minimum of the expectation θ_{mean} , other strategies have been derived, which consist first in removing the dependency on U . For instance in [Janusevskis and Le Riche \(2010\)](#), the authors integrate directly the GP with respect to U in order to get another GP, which models the function $\theta \mapsto \mathbb{E}_U[J(\theta, U)]$. This integrated GP will then be used to select a point $\hat{\theta}$, and then a couple of points (θ_{n+1}, u_{n+1}) is chosen in order to reduce a measure of uncertainty on the space $\{\hat{\theta}\} \times \mathbb{U}$. Alternatively, when introducing the variance in a multiobjective optimisation problem, [Miranda et al. \(2016\)](#) proposes a method based on polynomial chaos expansion in order to get both the value and the gradient of the expected value and the variance.

3.5.1 PEI for the conditional minimisers

In the previous chapter, we introduced the conditional minimums $u \mapsto J^*(u)$ and the conditional minimisers $u \mapsto \theta^*(u)$. Each evaluation of these functions require an optimisation procedure, thus can be expensive from a computational point of view. [Ginsbourger et al. \(2014\)](#) proposes a criterion, adapted from the EI, which aims at solving such a

problem: the *Profile Expected Improvement* PEI is defined as

$$\text{PEI}(\theta, u) = \mathbb{E} [[f_{\min}(u) - Z(\theta, u)]_+] \quad (3.48)$$

$$\text{with } f_{\min}(u) = \max_{x \in \mathcal{X}} (\min_{\theta \in \Theta} f(x), \min_{\theta \in \Theta} m_Z(\theta, u)) \quad (3.49)$$

for the design $\mathcal{X} = \{(x_i, f(x_i))\}_{1 \leq i \leq n}$ with the slight abuse of notation $J(\theta, u) = f(x)$ for $x = (\theta, u)$.

This allows us to see the similarity with the EI criterion: instead of having a fixed threshold, the PEI introduces a criterion that depends on u . Figure 3.6 shows an example of GP enriched using the PEI criterion. The criterion can then be used in a classical SUR strategy introduced in Algorithm 1, where at each iteration, the PEI is maximised over $\Theta \times \mathbb{U}$:

$$x_{n+1} = (\theta_{n+1}, u_{n+1}) = \arg \max_{(\theta, u) \in \Theta \times \mathbb{U}} \text{PEI}(\theta, u) \quad (3.50)$$

and the maximiser is evaluated by f and added to the design. Once a stopping criterion is met, such as the maximal variance attained at the conditional minimum is below a specified threshold ϵ_{stop} ,

$$\max_{u \in \mathbb{U}} \sigma_Z^2(m_{Z^*}(u), u) \leq \epsilon_{\text{stop}} \quad (3.51)$$

we can construct accurate surrogates of $J^*(u)$ and $\theta^*(u)$, using the GP in a plug-in approach:

$$m_{Z^*}(u) = \min_{\theta \in \Theta} m_Z(\theta, u) \quad \text{and} \quad \theta_Z^*(u) = \arg \min_{\theta \in \Theta} m_Z(\theta, u) \quad (3.52)$$

and with a set of i.i.d. samples $\{u_i\}_i$ of U , use the surrogates to find the MPE, based on the distribution of $\{\theta_Z^*(u_i)\}$.

3.5.2 Gaussian formulations for the relative and additive regret families of estimators

We are now going to detail how Gaussian processes can be used to estimate regret-based estimators. Recalling the previous chapter, the family of regret-based estimators are defined as

$$\{\theta_{\text{AR}, \beta} = \arg \max_{\theta \in \Theta} \Gamma_{\beta}(\theta) = \arg \max_{\theta \in \Theta} \mathbb{P}_U [J(\theta, U) \leq J^*(U) + \beta] \mid \beta \geq 0\} \quad (3.53)$$

in the case of additive-regret, and

$$\{\theta_{\text{RR}, \alpha} = \arg \max_{\theta \in \Theta} \Gamma_{\alpha}(\theta) = \arg \max_{\theta \in \Theta} \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)] \mid \alpha \geq 1\} \quad (3.54)$$

for the relative-regret.

As explained Section 2.4, in order to get an estimator, two starting points can be considered: either fixing the threshold α (or β), and computing the associated estimator $\theta_{\text{RR}, \alpha}$ (or $\theta_{\text{AR}, \beta}$), or instead, looking to get a probability of acceptability of level p , and then compute the associated estimator and threshold. For the first alternative, we are

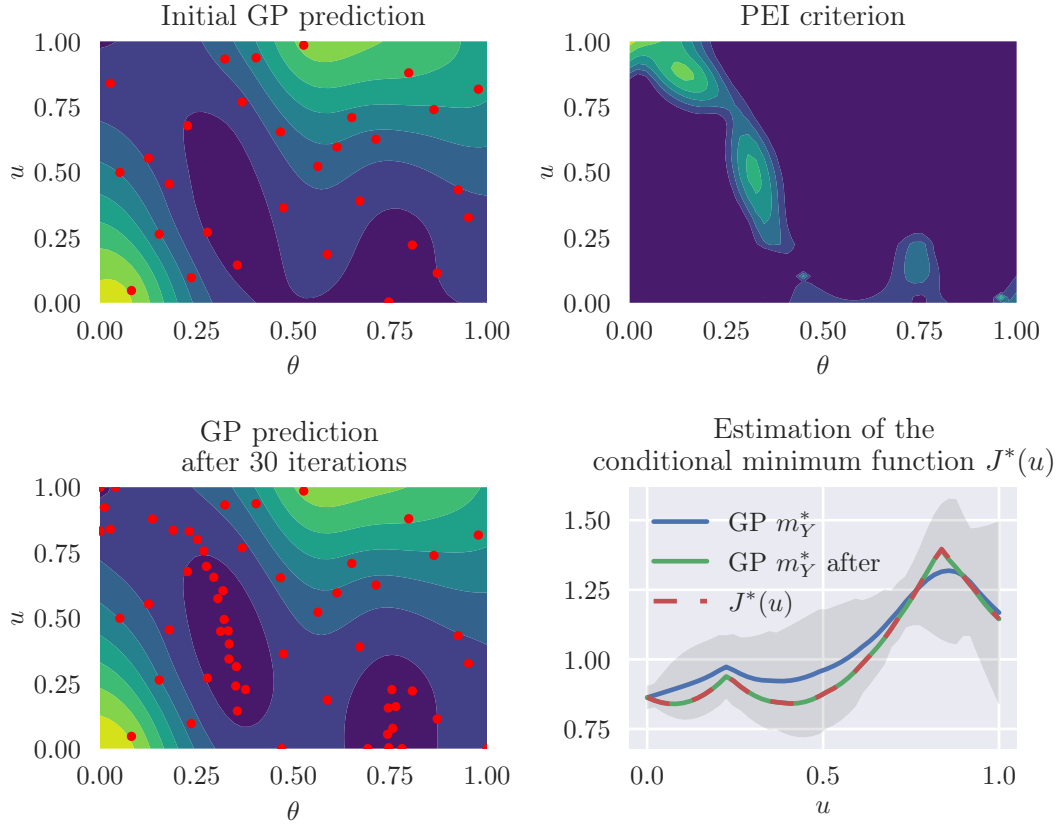


Figure 3.6 – GP after 30 additional iterations chosen using PEI. The GP prediction using the initial design is represented on the top left plot. The PEI criterion for this GP is evaluated over the whole space and plotted top right. Bottom left plot is the GP prediction after 30 iterations. Bottom right is the surrogate estimation of the conditional minimum. The shaded region corresponds to the initial surrogate conditional minimum, $\pm$ one standard deviation. After the iterations, the surrogate estimate and the true value coincide.

going to focus on improving the estimation of Γ using GP in [Section 3.5.3](#), while for the second, we will try to improve the estimation of the quantile of the ratio in [Section 3.5.4](#).

We first need to detail some notations with respect to the conditional minimum and conditional minimisers, and their equivalent for the GP.

We define Z^* as

$$Z^*(u) \sim \mathcal{N}(m_{Z^*}(u), \sigma_{Z^*}^2(u)) \quad (3.55)$$

where

$$m_{Z^*}(u) = \min_{\theta \in \Theta} m_Z(\theta, u) \quad (3.56)$$

$$\theta_Z^*(u) = \arg \min_{\theta \in \Theta} m_Z(\theta, u) \quad (3.57)$$

$$\sigma_{Z^*}^2(u) = \sigma_Z^2(\theta_Z^*(u), u) \quad (3.58)$$

We use then the minimiser and the minimum of the GP prediction as estimates of the conditional minimum and minimiser, similarly as in [Ginsbourger et al. \(2014\)](#). To generalize notions of the additive- and relative-regret, we can define $\Delta_{\alpha,\beta} = Z - \alpha Z^* - \beta$. This difference $\Delta_{\alpha,\beta}$ is a linear combination of correlated Gaussian processes, thus is a GP as well and its distribution can be derived by first considering the joint distribution of $Z(\theta, u)$ and $Z^*(u) = Z(\theta_Z^*(u), u)$:

$$\begin{bmatrix} Z(\theta, u) \\ Z^*(u) \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} m_Z(\theta, u) \\ m_{Z^*}(u) \end{bmatrix}; \begin{bmatrix} C((\theta, u), (\theta, u)) & C((\theta, u), (\theta_Z^*(u), u)) \\ C((\theta, u), (\theta_Z^*(u), u)) & C((\theta_Z^*(u), u), (\theta_Z^*(u), u)) \end{bmatrix} \right) \quad (3.59)$$

Multiplying by the 1×2 matrix $[1 \quad -\alpha]$ and translating by $-\beta$ yields

$$\Delta_{\alpha,\beta}(\theta, u) \sim \mathcal{N}(m_\Delta(\theta, u); \sigma_\Delta^2(\theta, u)) \quad (3.60)$$

$$m_\Delta(\theta, u) = m_Z(\theta, u) - \alpha m_{Z^*}(u) - \beta \quad (3.61)$$

$$\sigma_\Delta^2(\theta, u) = \sigma_Z^2(\theta, u) + \alpha^2 \sigma_{Z^*}^2(u) - 2\alpha C((\theta, u), (\theta_Z^*(u), u)) \quad (3.62)$$

We are then interested in the random set $\{(\theta, u) \mid \Delta_{\alpha,\beta}(\theta, u) \leq 0\}$. The different combinations of α and β dictate if we are either interested in the additive or the relative regret:

- $(\alpha, \beta) = (1, \beta)$ corresponds to the additive regret
- $(\alpha, \beta) = (\alpha, 0)$ corresponds to the relative regret

Decomposing the variance σ_Δ^2 in [Eq. \(3.62\)](#), 3 sources of uncertainty appear:

- σ_Z^2 is the prediction variance of the GP on J , that is directly reduced when additional points are evaluated
- $\sigma_{Z^*}^2$ is the variance of the predicted value of the minimisers.
- Assuming a stationary form of the covariance, the third term is dependent only on the distance between θ and $\theta_{Z^*}(u)$, (the estimated conditional minimiser). By separating the θ and u components, the covariance term can be written $C((\theta, u), (\theta', u')) = s\mathcal{K}_\theta(\|\theta - \theta'\|)\mathcal{K}_u(\|u - u'\|)$, and substituting $\theta_Z^*(u)$ for θ' gives

$$C((\theta, u), (\theta_Z^*(u), u)) = s\mathcal{K}_\theta(\|\theta - \theta_Z^*(u)\|)\mathcal{K}_u(0) \quad (3.63)$$

$$= s\mathcal{K}_\theta(\|\theta - \theta_Z^*(u)\|) \quad (3.64)$$

This decomposition highlights the fact that the prediction error σ_Δ^2 of the difference $\Delta_{\alpha,\beta}$ measured at a point (θ, u) will not be reduced completely by evaluating the function J

at this point, as only the prediction variance σ_Z^2 will be significantly affected in general. In other words, reducing the uncertainty at a point (θ, u) may require the evaluation of another point $(\tilde{\theta}, \tilde{u}) \neq (\theta, u)$.

In the following, unless explicitly stated, we will focus on the relative-regret family of estimators, so $\beta = 0$ and $\alpha \geq 1$, thus only keeping α in the subscript. The strategies introduced hereafter can be easily adapted for the additive-regret, by setting $\alpha = 1$, and translating the mean functions accordingly, thus defining $\{\Delta_{\alpha=1} \leq \beta\}$.

3.5.3 GP-based methods for the estimation of Γ_α

Let consider $\alpha \geq 1$ fixed. In order to compute $\theta_{RR,\alpha}$, we need to estimate and optimise the function Γ_α . For that purpose, we are going to explore the space to improve the estimation of Γ_α , and once sufficient knowledge is acquired, use the plug-in estimate $\hat{\Gamma}_\alpha$ (*i.e.* the one obtained by replacing J by m_Z) for the optimisation, to get $\hat{\theta}_{RR,\alpha}$.

3.5.3.a Estimation of Γ_α

For a given $\theta \in \Theta$, the coverage probability of the α -acceptable region, *i.e.* the probability for θ to be α -acceptable is

$$\Gamma_\alpha(\theta) = \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)] \quad (3.65)$$

$$= \mathbb{E}_U [\mathbb{1}_{J(\theta, U) \leq \alpha J^*(U)}] \quad (3.66)$$

As J is not known perfectly, it can be seen as a classification problem. This classification problem can be approached with a plug-in approach, that is to consider m_Z instead of J :

$$\mathbb{1}_{J(\theta, u) \leq \alpha J^*(u)} \approx \mathbb{1}_{m_Z(\theta, u) \leq \alpha m_{Z^*}(u)} \quad (3.67)$$

Alternatively we can also approximate this indicator function by consider the probability of the random variable Δ to be less than 0:

$$\mathbb{1}_{J(\theta, u) \leq \alpha J^*(u)} \approx \mathbb{P}_Z [\Delta_\alpha(\theta, u) \leq 0] = \pi_\alpha(\theta, u) \quad (3.68)$$

Based on those two approximation, we can define two different estimations of Γ_α , namely $\hat{\Gamma}_\alpha^{\text{PI}}$ with the plug-in approach of Eq. (3.67), and $\hat{\Gamma}_\alpha^\pi$ for the probabilistic one in Eq. (3.68), that we will detail shortly after. Based on those estimations, two strategies emerge:

- For $\hat{\Gamma}_\alpha^{\text{PI}}$, we are going to reduce the expected augmented IMSE of the GP $Z - \alpha Z^*$.
- For $\hat{\Gamma}_\alpha^\pi$, we are going to reduce the expected augmented integrated variance of probability of coverage IVPC.

3.5.3.b Reduction of the augmented IMSE

For the plug-in approach, the chosen estimator is defined by:

$$\Gamma_\alpha^{\text{PI}}(\theta) = \mathbb{P}_U [m_Z(\theta, u) \leq \alpha m_{Z^*}(u)] \quad (3.69)$$

The outer expectation operator is to be computed numerically, using quadrature rule, or Monte-Carlo methods. In general, from a set of i.i.d. samples $\{u_i\}_{1 \leq i \leq n_u}$ sampled from U , we define

$$\hat{\Gamma}_\alpha(\theta) = \frac{1}{n_u} \sum_{i=1}^{n_u} \mathbb{1}_{m_Z(\theta, u_i) - \alpha m_{Z^*}(u_i) \leq 0} \quad (3.70)$$

Due to the fact that the GP surrogate is cheap to evaluate, the computation of the outer expectation with respect to U is assumed to be performed without too much problems.

In order to improve the accuracy of this estimator, one need to improve the GP prediction m_Z of the objective function J . In this case, we propose to use the augmented IMSE of the GP $Z - \alpha Z^* = \Delta_\alpha$ as a learning function. The choice of the IMSE (instead of choosing the point of maximal variance for instance) comes from the decomposition of the variance [Eq. \(3.62\)](#).

Let $\hat{\Gamma}_{\alpha,n}^{\text{PI}}$ be the plug-in approximation of Γ_α , constructed using the Gaussian Process surrogate with n points added according to the augmented IMSE criterion. [Figure 3.7](#) illustrates the L^2 and L^∞ between the truth Γ_α and the estimation $\hat{\Gamma}_{\alpha,n}$. In this figure, and in what follows, the error of two different estimations of Γ_α are considered: the plug-in estimate of [Eq. \(3.70\)](#) with the label PI, and the probabilistic one with the label π , which will be introduced in the next section. Finally, the subscript n will indicate a quantity computed with n additional evaluations of the function.

3.5.3.c Reduction of the augmented IVPC

Recalling the definition of the probabilistic approach [Eq. \(3.68\)](#), given the GP Z , we can write an estimator of Γ_α :

$$\Gamma_\alpha^\pi(\theta) = \mathbb{E}_U [\mathbb{P}_Z [\Delta_\alpha(\theta, U) \leq 0]] = \mathbb{E}_U [\mathbb{P}_Z [Z(\theta, U) - \alpha Z^*(U) \leq 0]] \quad (3.71)$$

$$= \mathbb{E}_U [\pi_\alpha(\theta, U)] \quad (3.72)$$

The set $\{(\theta, u) \mid Z(\theta, u) - \alpha Z^*(u) \leq 0\}$ has a probability of coverage written π_α , which can be computed using the CDF of the standard normal distribution Φ , because Δ_α is a GP with parameters described [Eqs. \(3.60\) to \(3.62\)](#):

$$\pi_\alpha(\theta, u) = \Phi \left(-\frac{m_{\Delta_\alpha}(\theta, u)}{\sigma_{\Delta_\alpha}(\theta, u)} \right) \quad (3.73)$$

Finally, the outer expectation can be computed in a similar fashion as in [Eq. \(3.70\)](#); given a set of i.i.d. samples $\{u_i\}_{1 \leq i \leq n_u}$ from U , we define the probabilistic estimation of Γ_α as

$$\hat{\Gamma}_\alpha^\pi(\theta) = \frac{1}{n_u} \sum_{i=1}^{n_u} \pi_\alpha(\theta, u_i) \quad (3.74)$$

Recalling [Eq. \(3.32\)](#), the variance of the probability of coverage is $\pi_\alpha(\theta, u) (1 - \pi_\alpha(\theta, u))$, and by integrating this variance over the whole space $\Theta \times \mathbb{U}$, similarly as in [Bect et al.](#)

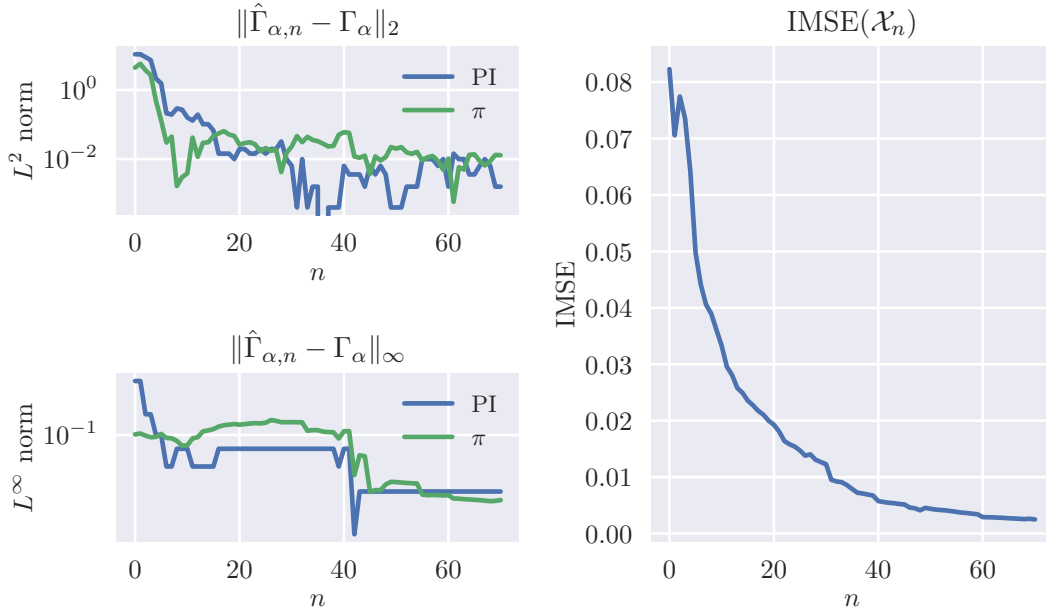


Figure 3.7 – Evolution of the L^2 and L^∞ error of the estimation of Γ_α and IMSE of the successive constructed GP. The points added are chosen by the augmented expected IMSE. The curve labelled PI refers to the plug-in estimation of Eq. (3.70), while the one labelled π refers to the probabilistic estimation, defined later in Eq. (3.74)

(2012), we can define the integrated variance of the probability of coverage IVPC.

$$\text{IVPC}(\mathcal{X}_n) = \int_{\Theta \times \mathbb{U}} \pi_\alpha(\theta, u) (1 - \pi_\alpha(\theta, u)) p_U(u) d\theta du \quad (3.75)$$

The IVPC is a measure of the uncertainty on the global coverage of the set $\{Z - \alpha Z^* \leq 0\}$, where Z is conditioned by the design $\mathcal{X}_n$. Instead of evaluating this integrated variance directly on the current design $\mathcal{X}_n$, we can once again augment the design at an (unevaluated) point (θ, u) , assuming that its evaluation is the random variable $Z(\theta, u)$:

$$\text{IVPC}(\mathcal{X}_n \cup \{((\theta, u), Z(\theta, u))\}) \quad (3.76)$$

Finally, we can define a new learning function, which is the expected IVPC with respect to the random variable $Z(\theta, u)$:

$$(\theta_{n+1}, u_{n+1}) = \arg \min_{(\theta, u) \in \Theta \times \mathbb{U}} \mathbb{E}_{Z(\theta, u)} [\text{IVPC}(\mathcal{X}_n \cup \{((\theta, u), Z(\theta, u))\})] \quad (3.77)$$

Figure 3.8 shows the evolution of the error in the estimation of Γ_α , with respect to the L^2 and L^∞ norm.

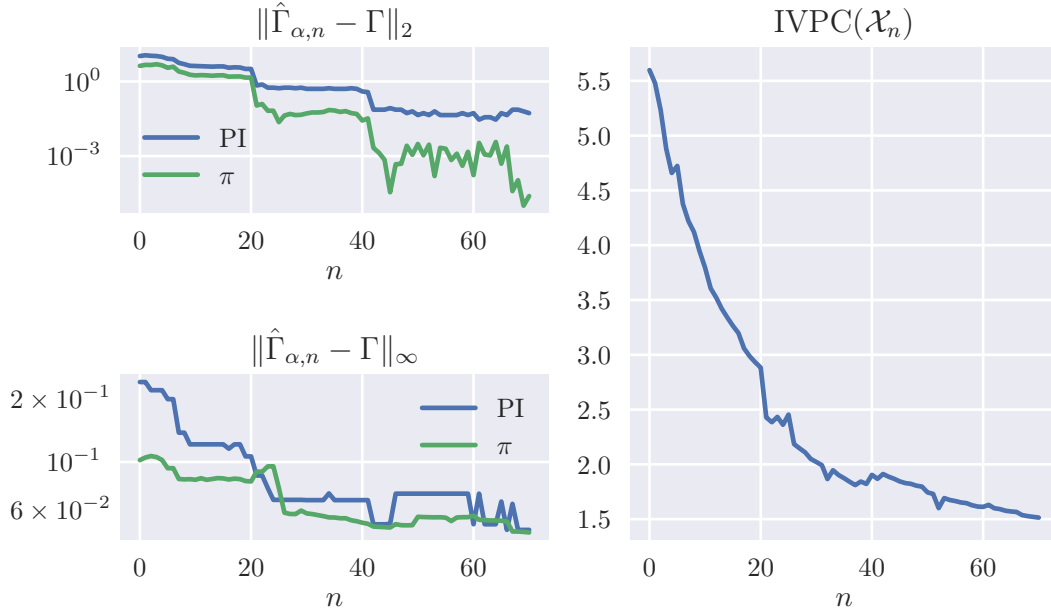


Figure 3.8 – Enriching the design according to the criterion of Eq. (3.77)

3.5.3.d Sampling in the margin of uncertainty

Sampling and translating to the source of uncertainties

We can also enrich the design by means of sampling within the margin of uncertainty defined Section 3.4.3.b Page 77, introduced in Echard et al. (2011); Schöbi et al. (2017); Razaaly (2019); Razaaly and Congedo (2020). The probability of coverage, defined using Δ_α is

$$\pi_\alpha(\theta, u) = \mathbb{P}_Z [\Delta_\alpha \leq 0] = \Phi \left(-\frac{m_{\Delta_\alpha}(\theta, u)}{\sigma_{\Delta_\alpha}(\theta, u)} \right) \quad (3.78)$$

the margin of uncertainty of level η is by definition

$$\mathbb{M}_\eta = \left\{ (\theta, u) \mid \frac{\eta}{2} \leq \pi_\alpha(\theta, u) \leq 1 - \frac{\eta}{2} \right\} \quad (3.79)$$

The margin of uncertainty of the random set $\{\Delta_\alpha \leq 0\}$ has been represented Fig. 3.9.

Sampling in this margin of uncertainty and evaluating some of those points makes sense if the points selected help reduce the margin of uncertainty. Let us assume that the current design used to construct Z comprises n points, and let us assume that we obtained K points which represent statistically the margin of uncertainty: $\tilde{x}_{n+i}$ for $1 \leq i \leq K$. Those points being in our case the centroids produced by the KMeans algorithms.

As mentioned before in Eq. (3.62) on Page 82, the uncertainty on the acceptability of those points can be reduced either by evaluating the function at this point (thus reducing σ_Z^2), or by improving the knowledge of the conditional minimiser (reducing $\alpha^2 \sigma_{Z^*}^2$). For each centroid $\tilde{x}_{n+i} = (\tilde{\theta}_{n+i}, \tilde{u}_{n+i})$, we can then adjust its θ -component if the uncertainty

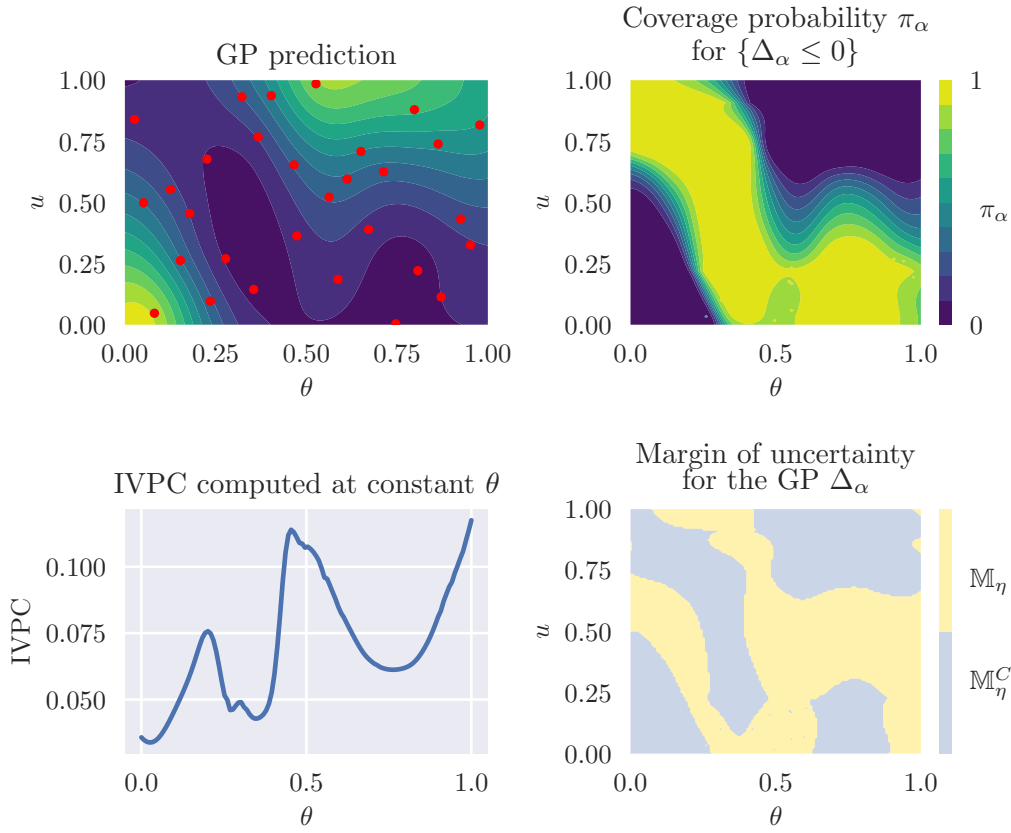


Figure 3.9 – Probability of coverage and Margin of uncertainty for $\{\Delta_\alpha \leq 0\}$ at level $\eta = 0.05$. The variance of the probability of coverage is integrated on $\mathbb{U}$, giving the IVPC at constant θ

on the conditional minimiser is large in order to reduce $\alpha^2 \sigma_{Z^*}^2$ (instead of σ_Z^2). This is done by introducing an adjustment function:

$$\text{adjust}(\theta, u) = \begin{cases} (\theta, u) & \text{if } \sigma_Z^2(\theta, u) \geq \alpha^2 \sigma_{Z^*}^2(u) \\ (\max_{\theta \in \Theta} \mathbb{E}_Z [[Z(\theta, u) - m_{Z^*}(u)]_+], u) & \text{otherwise} \end{cases} \quad (3.80)$$

where $m_{Z^*}(u) = \min_{\theta \in \Theta} m_Z(\theta, u)$ as defined Eq. (3.56). In this case, the adjusted value can be seen as an EGO iteration on $\Theta \times \mathbb{U}$. The adjusted centroids are then expected to reduce the most the uncertainty at the original centroids, given by the KMeans algorithm.

Hierarchical clustering for a second adjustment

However, as we can see on the left plot of Fig. 3.11, all centroids are adjusted independently, and since only their θ component is changed, the adjusted centroids may end up very close to each other. The issue is twofold: the evaluation of those points brings redundant information to the problem, and thus we do not take fully advantage of the batch evaluations. Moreover when adding close points to the design, the proximity

can lead to an ill-conditioned covariance matrix of the GP and in turn to numerical problems in the estimation of its hyperparameters.

In order to identify the nearby points, we can perform once again a clustering procedure such as Hierarchical clustering (Nielsen, 2016), as the number of points to consider is much more limited, and only a labelling is needed.

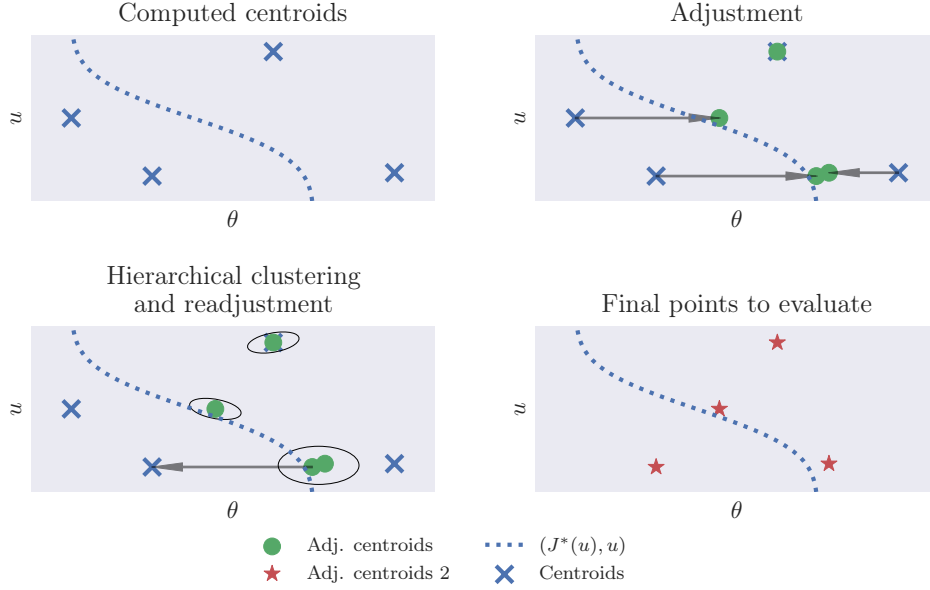


Figure 3.10 – Principle of the two consecutive adjustments: The first adjustment is done in order to reduce the most the uncertainty, and the second to select points far enough from each other

For each point in each cluster, as they are in close vicinity, we can assume that their uncertainties on the conditional minimiser $\alpha^2 \sigma_{Z^*}^2$ are sensibly the same. So instead, we are going to make the adjustment described Eq. (3.80) only for the point which presents the smallest σ_Z^2 . The whole algorithm is described Algorithm 3, and the double adjustment is described schematically in Fig. 3.10. In the end, the points added to the design are those labeled “Adjusted centroids 2”.

We apply an iteration of this method on the real function, in Fig. 3.11. At first, the $K = 10$ centroids obtained by KMeans are represented with the red stars. As the measured uncertainty is due to the conditional minimisers, all the centroids are adjusted, yielding the blue stars near the locus of the conditional minimisers. At $u \approx 0.6$, $u \approx 0.4$ and $u \approx 0.15$, the hierarchical clustering step detects close points, and a second adjustment is made, giving (back) the point at approximately $(0.8, 0.6)$, and $(0.95, 0.15)$ for instance.

On Fig. 3.12 is shown the evolution of the norm of the difference between the estimated $\hat{\Gamma}_{\alpha,n}$ (after having added n points to the design) and the true value Γ_α , when adding simultaneously $K = 5$ and $K = 10$ points.

Algorithm 3 Enrichment of the design using sampling to reduce the margin of uncertainty of $\{\Delta_\alpha \leq 0\}$

Require: Initial design $\mathcal{X}_0$,

Require: number of points to evaluate each iterations K , number of samples N

Require: Z , a GP using the design $\mathcal{X}_0$

Require: Probability of coverage π_α , and sampling pdf $\kappa_{\mathbb{M}} = \mathbb{1}_{\mathbb{M}_\eta}$

$n \leftarrow 0$

while (stopping criterion not met) **or** (evaluation budget not reached) **do**

 Sample N points $S = \{x_i^s\}_{1 \leq i \leq N}$ in $\mathbb{M}_\eta \subset \Theta \times \mathbb{U}$ using $\kappa_{\mathbb{M}}$

 Get the K centroids $\{x_i^{\text{center}}\}_{1 \leq i \leq K}$ using KMeans on S

for $i = 1$ to K **do**

$\tilde{x}_{n+i} = (\theta_{n+i}, u_{n+i}) = \arg \min_{x \in S \setminus \{\tilde{x}_{n+1}, \dots, \tilde{x}_{n+i-1}\}} \|x - x_i^{\text{center}}\|$

end for

 Hierarchical clustering of $\{\text{adjust}(\tilde{x}_{n+i})\}$ for $1 \leq i \leq K$

 Define $\{x_{n+i}\}_{1 \leq i \leq K}$ as the readjusted clusters using hierarchical clustering

 Evaluate $f(x_{n+i}) = J(\theta_{n+i}, u_{n+i})$ for $1 \leq i \leq K$

$\mathcal{X}_{n+K} \leftarrow \mathcal{X}_n \cup \{(x_{n+1}, f(x_{n+1})), \dots, (x_{n+K}, f(x_{n+K}))\}$

 Condition the GP according to $\mathcal{X}_{n+K}$

$n \leftarrow n + K$

end while

The added points are shown [Fig. 3.13](#), and the margin of uncertainty $\mathbb{M}_\eta$ after the $n = 70$ evaluations. We can see that the added point are distributed either toward the conditional minimisers, or toward the frontier of $\{(\theta, u) \mid J(\theta, u) = \alpha J^*(u)\}$, as expected.

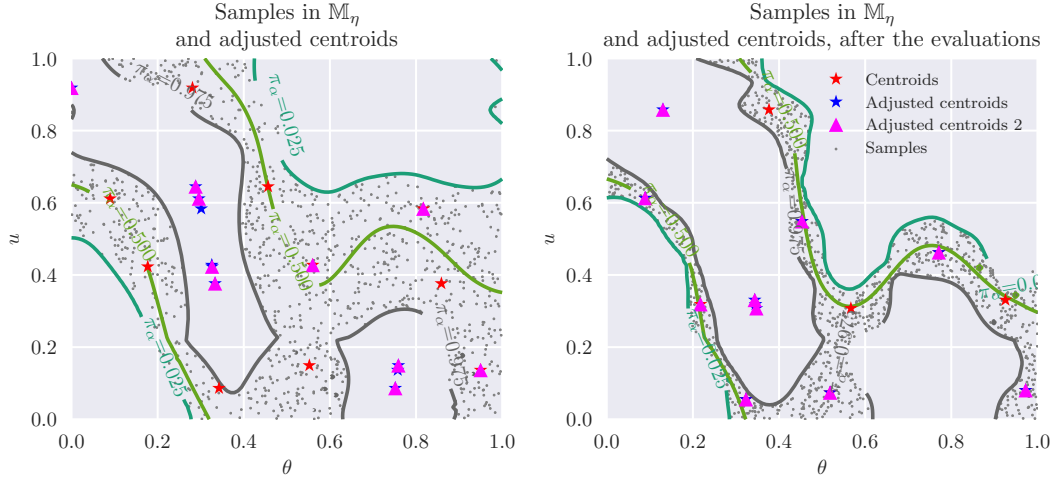


Figure 3.11 – Sampling in the margin of uncertainty, and adjustment of centroids. The centroids of the KMeans algorithm are in red, the adjusted centroids are in blue, and the adjusted centroids after the hierarchical clustering step are in magenta

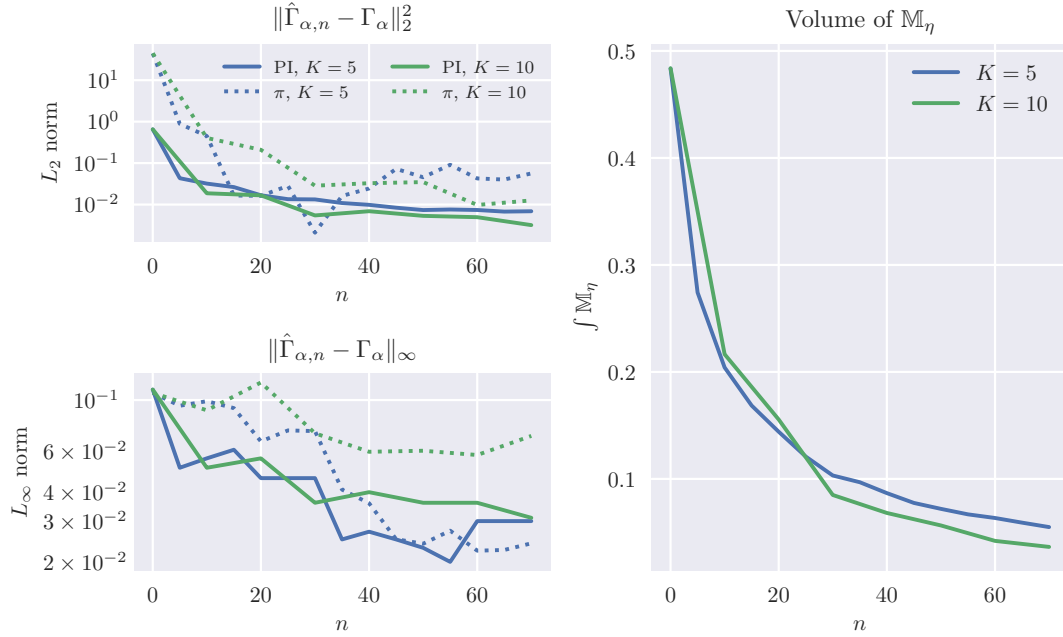


Figure 3.12 – Evolution of the error with the number of added points in the estimation of Γ_α when selecting points using the sampling based methods described Algorithm 3. π refers to the estimation of Γ using a probabilistic approach, while PI is the estimation using a plug-in approach

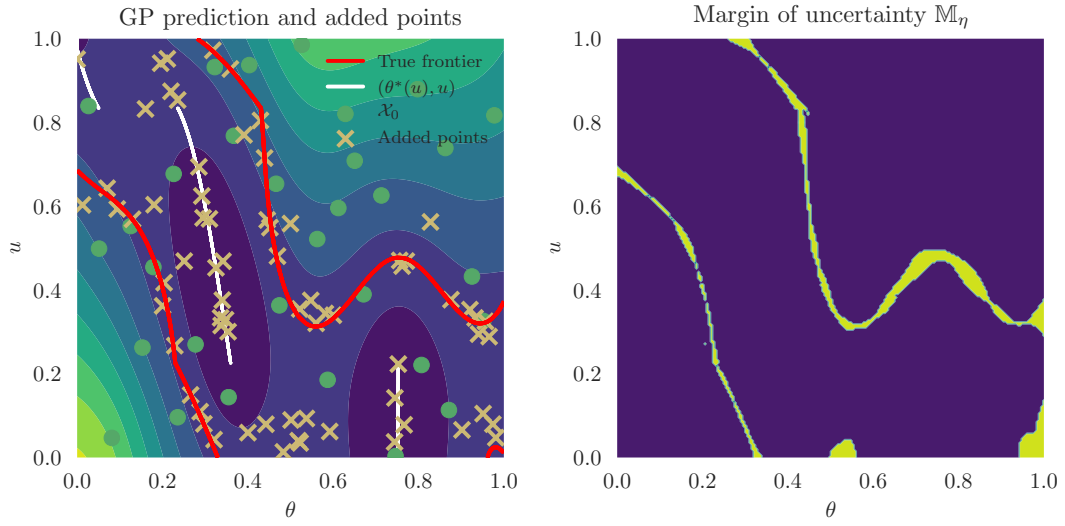


Figure 3.13 – GP prediction and final experimental design (left) and resulting margin of uncertainty $\mathbb{M}_\eta$

3.5.4 Optimisation of the quantile of the relative regret

In the previous section, we introduced methods in order to improve the estimation of $\Gamma_\alpha(\theta) = \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)]$, the probability of exceeding a threshold $\alpha \geq 1$ chosen and fixed beforehand. However, once properly estimated and optimised, the maximal probability $\max \Gamma_\alpha$ may not ensure a large enough reliability in terms of relative-regret, especially if the threshold has been chosen arbitrarily.

As mentioned in [Section 2.4.4](#), we can instead look to tune α such that $\max \Gamma_\alpha$ is large enough, *i.e.* exceeds a specified level of confidence p . Let us recall [Eq. \(2.69\)](#) which defines α_p as the smallest threshold giving a maximal probability of at least p :

$$\alpha_p = \inf_{\alpha \geq 1} \{ \max_{\theta \in \Theta} \Gamma_\alpha(\theta) \geq p \} \quad (3.81)$$

As $J^*(u) > 0$ for all u , we can define $q_p(\theta)$ as the quantile of level p of the ratio $J(\theta, U)/J^*(U)$:

$$q_p(\theta) = Q_U \left(\frac{J(\theta, U)}{J^*(U)}; p \right) \iff p = \mathbb{P}_U \left[\frac{J(\theta, U)}{J^*(U)} \leq q_p(\theta) \right] = \Gamma_{q_p(\theta)}(\theta) \quad (3.82)$$

where $Q_U(\cdot; p)$ is the quantile function of order p with respect to the random variable U . α_p verifies then

$$\alpha_p = \min_{\theta \in \Theta} q_p(\theta) \quad (3.83)$$

The relation between Γ_{α_p} and q_p is illustrated [Fig. 3.14](#).

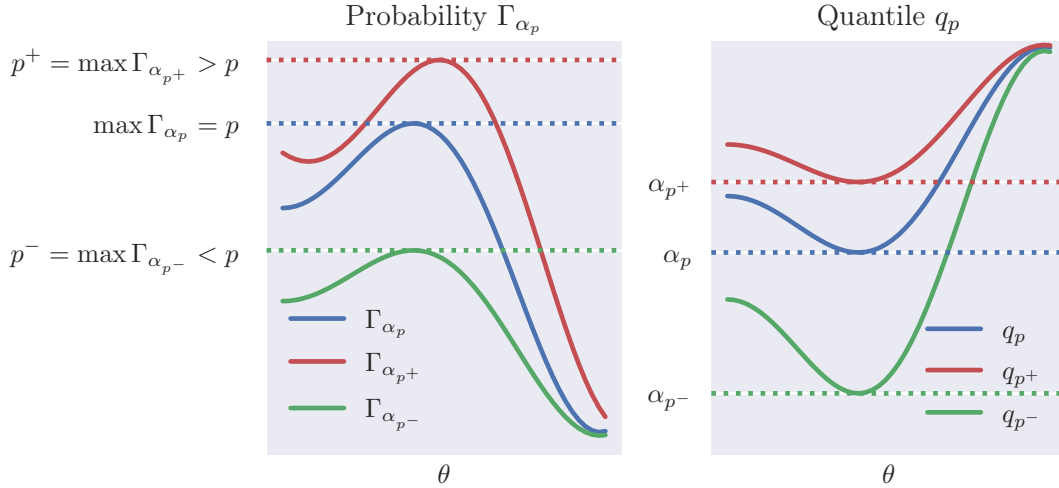


Figure 3.14 – Relation between the probability Γ_α and its maximum, and the quantile function of the ratio q_p

Given [Eq. \(3.83\)](#), we are going to focus on the estimation of the quantiles of order p of the ratio J/J^* defined [Eq. \(3.82\)](#) as $q_p(\theta)$ for $\theta \in \Theta$.

Numerically speaking and given N i.i.d. samples $x_i \sim X$, where X is a real-valued random variable, the estimation of a quantile is quite easy. A common estimation of the quantile of order p of X is $x_{([Np])}$, where $x_{(i)}$ is the i th order statistic of samples (*i.e.* the i th smallest value of the set of samples), and $[\cdot]$ is the rounding operator.

As done for $\Delta_\alpha = Z - \alpha Z^*$, and the estimation of Γ_α , using the GP Z constructed on J , we are going to approximate the ratio J/J^* , in order to have an estimation which depends on the properties of the surrogate Z . In [Section 3.5.4.a](#), we are going to derive some properties of the ratio Z/Z^* , before introducing a 1-step criterion in [Section 3.5.4.b](#), and a sampling method in [Section 3.5.4.c](#).

3.5.4.a Lognormal approximation of the ratio

Let $\theta \in \Theta$ and $u \in \mathbb{U}$. The true value of the ratio $J(\theta, u)/J^*(u)$ is modelled as the random variable $Z(\theta, u)/Z^*(u)$, where the joint distribution of $Z(\theta, u)$ and $Z^*(u)$ is

$$\begin{bmatrix} Z(\theta, u) \\ Z^*(u) \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} m_Z(\theta, u) \\ m_{Z^*}(u) \end{bmatrix}; \begin{bmatrix} \sigma_Z^2(\theta, u) & \rho \sigma_Z(\theta, u) \sigma_{Z^*}(u) \\ \rho \sigma_Z(\theta, u) \sigma_{Z^*}(u) & \sigma_{Z^*}^2(u) \end{bmatrix} \right) \quad (3.84)$$

where ρ is the correlation coefficient $\rho = \frac{\text{Cov}(Z(\theta, u), Z^*(u))}{\sigma_Z(\theta, u) \sigma_{Z^*}(u)}$. Under the condition that $m_{Z^*}(u)/\sigma_{Z^*}(u)$ is large, *i.e.* that $Z^*(u) \geq 0$ with high probability, we can make the approximation that at each (θ, u) , the ratio is lognormally distributed (see [Section A.1](#)):

$$\frac{Z(\theta, u)}{Z^*(u)} \sim \log \mathcal{N} \left(\log \frac{m_Z(\theta, u)}{m_{Z^*}(u)}; \frac{\sigma_Z^2(\theta, u)}{m_{Z^*}^2(u)} + \frac{\sigma_{Z^*}^2(u)}{m_{Z^*}^2(u)^2} - 2\rho \frac{\sigma_Z(\theta, u) \sigma_{Z^*}(u)}{m_Z(\theta, u) m_{Z^*}(u)} \right) \quad (3.85)$$

$$\sim \log \mathcal{N}(m_\Xi(\theta, u), \sigma_\Xi^2(\theta, u)) \quad (3.86)$$

By definition of the lognormal distribution, the logarithm of Z/Z^* can then be approximated by a normally distributed random variable, denoted $\Xi(\theta, u)$:

$$\Xi(\theta, u) = \log \left(\frac{Z(\theta, u)}{Z^*(u)} \right) \sim \mathcal{N}(m_\Xi(\theta, u), \sigma_\Xi^2(\theta, u)) \quad (3.87)$$

The mean value of the lognormally distributed random variable is used as an estimate of the ratio:

$$\mathbb{E}_Z \left[\frac{Z(\theta, u)}{Z^*(u)} \right] = \frac{m_Z(\theta, u)}{m_{Z^*}(u)} \exp \left(\frac{1}{2} \sigma_\Xi^2(\theta, u) \right) \quad (3.88)$$

and then the plug-in estimation of the quantile of order p is:

$$q_p^{\text{PI}}(\theta) = Q_U \left(\frac{m_Z(\theta, U)}{m_{Z^*}(U)} \exp \left(\frac{1}{2} \sigma_\Xi^2(\theta, U) \right), p \right) \quad (3.89)$$

$$= Q_U \left(\exp \left(m_\Xi(\theta, u) + \frac{1}{2} \sigma_\Xi^2(\theta, U) \right), p \right) \quad (3.90)$$

And based on a set of samples $\{u_i\}_{1 \leq i \leq n_u}$, the estimated quantile becomes

$$\hat{q}_p^{\text{PI}}(\theta) = \left(\exp \left(m_\Xi(\theta, u) + \frac{1}{2} \sigma_\Xi^2(\theta, u) \right) \right)_{([n_u p])} \quad (3.91)$$

We can use directly the Gaussian formulation of Ξ in order to derive strategies of enrichment. In the following, we will set the confidence level to $p = 0.95$.

3.5.4.b Reduction of the IMSE

As previously done for the random process $Z - \alpha Z^*$, we can look to reduce at each iteration the integrated prediction variance of the ratio, which is σ_{Ξ}^2 by defining the IMSE

$$\text{IMSE}(\mathcal{X}_n) = \int_{\Theta \times \mathbb{U}} \sigma_{\Xi}^2(\theta, u) d\theta du \quad (3.92)$$

where σ_{Ξ}^2 is constructed using Z/Z^* and the design $\mathcal{X}_n$.

We can then look to improve the expected prediction error on the log-ratio described Eq. (3.87), by minimising the augmented IMSE:

$$(\theta_{n+1}, u_{n+1}) = \arg \min_{(\theta, u) \in \Theta \times \mathbb{U}} \mathbb{E}_{Z(\theta, u)} [\text{IMSE}(\mathcal{X}_n \cup \{(\theta, u), Z(\theta, u)\})] \quad (3.93)$$

meaning that σ_{Ξ} is computed using Eq. (3.85), and the GP augmented by the couple $((\theta, u), Z(\theta, u))$. The result of this adaptive strategy is illustrated Fig. 3.15. Along with $\hat{q}_{p,n}^{\text{PI}}$, we plotted Monte-Carlo estimation of the quantile: $\hat{q}_{p,n}^{\text{MC}}$, which has been obtained as a Monte-Carlo estimation of $\mathbb{E}_Z \left[Q_U \left(\frac{Z(\theta, u)}{Z^*(u)}; p \right) \right]$ for comparison. We can see that this method is able to reduce steadily the error on the estimation of the quantile of order $p = 0.95$ in this case. The large discontinuities in the error are probably due to the evaluation of points in regions previously unexplored, and/or the effect of a significant change in the hyperparameters of the GP.

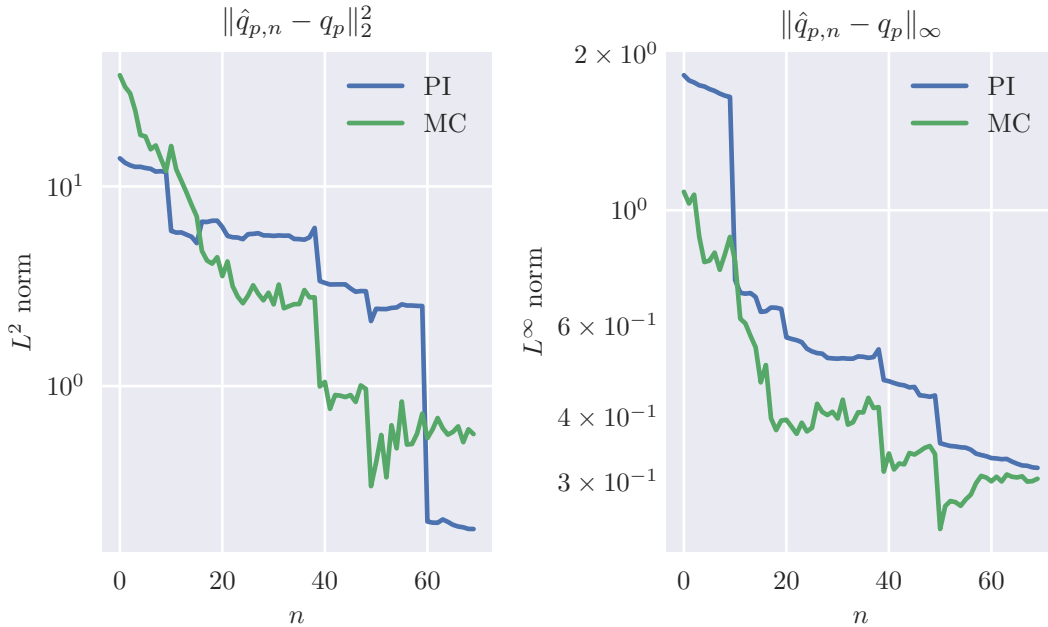


Figure 3.15 – Evolution of the error of the estimation when reducing the augmented IMSE of the log-ratio Ξ

Each step of this method requires the optimisation of the expected IMSE described Eq. (3.93), which can become expensive. We can also define a sampling-based method, which is based on plausible values of the quantile of order p .

3.5.4.c Sampling-based method, adaptation of the QeAK-MCS

We introduced previously a 1-step strategy, which looks to reduce the prediction error of the ratio Z/Z^* . In Razaaly (2019); Razaaly et al. (2020), the author introduce a sampling-based method, named QeAK-MCS, for the estimation of extreme quantiles, that we can adapt for the estimation of the quantiles of the ratio. Recalling that

$$\Xi(\theta, u) = \log \frac{Z(\theta, u)}{Z^*(u)} \sim \mathcal{N}(m_\Xi(\theta, u), \sigma_\Xi^2(\theta, u)) \quad (3.94)$$

we can get bounds of the true value of the quantile of order p at a given $\theta \in \Theta$ by using the quantiles of $\Xi(\theta, u)$:

$$\log q_p^+(\theta) = Q_U(m_\Xi(\theta, U) + k\sigma_\Xi(\theta, U), p) \quad (3.95)$$

$$\log q_p^-(\theta) = Q_U(m_\Xi(\theta, U) - k\sigma_\Xi(\theta, U), p) \quad (3.96)$$

where k is a quantile of the standard Gaussian random variable. Let us define $\tilde{\theta} = \arg \min_{\theta \in \Theta} \hat{q}_p(\theta)$, that is the value that minimises the quantile of order p of the plug-in estimate of the ratio, and $\hat{\alpha}_p = \hat{q}_p(\tilde{\theta})$.

By linearly discretizing the interval $[q_p^-(\tilde{\theta}), q_p^+(\tilde{\theta})]$, we can define $\mathbf{q}_l$ for $1 \leq l \leq K_q$, with $\mathbf{q}_1 = q_p^-$ and $\mathbf{q}_{K_q} = q_p^+$. The $\{\mathbf{q}_l\}_{1 \leq l \leq K_q}$ are then expected to cover the plausible range of values that the sought quantile $\alpha_p = \min q_p(\theta)$ may take.

For each of those K_q quantiles, we can look to sample $K_{\mathbb{M}}$ points in the margin of uncertainty of the log-ratio, which is defined as

$$\mathbb{M}_\eta(\mathbf{q}_l) = \left\{ (\theta, u) \mid \frac{m_\Xi(\theta, u) - \log \mathbf{q}_l}{\sigma_\Xi(\theta, u)} < k \text{ and } -\frac{m_\Xi(\theta, u) - \log \mathbf{q}_l}{\sigma_\Xi(\theta, u)} < k \right\} \quad (3.97)$$

$$= \{(\theta, u) \mid m_\Xi(\theta, u) - k\sigma_\Xi(\theta, u) < \log \mathbf{q}_l < m_\Xi(\theta, u) + k\sigma_\Xi(\theta, u)\} \quad (3.98)$$

with k an appropriate quantile of the standard Gaussian random variable. Equivalently, by defining the coverage probability of the set $\{\log \frac{Z(\theta, u)}{Z^*(u)} \leq \log \mathbf{q}_l\}$, which is $\pi_{\Xi, l}(\theta, u) = \Phi\left(-\frac{m_\Xi(\theta, u) - \log \mathbf{q}_l}{\sigma_\Xi(\theta, u)}\right)$,

$$\mathbb{M}_\eta(\mathbf{q}_l) = \left\{ (\theta, u) \mid \frac{\eta}{2} \leq \pi_{\Xi, l}(\theta, u) \leq 1 - \frac{\eta}{2} \right\} \quad (3.99)$$

In other words, this margin of uncertainty is the set of points whose confidence interval of level η of the log-ratio comprises the targeted value $\log \mathbf{q}_l$.

In the end, we have $K = K_q K_{\mathbb{M}}$ points to add to the design, that we can adjust as in Algorithm 3, in order to reduce the uncertainty on the value of the conditional minimum when required. This procedure is illustrated Fig. 3.16, while on Fig. 3.17, the evolution of the error on the estimation of q_p is shown.

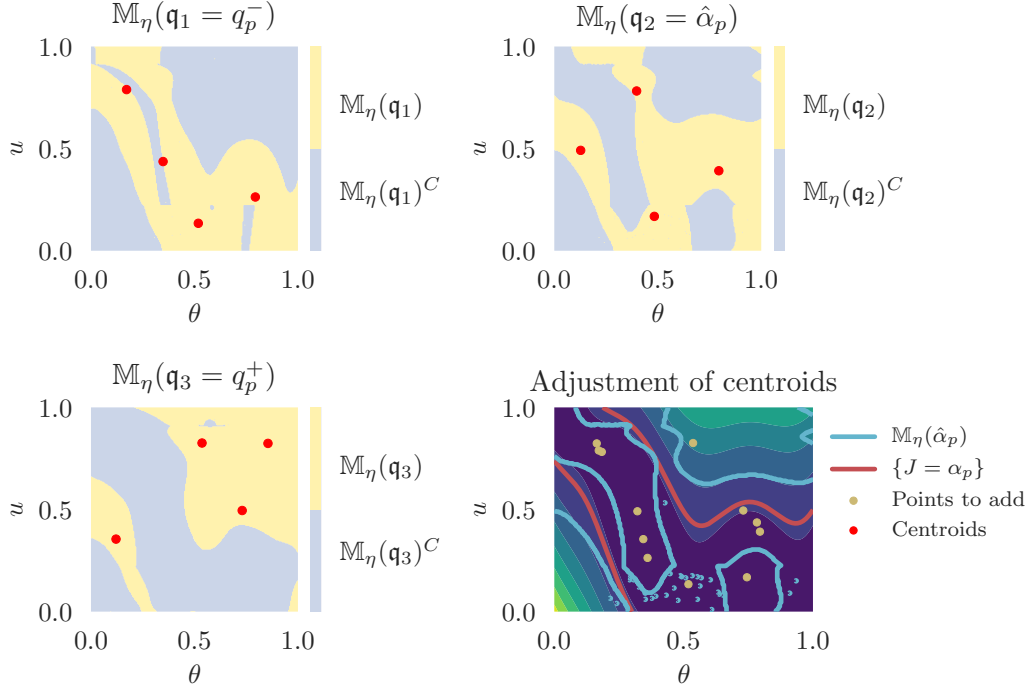


Figure 3.16 – One iteration of the QeAK-MCS procedure, with $K_q = 3$ and $K_{\mathbb{M}} = 4$

We can see that we can reduce globally the estimation error, but it seems slower than the augmented IMSE, as shown [Fig. 3.15](#).

An issue that may arise is that when the GP is accurate enough, q_p^- and q_p^+ become close, hence the different margin of uncertainties produce centroids close to each others, which may lead to numerical difficulties. A crude way to deal with this technical issue is to discard randomly all but one point when some are close to each other, as done with hierarchical clustering.

3.6 Partial conclusion

In this chapter, we introduced Gaussian Processes as surrogates for the expensive-to-evaluate objective function J . This metamodel, once constructed using an initial design of points on $\Theta \times \mathbb{U}$, can then be used to compute the robust estimators introduced in [Chapter 2](#).

However, due to the limited size of the initial design, and the impossibility to perform exhaustive computations of the objective function, this surrogate can be inaccurate in some regions, leading to a bad estimation, and thus a questionable calibration. Using the properties of the GP, we can construct adaptive strategies in order to enrich the design of experiment sequentially, where the points chosen improve significantly the various estimations, and in this thesis, the estimation of the regret-based estimators.

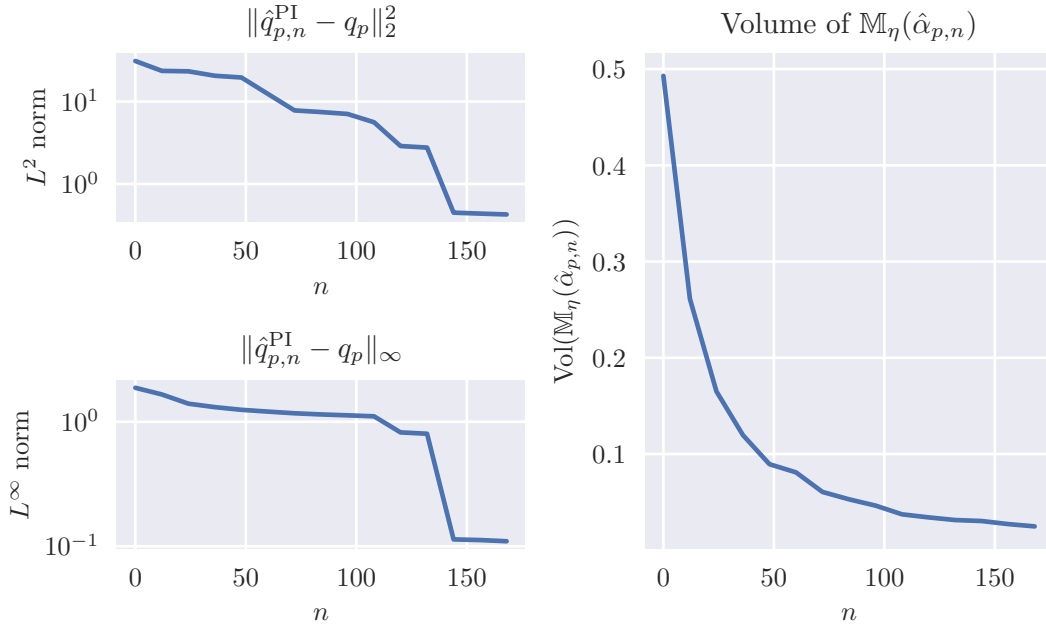


Figure 3.17 – Evolution of the estimation error as $K = K_q K_{\mathbb{M}} = 12$ are added each iterations

Several improvements can be considered. First, all the methods introduced here focus on the improvement of the estimation of a function, which will be optimised in order to get the estimator. The adaptive strategy could be adapted to take into account the subsequent optimisation. Similarly as in [Janusevskis and Le Riche \(2010\)](#), we could develop strategies in two parts: first, identify the point in Θ which has the most potential to be the wanted estimator, and second, find the point in the joint space $\Theta \times \mathbb{U}$ which reduces the most the uncertainty on the candidate.

Another issue encountered is that Gaussian Processes are suited to approximate functions of a moderate number of variables, because the optimisation of the hyperparameters becomes increasingly difficult. In order to apply the algorithms based on GP, the dimension may first have to be reduced. Moreover, we did not take advantage of the gradient information, which may be available through the adjoint method and automatic differentiation tools, as has been done in [Bouhlef and Martins \(2019\)](#); [Laurent et al. \(2019\)](#); [Miranda et al. \(2016\)](#); [Marmin et al. \(2015\)](#).

In the next chapter, we will study the problem of robust calibration of a realistic numerical model, using some of the methods we introduced in this chapter.

* * *

CHAPTER 4

APPLICATION TO THE NUMERICAL COASTAL MODEL CROCO

Contents

4.1	Introduction	101
4.2	CROCO and bottom friction modelling	101
4.2.1	Parameters and configuration of the model	102
4.2.2	Modelling of the bottom friction	103
4.2.3	Definition of the control and environmental parameters	103
4.2.3.a	Bottom roughness and sediments size	103
4.2.3.b	Tidal modelling and uncertainties	105
4.3	Deterministic calibration of the bottom friction	106
4.3.1	Twin experiment setup	107
4.3.2	Cost function definition	107
4.3.3	Gradient-descent optimisation	107
4.4	Sensitivity analysis of the objective function	112
4.4.1	Global Sensitivity Analysis: Sobol' indices	112
4.4.2	SA of the objective function for the calibration of CROCO	113
4.4.2.a	SA on the regions defined by the sediments	113
4.4.2.b	SA on the tide components	113
4.5	Robust Calibration of the bottom friction	114
4.5.1	Objective function and global minimum	116
4.5.2	Conditional minimums and conditional minimisers	118
4.5.3	Relative-regret based estimators	121
4.5.3.a	Optimisation of the probability of exceeding a threshold	121
4.5.3.b	Optimisation of the quantile of the relative-regret	125

4.6	Partial conclusion	128
------------	---------------------------	------------

4.1 Introduction

In this chapter, we will study the problem of calibration under uncertainties of the bottom friction of the ocean bed off the coast of France, from the English Channel to the Bay of Biscay. This will be realised using a realistic numerical model based on CROCO¹ (Coastal and Regional Ocean COmmunity model).

Since the bottom friction depends directly on the size of the asperities on the ocean bed, the length scales involved in this process are way smaller than the scales of the computational grid. In consequence, the numerical model does not solve the equations of motions of the fluid around those asperities. Instead, the dissipation coming from the associated rugosity is parametrized at every cell of the mesh.

The bottom friction has been identified as a crucial parameter that, if ill-specified, limits the accuracy of the predictions (Sinha and Pingree, 1997; Kreitmair et al., 2019), especially in shallow regions; consequently, there has been an effort to control this parameter in various studies, for instance in Das and Lardner (1992, 1991); Boutet (2015). We will detail how bottom friction affects the oceanic circulation in Section 4.2.2, in order to get a first insight on the regions that may influence the most the calibration.

The deterministic problem of calibration will then be addressed in Section 4.3, by first defining the objective function and the input space. We will then calibrate the model without external uncertainties using adjoint-based gradient, in high-dimension ($\approx 15\,000$).

However for such problems, as the parameter may be spatially distributed and thus high-dimensional, any estimation procedure may become quickly expensive. In consequence, instead of considering each grid cell individually, we will segment the geographical input space in different independent regions, which are based on the type of sediments listed at the bottom of the water.

In this problem of calibration we will assume that the uncertainties take the form of an environmental parameter which perturbs the amplitude of some tidal constituents. In order to quantify the influence of each of the sediment-based regions and the influence of each of the components of the environmental variable, we will carry a global sensitivity analysis in Section 4.4.

Finally based on this study, we will reduce significantly the input space, and then apply some of the methods proposed in the previous chapter, in order to get a robust estimation of the bottom friction using Gaussian processes in Section 4.5.

4.2 CROCO and bottom friction modelling

CROCO¹ (Coastal and Regional Ocean COmmunity model) is a numerical model that describes the motion of the ocean by solving the *primitive equations*, which are simplified versions of the Navier-Stokes equations, taking into account the particular scales at play at the surface of the Earth. CROCO has been developed upon ROMS_AGRIF (Regional Ocean Modeling System, Adaptive Grid Refinement in Fortran Debreu et al. (2012)), and

¹CROCO and CROCO_TOOLS are provided by <https://www.croco-ocean.org>

is designed to be coupled with other modelling systems, such as atmospheric, biological or ecosystem models.

4.2.1 Parameters and configuration of the model

The configuration used in this thesis is based on the one used in [Boutet \(2015\)](#). The spatial domain ranges from 9°W to 1°E (comprising 139 internal points in this direction) and from 43°N to 51°N (164 in this direction), and spans most of the Bay of Biscay, the English Channel and the eastern part of the Celtic Sea. The resolution is 1/14°, which leads to a mesh size between 5 km and 6 km. The bathymetry map and the spatial domain is shown [Fig. 4.1](#). The ocean can be split roughly in two regions, based on its depth : the region near the coasts which corresponds to the continental shelf, where the water depth is less than 200 m, and the offshore region of the Bay of Biscay, where the depth is closer to 5000 m.

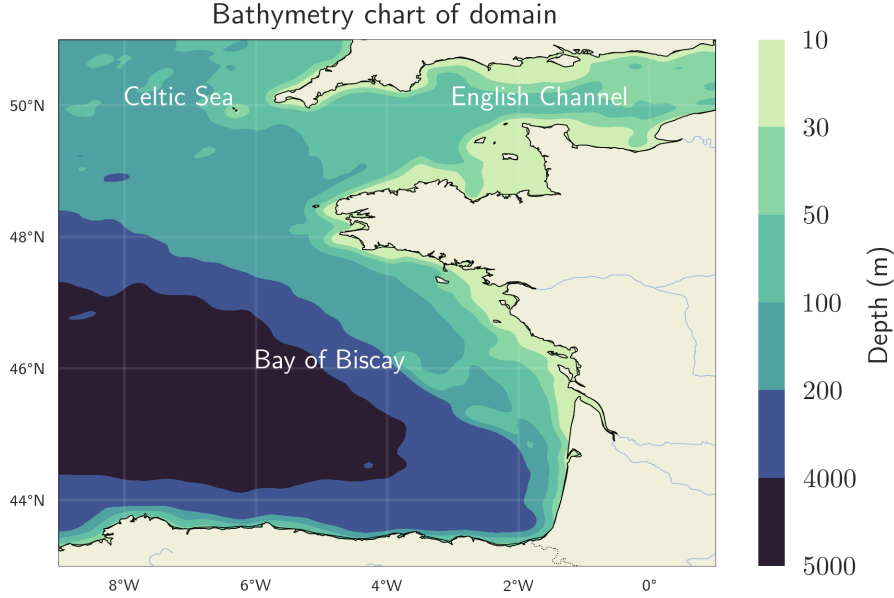


Figure 4.1 – Bathymetry used in CROCO, and geographical landmarks. The continental shelf corresponds roughly to the area with depth less than 200 m (green hue), while the abyssal plain has a depth larger than 4000 m (blue hue)

CROCO can solve the fluid motion equations in 3D, but in this configuration, solves the rotating shallow water equations instead, which are obtained by vertically averaging the primitive equations, leading to:

$$\begin{cases} \frac{\partial \mathbf{v}}{\partial t} + (\mathbf{v} \cdot \nabla) \mathbf{v} + 2\mathbf{\Omega} \wedge \mathbf{v} &= -g\nabla H + \frac{\tau_b}{\rho H} + F \\ \frac{\partial \zeta}{\partial t} + \nabla(H \cdot \mathbf{v}) &= 0 \end{cases} \quad (4.1)$$

where $\mathbf{v} = (v_x, v_y)$ is the velocity field of the fluid, $\mathbf{\Omega}$ is the rotational angular vector of Earth. H , the total water column height and ζ , the free-surface height (relative to the

geoid) satisfy the relation $H = \zeta + b$ where b is the bathymetry. The effect of the bottom topography and the friction are modelled using g the gravitational constant, ρ the fluid density, and τ_b the shear stress at the bottom. Finally, F represents the external forcing of the model. The forcing due to the tides is done at the boundaries. The bottom friction affects the circulation through τ_b , and different parameterizations of this stress can be derived.

4.2.2 Modelling of the bottom friction

In CROCO, the shear stress at the bottom is modelled using a quadratic drag coefficient C_d :

$$\tau_b = -C_d \|\mathbf{v}_b\| \mathbf{v}_b \quad (4.2)$$

where $\mathbf{v}_b$ is the velocity at the bottom, so in the case of the Shallow Water equations, $\mathbf{v}_b = \mathbf{v}$. The drag coefficient can in turn be formulated as a function of the water column height and the *bottom roughness* z_b by assuming a logarithmic profile of the velocity at bottom (a derivation can be found in [Le Bars et al. \(2010\)](#) for instance)

$$C_d = \left(\frac{\kappa}{\log\left(\frac{H}{z_b}\right) - 1} \right)^2 \quad (4.3)$$

where κ is the Von Kármán constant, usually taken equal to 0.41. The bottom roughness z_b , or *rugosity* in this document, can be interpreted as the size of the turbulent layer at the bottom, induced by the asperities of the sediments. [Boutet \(2015\)](#) shows that in a calibration context, controlling the rugosity z_b yields better result than controlling the drag coefficient C_d due to the influence of the water column height H . On [Fig. 4.2](#) is shown the drag coefficient C_d as a function of the roughness z_b of the ocean floor, for different heights of the water column H .

We can see that the higher the water column height, the less variation appears when adjusting the bottom roughness z_b . Considering the physical properties of the bottom friction and the types of sediments, it can be expected that the English Channel, and at a lesser extent the rest of the continental shelf are the areas which are the most influential for the calibration.

We are now going to develop more precisely which inputs we are going to consider for the numerical problem of calibration.

4.2.3 Definition of the control and environmental parameters

4.2.3.a Bottom roughness and sediments size

In this work, we are going to use a twin experiment setup: the calibration will be performed with respect to some observation generated using CROCO. This observation $y \in \mathbb{R}^{N_{\text{obs}}}$ is computed using a specific configuration of the forward model, meaning that we are going to define a *truth* value for the bottom friction.

To do so, we are going to make the assumption that the size of the turbulent layer at the bottom is equal to the size of the sediments there, so the rugosity is directly linked to

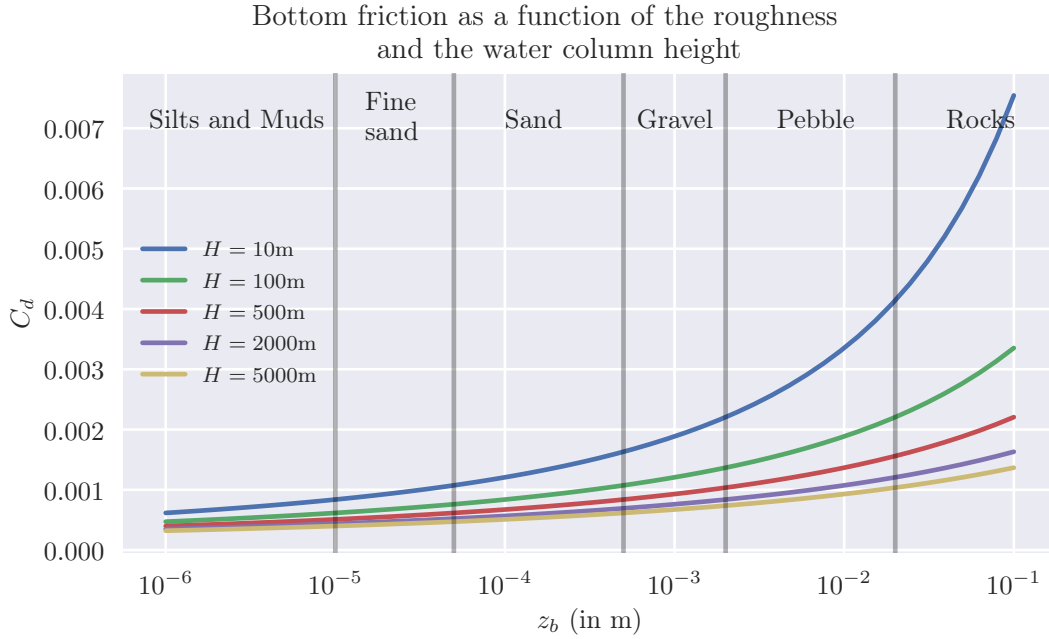


Figure 4.2 – Drag coefficient C_d as a function of the column water height and the roughness at the bottom

the type of sediment found on the ocean bed. Table 4.1 presents a coarse classification, along with the typical size of the sediments that can be found, that will serve as *truth value*: z_b^{truth} .

Code	Description	Size of the majority of particles	z_b^{truth}
R	Rock	Larger	50 mm
C	Pebble	>20 mm	25 mm
G	Gravel	[20 mm, 2 mm]	7 mm
S	Sand	[2 mm, 0.5 mm]	1 mm
SF	Fine Sand	[0.5 mm, 0.05 mm]	1.5×10^{-1} mm
Si	Silt	[0.05 mm, 0.01 mm]	2×10^{-2} mm
V	Muds	< 0.05 mm	2×10^{-2} mm

Table 4.1 – Type of sediments and size of the majority of particles for each type of sediment. Data source: SHOM, used under CC BY-SA 4.0 license

Based on the documentation of the SHOM², Fig. 4.3 shows a map of the repartition of the different types of sediments introduced there. A more complete chart with a finer classification of the types of sediments can be found in the appendix, on Fig. A.1.

Based on this classification, we can see that most of the ocean floor of the studied domain is composed of sand. Even though siltic soil is listed, it is only scarcely present.

²Service hydrographique et océanographique de la Marine, <https://www.shom.fr/fr>

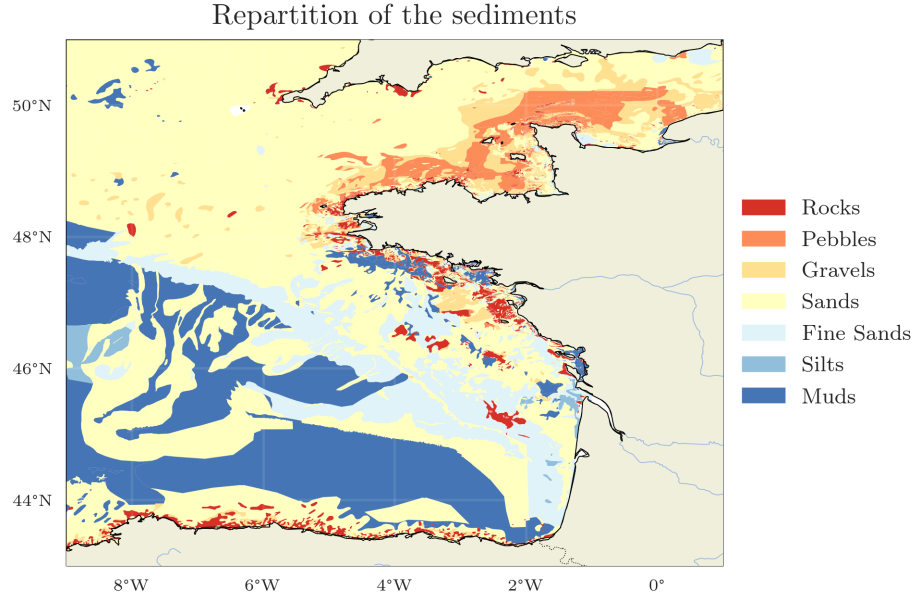


Figure 4.3 – Repartition of the sediments on the ocean floor. Data source: SHOM, used under [CC BY-SA 4.0](#) license

The figure also shows that the largest sediments are rocks but are mostly located in the Bay of Biscay, near the boundary of the continental shelf. Pebbles however are mostly located in the shallow region in the English Channel, thus it may be expected that controlling the roughness in the regions listed as pebbles will affect significantly the water circulation, and thus the sea surface height. Incidentally, we can notice the inverse correlation between the size of the sediments, and the depth at which they are found.

According to [Table 4.1](#), the rugosity z_b spans multiple orders of magnitude, so we are then going to define the control variable θ as

$$\theta = \log z_b \in \Theta = [\theta_{\min}, \theta_{\max}]^p \quad (4.4)$$

where $\theta_{\min} = \log(10^{-5}) \approx -11.5$, and $\theta_{\max} = \log(5 \cdot 10^{-2}) \approx -3$. The dimension of Θ is noted p , and will be specified later depending on the chosen segmentation.

4.2.3.b Tidal modelling and uncertainties

The ocean, especially near the English Channel is driven by tidal forces that produce currents at the surface. As a periodic signal, the tidal forcing is usually analysed harmonically, in order to separate its influence by frequency. In CROCO, this forcing can come from the TPXO model of tides ([Egbert and Erofeeva, 2002](#)), and in our configuration, we use the 5 primary harmonic constituents as described [Table 4.2](#).

By perturbing some properties of those tide components, we can artificially introduce some error in the numerical model, *i.e.* a parametric misspecification that we will

Darwin Symbol	Period (h)	Species
M_2	12.4206	Principal Lunar Semidiurnal
S_2	12	Principal Solar Semidiurnal
N_2	12.65834751	Larger Lunar Elliptic Semidiurnal
K_2	11.96723606	Lunisolar Semidiurnal
K_1	23.93447213	Lunar Diurnal

Table 4.2 – Harmonic constituents used in the configuration

define as the environmental variable. In this work, this takes the form of a small multiplicative error on the amplitude of the different components of the tide. Let $u = (u_1, \dots, u_5) \in \mathbb{U} = [0, 1]^5$ be the environmental variable, that will be considered as a random variable later, and let A_k be the amplitude of the k th component of the tide. The perturbed amplitude $\tilde{A}_k$ is defined as

$$\tilde{A}_k(u_k) = A_k(1 + 0.01(2u_k - 1)) \quad (4.5)$$

for $1 \leq k \leq 5$. Based on this definition, the perturbed amplitude varies from $\tilde{A}_k(0) = 0.99A_k$ to $\tilde{A}_k(1) = 1.01A_k$, for every $1 \leq k \leq 5$. We define also u^{truth} as the vector whose all its components are set to 0.5, so when no amplitude is perturbed: $\tilde{A}_k(u_k^{\text{truth}}) = A_k$ for $1 \leq k \leq 5$.

In the next section, we are going to estimate the parameter θ by minimising the objective function using a gradient-descent algorithm.

4.3 Deterministic calibration of the bottom friction

The calibration of the bottom friction will first be studied without external uncertainties, which corresponds to an unperturbed tidal forcing. In consequence, the environmental variable is set to u^{truth} .

In ocean modelling, the free-surface height ζ is often used as an observable quantity, because it can be measured using satellites or tide gauges near the coasts. We consider then $\zeta \in \mathbb{R}^{N_{\text{obs}}}$ as the output of the numerical model.

Following the notations introduced in [Chapter 1](#) for the model, let us define $(\mathcal{M}(\cdot, u^{\text{truth}}), \Theta)$ as the numerical model to calibrate. The forward operator $\mathcal{M}(\cdot, u^{\text{truth}})$ is defined by

$$\begin{aligned} \mathcal{M} : \Theta &\longrightarrow \mathbb{R}^{N_{\text{obs}}} \\ \theta &\longmapsto \mathcal{M}(\theta, u^{\text{truth}}) = (\zeta_{t,i}(\theta, u^{\text{truth}}))_{\substack{1 \leq i \leq N_{\text{mesh}} \\ 1 \leq t \leq N_{\text{time}}}} \end{aligned} \quad (4.6)$$

where $\zeta_{t,i}(\theta, u)$ is the free-surface height of the ocean at the mesh point i , and at the time-step t , obtained using the model and the bottom friction associated with θ and the environmental variable u , and $N_{\text{mesh}} \cdot N_{\text{time}} = N_{\text{obs}}$. In this configuration, $N_{\text{time}} = 49$, and corresponds to the number of records saved, while $N_{\text{mesh}} = 15\,684$ is the number of cells of the computational grid not located on land. The time step of the simulation in itself is 10 s, and its total duration is 24 h, meaning that the water height ζ is saved every 30 min.

4.3.1 Twin experiment setup

Recalling the definition of a twin experiment setup, the observation y is generated using the numerical model, and a predefined truth value θ^{truth} . This means that the “physical model” is defined using the forward operator $\mathcal{M}$, based on the forward numerical model, evaluated with the environmental parameter u^{truth} .

$$\begin{aligned}\mathcal{M} : \Theta &\longrightarrow \mathbb{R}^{N_{\text{obs}}} \\ \theta &\longmapsto \mathcal{M}(\theta) = \mathcal{M}(\theta, u^{\text{truth}})\end{aligned}\tag{4.7}$$

Based on the forward operator $\mathcal{M}$, we can generate the observations $y \in \mathbb{Y} = \mathbb{R}^{N_{\text{obs}}}$ using the truth value $\theta^{\text{truth}} = \log z_b^{\text{truth}}$, as defined [Table 4.1](#):

$$y = \mathcal{M}(\theta^{\text{truth}}) = \mathcal{M}(\theta^{\text{truth}}, u^{\text{truth}})\tag{4.8}$$

In the following, if the u argument is omitted, it means that the model, or subsequent functions are evaluated with u^{truth} .

4.3.2 Cost function definition

Once the observation $y \in \mathbb{R}^{N_{\text{obs}}}$ has been generated, we can define the objective function J :

$$J(\theta) = \sum_{t=1}^{N_{\text{time}}} \sum_{i=1}^{N_{\text{mesh}}} \left(\zeta_{t,i}(\theta, u^{\text{truth}}) - y_{t,i} \right)^2\tag{4.9}$$

$$= \|\mathcal{M}(\theta, u^{\text{truth}}) - y\|_2^2\tag{4.10}$$

Equivalently, as mentioned in [Chapter 1](#), by assuming that the distribution of the (random) observation vector is known and $Y \mid \theta \sim \mathcal{N}(\mathcal{M}(\theta), I)$ (with I being the identity matrix of dimension p), J is proportional to the negative log-likelihood of the data.

4.3.3 Gradient-descent optimisation

The optimisation is carried using M1QN3, a version of a gradient-descent procedure, as described in [Gilbert and Lemaréchal \(1989\)](#). We can first look to control z_b at every cell of the mesh: $\theta = (\theta_1, \dots, \theta_p)$ where $\theta_i = \log z_b^i$ and $p = 15\,684$.

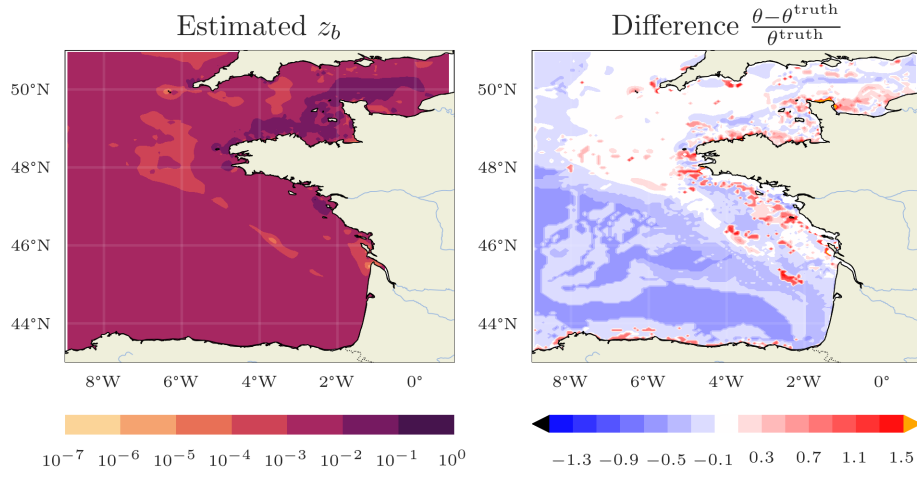
Due to the large number of points whose friction can be controlled, a finite-difference method to get the gradient is unfeasible. Instead, Tapenade ([Hascoet and Pascual, 2013](#)), an Automatic Differentiation tool has been used in order to get the gradient (with respect to θ) of the cost function J using the adjoint method, as described [Section 1.4](#). The starting point of the optimisation is $\theta = \log 5 \times 10^{-3}$, and the procedure is stopped after 400 iterations. The estimated control parameter (*i.e.* the minimiser found) is shown [Fig. 4.4a](#). The evolution of the cost function and the squared norm of the gradient during the optimisation procedure is shown [Fig. 4.4b](#).

By comparing the result of the optimisation [Fig. 4.4a](#) with the sediment chart [Fig. 4.3](#) and the bathymetry map [Fig. 4.1](#), we can have a first overview on which regions of the

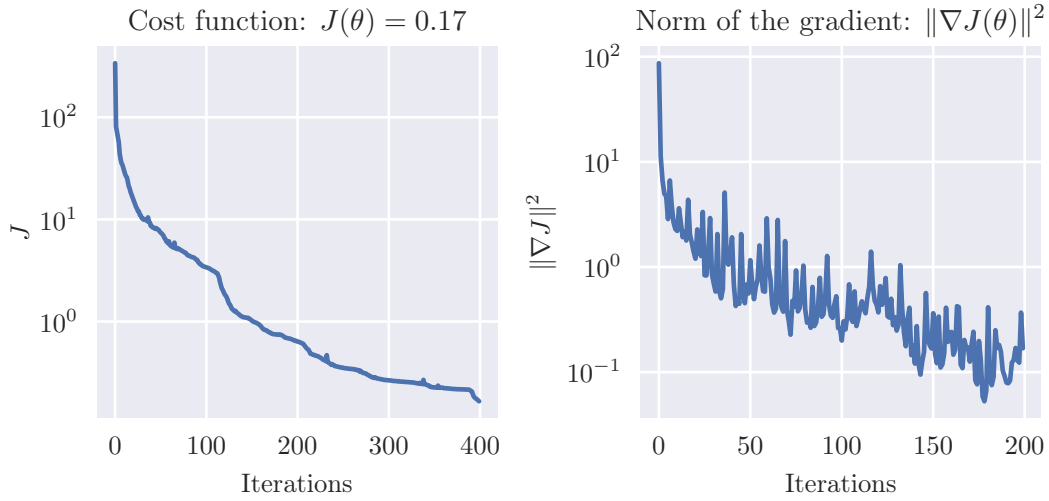
domain are properly estimated (*i.e.* where the estimation is close to the truth value). On a first look, we can see that the abyssal plain (the deep region off the Bay of Biscay) remains mostly unaffected by the optimisation, while the continental shelf, except for some parts of the English Channel, is well retrieved.

On Fig. 4.5 is shown the value of the optimised rugosity z_b , depending on the type of sediment associated. Each point on this figure corresponds then to a mesh point. We can see that indeed, as Fig. 4.4a shows, points of the mesh corresponding to sand, *i.e.* most of the continental shelf, tends to get closer to the truth value z_b^{truth} . For silts and muds however, the procedure did not change significantly their roughness, and thus stays close to the initial value.

This can be probably explained by the great depth at which those sediments lay, and thus it mitigates their influence on the drag coefficient per Eq. (4.3). In the English Channel, the size of the pebbles is quite well retrieved albeit a bit underestimated, but the points mapped to gravel do seem to compensate: on the northern part of the channel the size of the gravel is overestimated, while it is underestimated on the southern part. Finally the rocks appear to be hard to estimate properly, as only about 3% of the domain is listed as rocks, and their corresponding z_b is quite large in contrast to the other sediments.



(a) Optimisation of z_b on the whole space using gradient obtained via adjoint method, after 400 iterations.



(b) Evolution of the cost function and the squared norm of the gradient

Figure 4.4 – Calibration of the bottom friction using gradient-descent with well-specified environmental variables

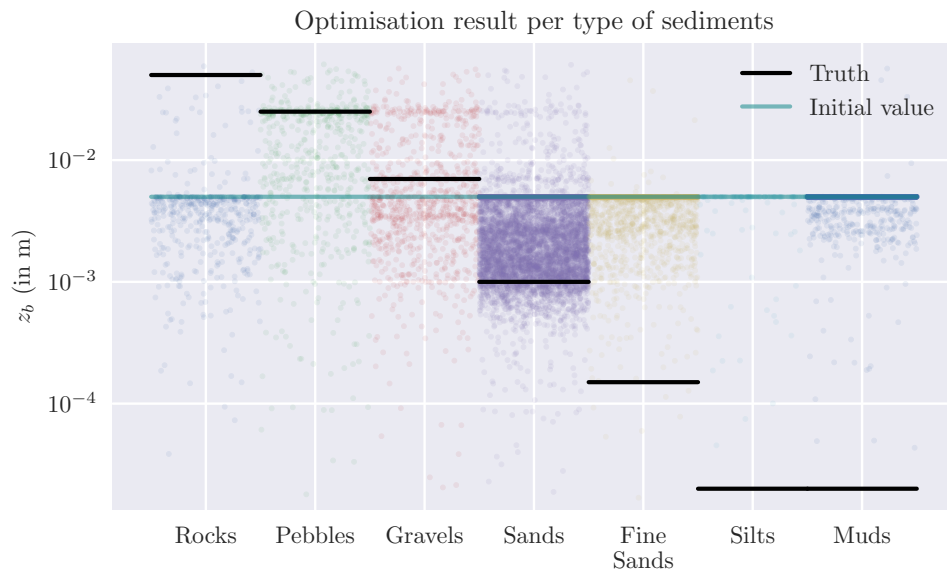


Figure 4.5 – Results of the optimisation procedure, depending on the type of sediments. The starting value in the optimisation procedure for z_b is 5×10^{-3} m

The optimisation procedure has been done also in the misspecified case, *i.e.* $u^b \neq u^{\text{truth}} = (0.5, 0.5)$ where the first component refers to the error on M_2 , and the second on S_2 (the others are supposed to be well-specified).

Figure 4.6 shows the resulting optimisation for $u^b = (0, 0)$ and for $u^b = (1, 1)$.

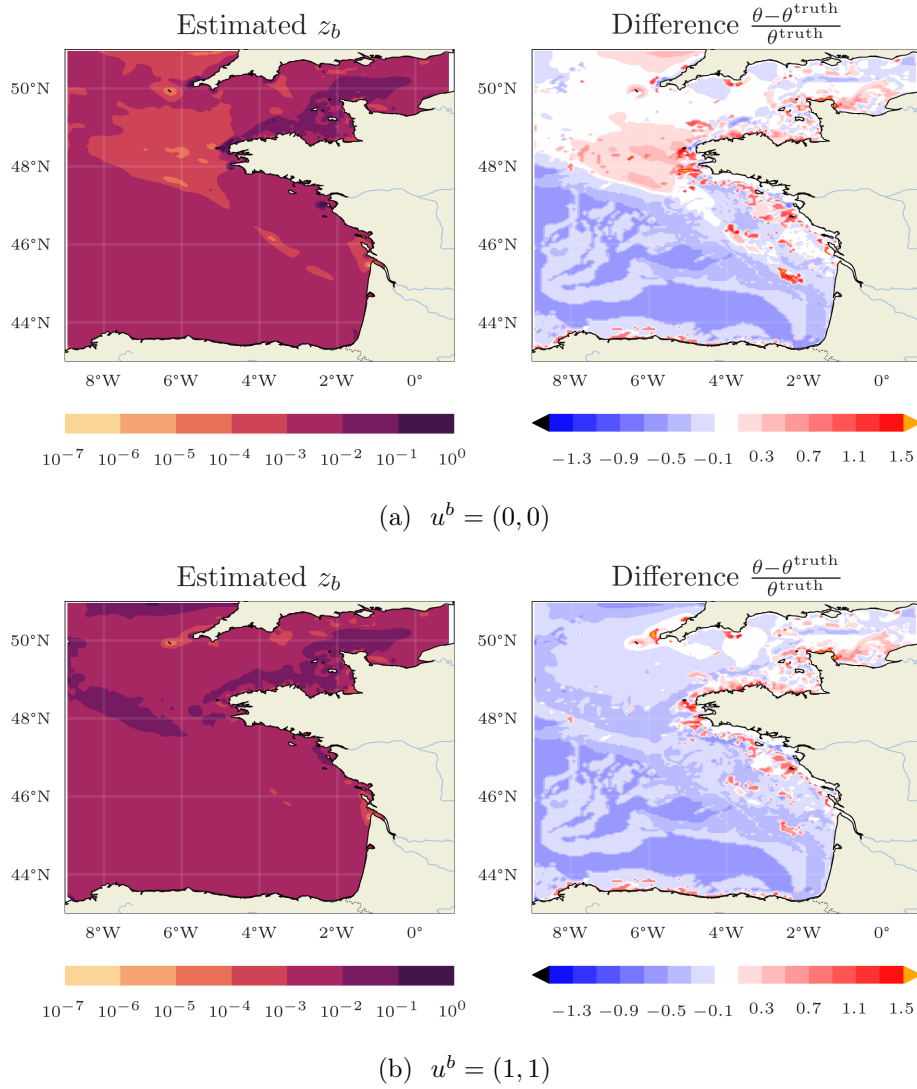


Figure 4.6 – Optimisation of z_b on the whole space, with a misspecification of the environmental variable

We can see that while the Bay of Biscay is left unaffected by the optimisation procedure, the values in the English Channel seem to be well estimated, no matter the misspecification. However, the bottom friction near the Celtic Sea seems to compensate the error due to the misspecification. In the appendix [Section A.3](#) is presented the result of the optimisation over the whole domain, for other choices of the environmental parameter.

This optimisation showed that all the regions are not crucial for a significant reduction of the objective. So, in order to clarify this, we are now going to study the influence of the different inputs of the objective function. More specifically, we are going to quantify the influence of each region defined by its sediment type, and quantify the influence of the uncertainties defined in [Section 4.2.3.b](#) on the output of the objective function, by means of sensitivity analysis.

4.4 Sensitivity analysis of the objective function

Sensitivity analysis (often abbreviated as *SA*), aims at quantifying the effect of the variation of some input variable to the output of the model ([Iooss, 2011](#); [Janon, 2012](#)). Intuitively, SA aims at understanding how much the variations of each input or combination of inputs explain the variations of the output.

It can then be approached at two different scales: around a nominal value, using the gradient, and at a global scale, by considering the inputs as random variables, and by measuring the variance of the output. In this work, we are going to focus exclusively on global sensitivity analysis.

Here, sensitivity analysis is performed as a dimension-reduction method, because it will be used to reduce the input space, based on the prior assumption that the rugosity z_b is considered constant for each regions. Indeed, we will make the link between *sensitivity* and *identifiability* ([Dobre et al., 2010](#)). If a parameter shows a very small influence on the output of the objective function, any choice of its value will yield sensibly the same output of the objective, provided that the other input parameters are the same, justifying their overlook.

4.4.1 Global Sensitivity Analysis: Sobol' indices

As global SA calls for a probabilistic framework, we are going to consider a real-valued random vector $X = (X_1, \dots, X_p)$, whose components are assumed independent, to represent the inputs (*i.e.* θ and u) of a real function $f : \mathbb{R}^p \rightarrow \mathbb{R}$ (the objective function in this thesis). As X is a random vector, we can introduce Y , the real-valued random variable defined as $Y = f(X)$.

The i -th Sobol' indice of order 1 is defined as ([Sobol, 1993, 2001](#))

$$S_i = \frac{\text{Var}_{X_i} [\mathbb{E}_Y [Y \mid X_i]]}{\text{Var}_Y [Y]} \quad (4.11)$$

and can be interpreted as the fraction of variance of the output Y explained by the variation of X_i *alone*. Indices of order 2 are defined as

$$S_{i \times j} = \frac{\text{Var}_{X_i, X_j} [\mathbb{E}_Y [Y \mid X_i, X_j]]}{\text{Var}_Y [Y]} - S_i - S_j \quad (4.12)$$

and account for the interactions of the inputs labelled i and j . Higher-order Sobol' indices can then be defined sequentially. Total-effect indices are also central in global sensitivity analysis: those indices measure the contributions of a single input X_i through all its

possible interactions, *i.e.* by considering the effect of its own variability (as the order 1) and the effect of all its interactions (order 2 and above). Those total-effect indices can be expressed as

$$S_{T_i} = 1 - \frac{\text{Var}_{X_{-i}} [\mathbb{E}_Y [Y \mid X_{-i}]]}{\text{Var}_Y [Y]} \quad (4.13)$$

where $X_{-i} = (X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_p)$ is a random vector of $p - 1$ components.

In our study, the Sobol' indices of order 1, 2 and total ones are computed using a replicated method (Gilquin et al., 2019; Gilquin, 2016), allowing for bootstrap confidence intervals for the first and second order effects.

4.4.2 SA of the objective function for the calibration of CROCO

4.4.2.a SA on the regions defined by the sediments

The drag coefficient, which affects the ocean circulation, is the result of two factors as shown in Eq. (4.3): the bottom roughness z_b and the ocean depth H . So depending on the depth, the influence of the rugosity at the bottom on the objective function may change.

We are first going to perform a sensitivity analysis in order to quantify the role of each sediment-based *region*, without incorporating the knowledge on the typical size of the sediment there. Considering the similar expected size of silts (Si) and muds particles (V) in Table 4.1, and the limited amount of silts, we will merge those regions, and label the result with the code Si, V. Finally, the function which will be analyzed is

$$\theta \mapsto \|\mathcal{M}(\theta, u^{\text{truth}}) - y\|_2^2 \quad (4.14)$$

with

$$\theta = (\theta_R, \theta_C, \theta_G, \theta_S, \theta_{SF}, \theta_{Si,V}) \in \Theta, \text{ with } \Theta = [\theta_{\min}, \theta_{\max}]^6 \quad (4.15)$$

In the context of the sensitivity analysis, all the inputs are assumed to be independent, and to be uniformly distributed on their support $[\theta_{\min}, \theta_{\max}] \approx [-11.5, -3]$. The first, second and total-order effects are displayed Fig. 4.7, where the regions defined by each sediment is Fig. 4.3. The experimental design used here comprises 7888 points.

We can see that the most influential region is the one defined by the pebbles (with code C). More generally, except for the rocks, we can see that the sediments that lies in shallower regions have a larger impact on the objective function. Based on geographic considerations and the result of this sensitivity analysis, we can adopt a new segmentation: the regions of Pebbles (C), and Gravel (G) will be kept intact, but the remaining regions (R, S, SF, Si, V) will be bundled together.

4.4.2.b SA on the tide components

As introduced Section 4.2.3.b, CROCO incorporates different tide constituents that we perturbate through the uncertain variable $u \in \mathbb{U}$. In order to quantify the influence of each component of the environmental parameter, we performed also a sensitivity analysis. As before, the SA is performed on the sum of squares of the difference between the

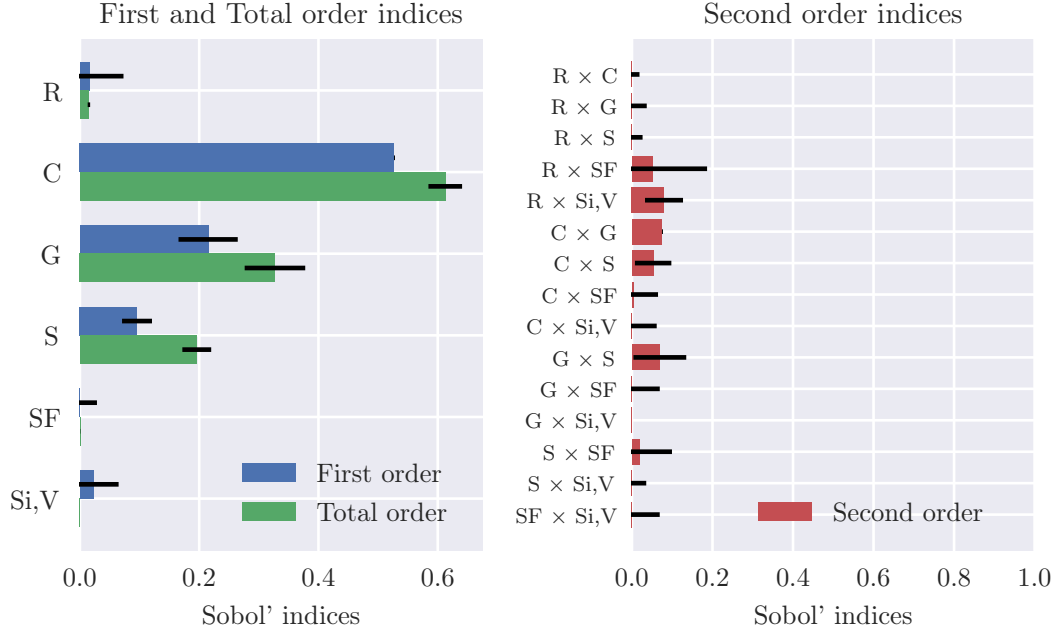


Figure 4.7 – Global SA on the regions defined by the sediment type, and bootstrap confidence intervals

forward operator and the observations. This time however, the control parameter θ is set to its truth value, and the sum of squares depends only on u :

$$u \mapsto \|\mathcal{M}(\theta^{\text{truth}}, u) - y\|_2^2 \quad (4.16)$$

for $u = (u_1, \dots, u_5)$, and $\mathbb{U} = [0, 1]^5$. Once again, in the SA context, every component is assumed to be uniformly distributed on $[0, 1]$ and independent. Figure 4.8 shows the Sobol' indices of first order (left), second (right), and the total effect indices, along with bootstrap confidence intervals, obtained with a design of experiments of 2888 points.

We can see that the component of the vector affecting the amplitude of the M_2 component of the tide has the most impact on the cost function, and the S_2 component seems to have a non negligible effect as well. For the other tide constituents, the SA reveals that perturbing their amplitude has little to no-effect in this configuration, and thus those variables will be discarded in further analysis. The uncertain variable can then be redefined in what follows as

$$U = (U_1, U_2) \quad (4.17)$$

where U_1 is the error on the M_2 amplitude, and U_2 is the error on the S_2 amplitude, U_1 and U_2 are independent, and $U \sim \mathcal{U}(\mathbb{U})$ with $\mathbb{U} = [0, 1]^2$.

4.5 Robust Calibration of the bottom friction

In this section, we are going to study the robust calibration of the numerical model, on the space defined according to the sensitivity analysis performed before. On this reduced

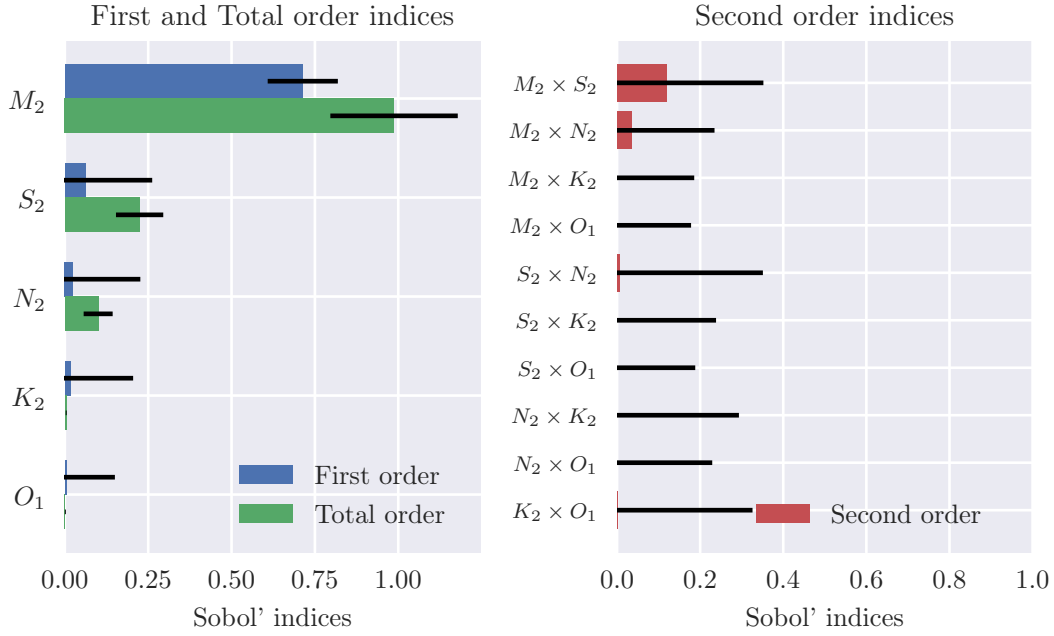


Figure 4.8 – Global SA on the different components of the tide, and bootstrap confidence intervals

space, the objective function J will first be optimised globally on $\Theta \times \mathbb{U}$ in [Section 4.5.1](#), and its results will be compared with the true value (used to generate the observations).

Afterwards, we are going to estimate relative-regret based estimates of the bottom friction as described in the previous chapters. The general approach is summarised [Fig. 4.9](#). We are first going to define and fit a Gaussian Process Z with respect to an initial design evaluated by the objective function J , as introduced in [Chapter 3](#). As J is supposed to be positive, we are first going to ensure that the surrogate constructed using J , namely m_Z is positive as well, by enriching the design using the PEI criterion in [Section 4.5.2](#). This will allow us to estimate the conditional minimums and conditional minimums. This initialisation step is represented as the top row of [Fig. 4.9](#).

The second row of the figure represents the enrichment step, which consists in adding points to the design according to some adaptive strategies. These strategies are implemented in order to improve the estimation of some quantities of interest linked to the relative-regret estimates: estimation of Γ_α using the augmented IMSE in [Section 4.5.3.a](#), and sampling in the margins of uncertainty for q_p in [Section 4.5.3.b](#).

Finally, using the Gaussian Process Z fitted based on this final enriched design, we then construct the associated surrogate m_Z to emulate J and to compute the estimations of Γ_α or q_p , and optimise them to get the members of the relative-regret family of estimators (as summarised in the last row of [Fig. 4.9](#)).

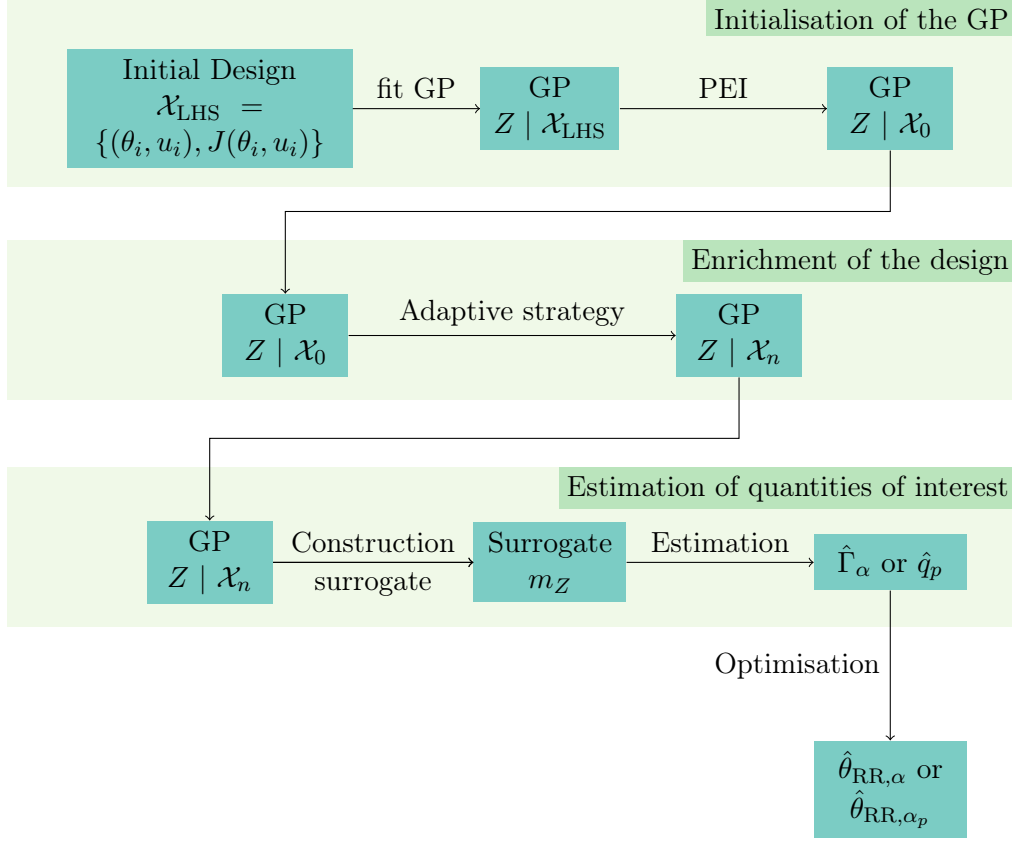


Figure 4.9 – Representation of the different steps for the robust calibration of the numerical model using relative-regret estimates

4.5.1 Objective function and global minimum

Based on the sensitivity analysis carried in the previous section, we are going to consider the following setting for robust calibration. The control variable is defined as

$$\theta = \left(\theta^{(1)}, \theta^{(2)}, \theta^{(3)} \right) \in \Theta \quad \Theta = [\theta_{\min}, \theta_{\max}]^3 \quad (4.18)$$

where the superscript (1) corresponds to the region defined as Pebbles, (2) is for the regions defined as Gravel, and (3) corresponds to the merged regions of Rocks, Sand, Silts, Mud, and Fine Sands.

Following the SA on the tide constituents, the uncertain variable is a random vector with two components, representing the error on amplitude of the M_2 and the S_2 tide constituents:

$$U = (U_1, U_2), \quad U_i \sim \mathcal{U}([0, 1]) \text{ for } i = 1, 2 \quad (4.19)$$

As previously, the forward operator of the numerical model yields the sea-surface height ζ :

$$\begin{aligned} \mathcal{M} : \Theta \times \mathbb{U} &\longrightarrow \mathbb{R}^{N_{\text{obs}}} \\ (\theta, u) &\longmapsto \mathcal{M}(\theta, u) = (\zeta_{t,i}(\theta, u))_{\substack{1 \leq i \leq N_{\text{mesh}} \\ 1 \leq t \leq N_{\text{time}}}} \end{aligned} \quad (4.20)$$

and the objective function J is the squared difference between the observations y and the forward operator.

$$\begin{aligned} J : \Theta \times \mathbb{U} &\longrightarrow \mathbb{R} \\ (\theta, u) &\longmapsto \|\mathcal{M}(\theta, u) - y\|_2^2 \end{aligned} \quad (4.21)$$

For the sake of the study, we can optimise the original function J over $\Theta \times \mathbb{U} = [\theta_{\min}, \theta_{\max}]^3 \times [0, 1]^2$. Because we are controlling the calibration parameter on a space of dimension 3 while the observation y has been generated without the dimension reduction, the result of the optimisation can be different from the truth values. We obtain:

$$\min_{(\theta, u) \in \Theta \times \mathbb{U}} J(\theta, u) = J(\hat{\theta}_{\text{global}}, \hat{u}_{\text{global}}) = 29.749 \quad (4.22)$$

for

$$\hat{\theta}_{\text{global}} = (-3.5162, -5.0776, -6.3459) \quad (4.23)$$

$$\hat{u}_{\text{global}} = (0.6348, 0.2989) \quad (4.24)$$

We can then indeed notice that $u^{\text{truth}} = (0.5, 0.5) \neq \hat{u}_{\text{global}}$. The difference in the environmental variables $\hat{u}_{\text{global}}$ and u^{truth} compensates for the difference of dimensionality between $\theta \in \Theta$, the control variable, and θ^{truth} .

In order to compare the result of this optimisation with the truth value, [Table 4.3](#) shows the different values of the different components of θ^{truth} and $\hat{\theta}_{\text{global}}$.

Component	θ^{truth}	Component	$\hat{\theta}_{\text{global}}$
θ_{C}	-3.6889	$\theta^{(1)}$	-3.5162
θ_{G}	-4.9618	$\theta^{(2)}$	-5.0776
θ_{R}	-2.9957	} $\theta^{(3)}$	-6.3459
θ_{S}	-6.9078		
θ_{SF}	-8.8049		
$\theta_{\text{Si,V}}$	-10.8198		

Table 4.3 – Values of the θ component of the global optimiser, and truth value. The region associated with $\theta^{(3)}$ is the union of the regions defined with code R, S, SF, Si and V.

As mentioned in the previous chapter, J can be expensive to evaluate computational-wise, so we propose to use GP in order to model it, and to enrich its design for the computation of members of the relative-regret family of estimators.

We will first define the initial design, which will be evaluated by J . A Latin Hyper-square of 100 points on $\Theta \times \mathbb{U}$ is first sampled in order to construct a GP that can be used as a surrogate. In this work, the GP will be constructed using the Python module Scikit-learn ([Pedregosa et al., 2011](#)).

We denote $\mathcal{X}_{\text{LHS}}$ the initial LHS, and Z the Gaussian Process constructed and fitted using $\mathcal{X}_{\text{LHS}}$. We will write

$$Z \sim \text{GP}(m_Z, C_Z) \quad \text{and} \quad C_Z((\theta, u), (\theta, u)) = \sigma_Z^2(\theta, u) \quad (4.25)$$

We then have, for any $(\theta, u) \in \Theta \times \mathbb{U}$,

$$Z(\theta, u) \sim \mathcal{N}(m_Z(\theta, u), \sigma_Z^2(\theta, u)) \quad (4.26)$$

By definition, for any $\theta \in \Theta$ and any $u \in \mathbb{U}$, we have $J(\theta, u) \geq 0$. However, the positivity of the objective function is not necessarily verified by the surrogate m_Z , and thus the notion of relative-regret is not defined. We have to first ensure that the GP Z (and consequently Z^*) is positive with large enough probability. To do so, we are first going to enrich the design near the conditional minimisers, using the PEI criterion. Alternatively, one could have added points to the design according to the reliability index, defined Eq. (3.40), and thus look and evaluate points having the largest probability of being negative.

4.5.2 Conditional minimums and conditional minimisers

In the previous chapter, we defined the conditional minimum as the minimum of the objective function at a given $u \in \mathbb{U}$:

$$J^* : u \mapsto J^*(u) = \min_{\theta \in \Theta} J(\theta, u) \quad (4.27)$$

The conditional minimisers function is defined as

$$\theta^* : u \mapsto \theta^*(u) = \arg \min_{\theta \in \Theta} J(\theta, u) \quad (4.28)$$

Since both of these functions require an optimisation of the objective function, they are quite expensive to compute, so, as done before, we use the GP prediction of Z in order to approximate them:

$$m_{Z^*} : u \mapsto \min_{\theta \in \Theta} m_Z(\theta, u) \quad (4.29)$$

$$\theta_Z^* : u \mapsto \arg \min_{\theta \in \Theta} m_{Z^*}(u) \quad (4.30)$$

As mentioned before, in order to improve the accuracy of those two functions and to ensure the positivity of m_Z , we are first going to enrich the design using the PEI criterion (Ginsbourger et al., 2014), as introduced Section 3.5.1. Choosing 200 additional points this way, the new obtained design is denoted $\mathcal{X}_0$, that will serve as the “initial” design for the enrichment procedures described later.

For the sake of this work, we are going to estimate the conditional minimisers as accurately as possible, using in total 750 evaluations of the objective function J . By sampling u_i from U for $1 \leq i \leq n_u = 2000$, we can first estimate $m_{Z^*}(u)$ using this set of samples, as shown in Fig. 4.10. As expected from the result of the optimisation carried in Section 4.5.1, we can see that the minimum of $m_{Z^*}(u)$ (i.e. the global minimum) is not attained at $(0.5, 0.5) = u^{\text{truth}}$ but rather at approximatively $(0.6, 0.3)$.

Based on these samples, we can also approximate the distribution of the random variable $\theta^*(U)$ by $\theta_Z^*(U) = (\theta_Z^{(1)*}(U), \theta_Z^{(2)*}(U), \theta_Z^{(3)*}(U))$. Figure 4.11 shows the pairwise relations between the components of the random variable $\theta_Z^*(U)$, based on the samples

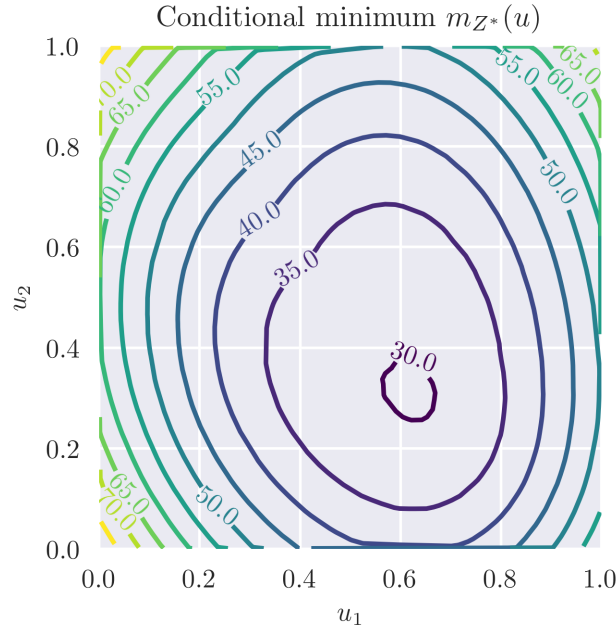


Figure 4.10 – Conditional minimum $m_{Z^*}(u)$ estimated using the GP Z

$\theta_Z^*(u_i)$ for $1 \leq i \leq n_u$. On the diagonal plots are shown the Kernel Density Estimation of the marginal distributions of $\theta^*(U)$, and the corresponding samples from $m_{Z^*}(U)$. The off-diagonal plots show the relations between the components. We can then observe a negative correlation between $\theta_Z^{(1)*}(U)$ and $\theta_Z^{(2)*}(U)$ for instance, indicating that an increase in one of those value of the friction is compensated by the decrease in the other component. We can also notice that the marginal distribution of the first component $\theta_Z^{(1)*}(U)$ seems to exhibit two modes: one at -3.6 , and one at around -3.45 , and both are quite close to the value found when optimising globally: $\hat{\theta}_{\text{global}}^{(1)} \approx -3.5$, which makes it a clear candidate for robust optimisation. For the two other components, we can observe that the range of values taken by the samples $\theta_Z^*(u_i)$ for $1 \leq i \leq n_u$ is larger, indicating that for these components, a robust candidate may be less *identifiable*, as discussed [Section 2.4.1, Page 49](#), but then also less influential.

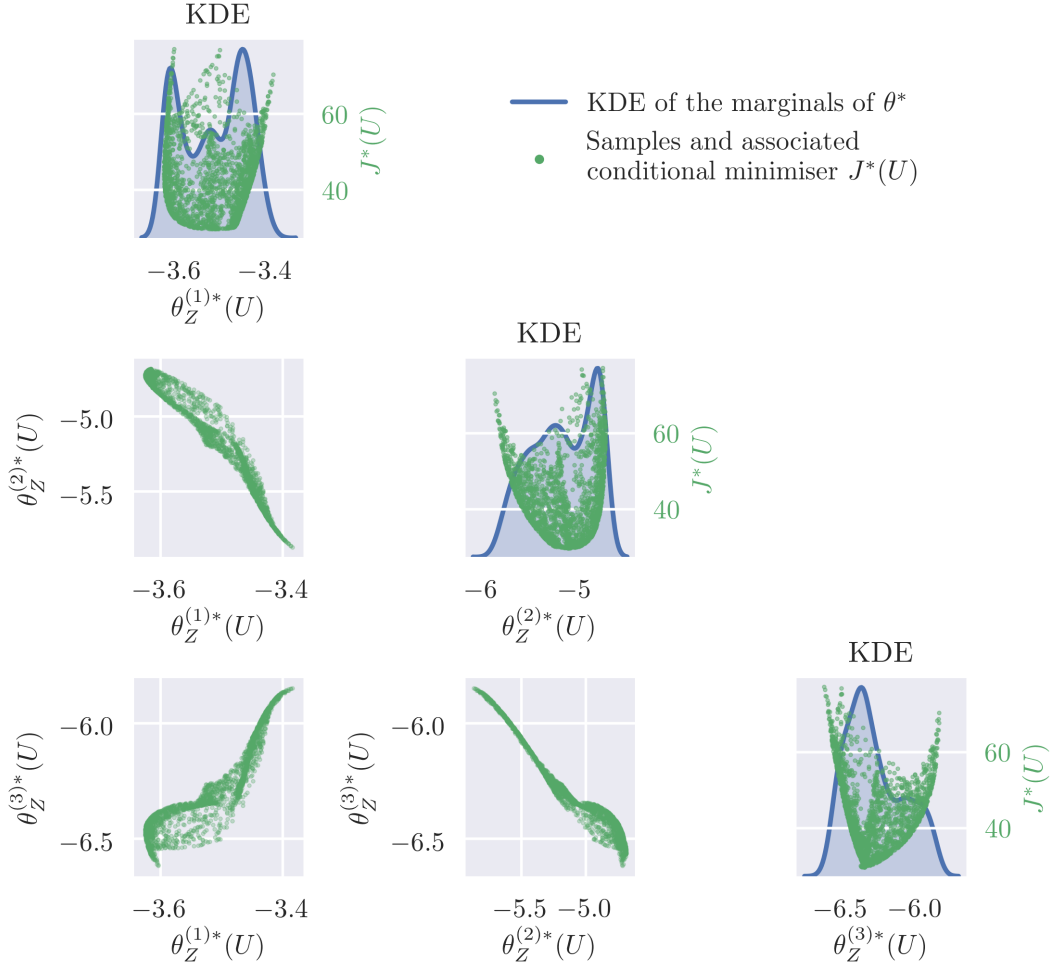


Figure 4.11 – Estimated distribution of the minimisers $\theta_Z^*(U)$. The diagonal plots show the KDE of the marginal distributions of the minimisers. Each point represents a sample, and its vertical component is the conditional minimum associated. The non-diagonal plots show the pairwise relation between the different components of the minimisers

4.5.3 Relative-regret based estimators

We are now going to estimate members of the relative-regret family of estimators as introduced [Definition 2.4.6](#), on [Page 55](#):

$$\left\{ \theta_{\text{RR},\alpha} = \arg \max_{\theta \in \Theta} \Gamma_{\alpha}(\theta) \mid \alpha \geq 1 \right\} \quad (4.31)$$

where the function to be maximised,

$$\Gamma_{\alpha}(\theta) = \mathbb{P}_U [J(\theta, U) \leq \alpha J^*(U)] \quad (4.32)$$

is the probability that θ is α -acceptable. In the following, we will assume also that we have at disposal a budget of 500 runs of the forward model, so after the initial design and the initial iterations of the PEI criterion, we assume available 200 additional runs.

4.5.3.a Optimisation of the probability of exceeding a threshold

Plug-in approximation of the probability of acceptability

Let $\alpha > 1$. The Γ_{α} function is quite expensive to compute, so we are going to use a plug-in approach (see [Definition 3.1.2](#)) in order to avoid exhaustive computations of the objective function J . Instead, we will use m_Z and m_{Z^*} as surrogates of J and J^* . For the approximation of the probability, we are going to use a Sample Average Approximation (SAA) approach, by using a set of n_u i.i.d. samples $\{u_i\}_{1 \leq i \leq n_u}$. All in all, both approximations leads to

$$\hat{\Gamma}_{\alpha}^{\text{PI}}(\theta) = \frac{1}{n_u} \sum_{i=1}^{n_u} \mathbb{1}_{\{m_Z(\theta, u_i) \leq \alpha m_{Z^*}(u_i)\}} \quad (4.33)$$

Maximising this expression yields

$$\hat{\theta}_{\text{RR},\alpha} = \arg \max_{\theta \in \Theta} \hat{\Gamma}_{\alpha}^{\text{PI}}(\theta) = (\hat{\theta}_{\text{RR},\alpha}^{(1)}, \hat{\theta}_{\text{RR},\alpha}^{(2)}, \hat{\theta}_{\text{RR},\alpha}^{(3)}) \quad (4.34)$$

In order to get a relevant value of $\hat{\theta}_{\text{RR},\alpha}$, we are going to reduce sequentially the augmented IMSE, as introduced in [Section 3.5.3](#).

Stepwise reduction of the augmented IMSE

We will first look to improve the estimation of Γ_{α} , by improving the plug-in approximation. We define $\Delta_{\alpha} = Z - \alpha Z^* \sim \text{GP}(m_{\Delta_{\alpha}}, C_{\Delta_{\alpha}})$, which is a GP as defined in the previous chapter. The prediction variance, or mean square error, is $\sigma_{\Delta_{\alpha}}^2 : (\theta, u) \mapsto C_{\Delta_{\alpha}}((\theta, u), (\theta, u))$. In the following, we are going to estimate the relative-regret, associated with the value $\alpha = 1.3$, meaning that we are looking to maximise the probability of being within 30% of the optimal value. Such a value has been chosen based on a preliminary analysis, using the GP constructed on the initial design.

We are then going to reduce the augmented IMSE of the random process Δ_{α} . Recalling that the IMSE associated with the GP $\Delta_{\alpha} = Z - \alpha Z^*$ constructed using the design $\mathcal{X}_n$ is defined as

$$\text{IMSE}(\mathcal{X}_n) = \int_{\Theta \times \mathbb{U}} \sigma_{\Delta_{\alpha}}^2(x) \, dx \quad (4.35)$$

we select the next point to evaluate as:

$$(\theta_{n+1}, u_{n+1}) = \arg \min_{(\theta, u) \in \Theta \times \mathbb{U}} \mathbb{E}_{Z(\theta, u)} [\text{IMSE}(\mathcal{X}_n \cup ((\theta, u), Z(\theta, u)))] \quad (4.36)$$

In practice, we will choose the new point (θ_{n+1}, u_{n+1}) in a stochastic manner, as the augmented IMSE will be estimated using a Monte-Carlo method. Let us consider the design augmented with the point $((\theta, u), z(\theta, u))$, which is equivalent to making the assumption that the function takes the value $z(\theta, u) > 0$ at the point (θ, u) , we note $\sigma_{\Delta_\alpha|z}^2$ the prediction variance associated with this design. The estimation of the augmented IMSE is then

$$\text{IMSE}(\mathcal{X}_n \cup ((\theta, u), z(\theta, u))) \approx \frac{1}{n_{\text{IMSE}}} \sum_{i=1}^{n_{\text{IMSE}}} \sigma_{\Delta_\alpha|z}^2(\theta_i, u_i) \quad (4.37)$$

where the points (θ_i, u_i) for $1 \leq i \leq n_{\text{IMSE}}$ are resampled each iteration using a LHS on $\Theta \times \mathbb{U}$ for $n_{\text{IMSE}} = 150$.

The expectation operator of Eq. (4.36) is also approximated by choosing $z_j(\theta, u)$ for $1 \leq j \leq n_Z$ as different quantiles of $Z(\theta, u)$ (which is normally distributed). Finally, we have

$$\mathbb{E}_{Z(\theta, u)} [\text{IMSE}(\mathcal{X}_n \cup ((\theta, u), Z(\theta, u)))] \approx \frac{1}{n_Z} \sum_{j=1}^{n_Z} \frac{1}{n_{\text{IMSE}}} \sum_{i=1}^{n_{\text{IMSE}}} \sigma_{\Delta_\alpha|z_j}^2(\theta_i, u_i) \quad (4.38)$$

$$\propto \sum_{j=1}^{n_Z} \sum_{i=1}^{n_{\text{IMSE}}} \sigma_{\Delta_\alpha|z_j}^2(\theta_i, u_i) \quad (4.39)$$

and finally we choose the point with the lowest augmented IMSE among 100 randomly sampled points in $\Theta \times \mathbb{U}$.

Figure 4.12 shows the reduction of the augmented-IMSE with the number of additional iterations, until reaching a number of 500 evaluations of the model in total. Abrupt changes in the IMSE can be explained by a significant changes in the hyperparameters of the GP.

Based on the enriched GP, we can maximise the plug-in approximation $\hat{\Gamma}_\alpha^{\text{PI}}$. We chose here $n_u = 500$, and the maximum found is

$$\max_{\theta \in \Theta} \hat{\Gamma}_\alpha^{\text{PI}}(\theta) = 0.9300 \pm 0.0224 \quad (4.40)$$

where the given confidence interval is computed using the normal approximation of the binomial proportion, at a 95% level. As we are using a plug-in approximation, we overlook the intrinsic uncertainty which is represented by the GP, and thus use directly m_Z instead of J .

Even though the design has been enriched for $\alpha = 1.3$, the resulting GP can be used to evaluate quantities associated with other thresholds. Figure 4.14 shows the maximal probability reached by $\hat{\Gamma}_\alpha^{\text{PI}}$ as a function of α , and the different components of the

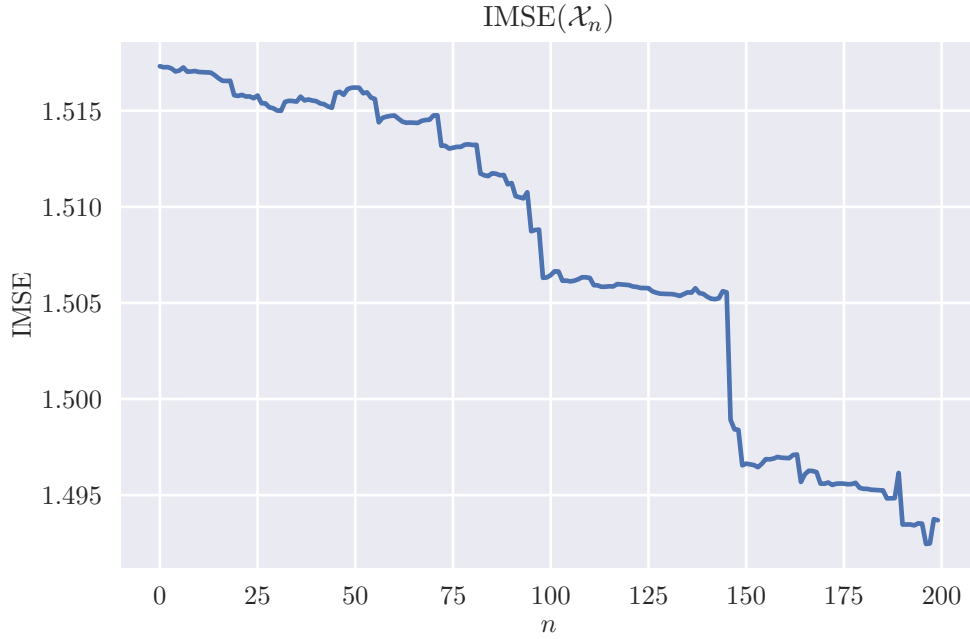


Figure 4.12 – Evolution of the IMSE during the enrichment strategy

control variable $\hat{\theta}_{\text{RR},\alpha}$, all these quantities computed using the metamodel m_Z after the additional iterations.

We can see that the estimated components of $\hat{\theta}_{\text{RR},\alpha}(U)$ stay in a rather small range for all $\alpha > 1$. Recalling the the global optimiser whose values are introduced [Table 4.3 Page 117](#), $\hat{\theta}_{\text{global}} = (-3.516, -5.078, -6.3459)$, we can observe that $\hat{\theta}_{\text{RR},\alpha}^{(1)}$ is slightly higher than $\hat{\theta}_{\text{global}}^{(1)}$ for all $\alpha > 1$. Globally however, we can see that the different values of the components do not seem to change a lot for different α .

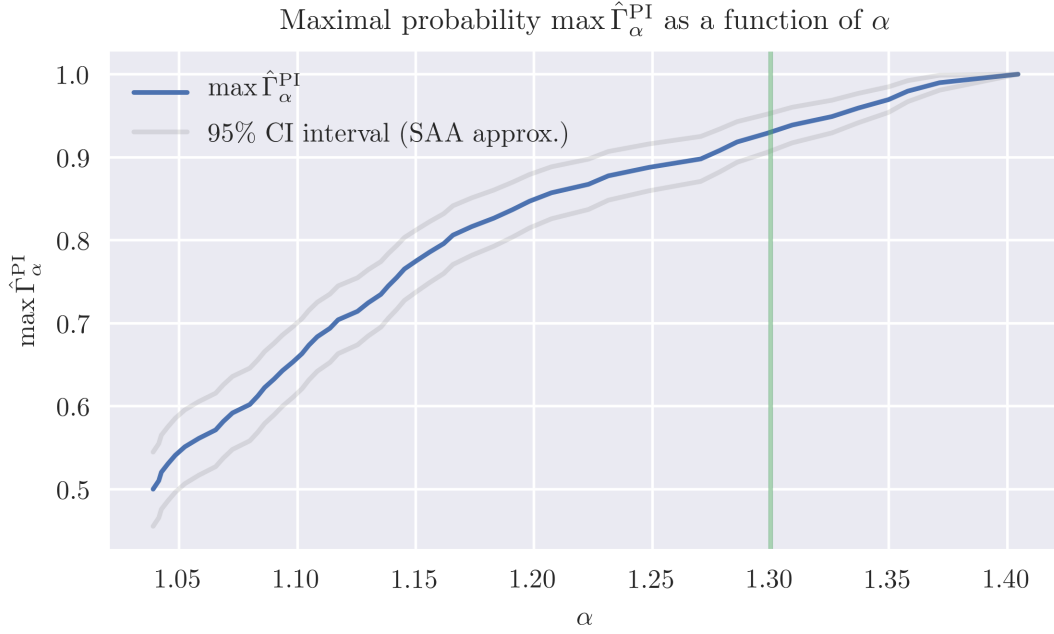


Figure 4.13 – Evolution of the maximal probability of acceptability $\max \Gamma_{\alpha}^{\text{PI}}$, and 95% CI interval associated with the SAA approximation of the estimation of the probability.

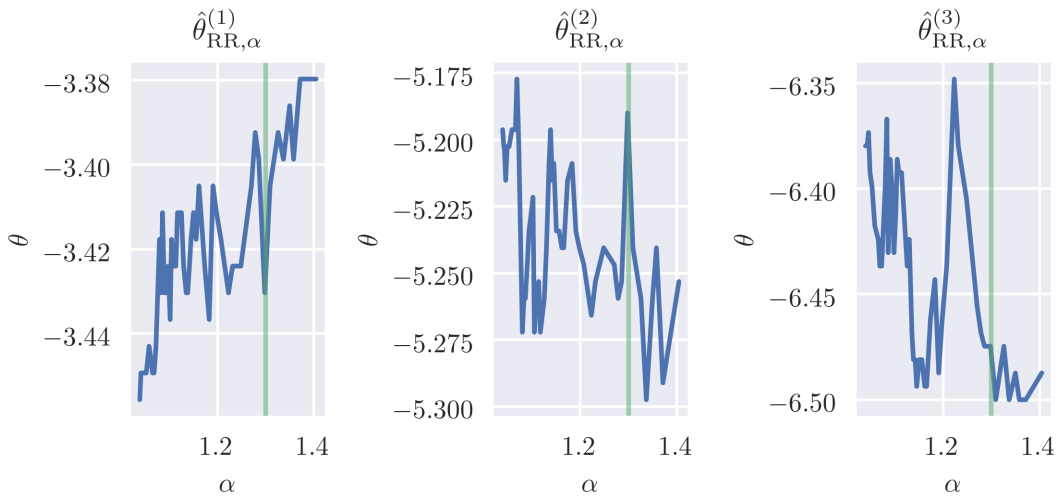


Figure 4.14 – Components of $\hat{\theta}_{\text{RR},\alpha}$, using the GP enriched with 200 additional points minimising the augmented IMSE.

4.5.3.b Optimisation of the quantile of the relative-regret

Plug-in approximation of the quantile

Alternatively, for a level of confidence $p \in [0, 1]$, we can define the quantile function of the ratio:

$$q_p(\theta) = Q_U(\text{RR}(\theta, U); p) \quad (4.41)$$

where $\text{RR}(\theta, U) = \frac{J(\theta, U)}{J^*(U)}$ is the relative-regret, that we introduce for notational convenience. As RR is unknown directly, we can also apply the plug-in approach, and define

$$\text{RR}^{\text{PI}}(\theta, u) = \exp \left[m_{\Xi}(\theta, u) + \frac{1}{2} \sigma_{\Xi}^2(\theta, u) \right] \quad (4.42)$$

where Ξ is the lognormal approximation as defined in [Eq. \(3.87\)](#), [Page 93](#). Once again, the estimation of the quantile of order p of RR is done using a set of i.i.d. samples of U : $\{u_i\}_{1 \leq i \leq n_u}$

$$\hat{q}_p^{\text{PI}}(\theta) = \text{RR}^{\text{PI}}(\theta, u)_{([n_u p])} \quad (4.43)$$

where the subscript indicates the order statistic, *i.e.* the $[n_u p]$ smallest value of $\{\text{RR}^{\text{PI}}(\theta, u_i)\}_{1 \leq i \leq n_u}$ (with $[\cdot]$ as the rounding operator).

As defined in the previous chapter, the minimiser of $\hat{q}_p^{\text{PI}}$ is the estimated member of the relative-regret family associated with α_p :

$$\hat{\theta}_{\text{RR}, \alpha_p} = \arg \min_{\theta \in \Theta} \hat{q}_p^{\text{PI}}(\theta) = (\hat{\theta}_{\text{RR}, \alpha_p}^{(1)}, \hat{\theta}_{\text{RR}, \alpha_p}^{(2)}, \hat{\theta}_{\text{RR}, \alpha_p}^{(3)}) \quad (4.44)$$

Sampling-based method for the estimation of the quantile: QeAK-MCS

We are now going to treat this problem using a sampling-based method, as introduced [Section 3.5.4.c](#), which is derived from [Razaaly \(2019\)](#). The log-normal approximation of the relative-regret, as introduced [Section 3.5.4.a](#), allows us to define the random process Ξ as

$$\log \frac{Z(\theta, u)}{Z^*(u)} \approx \Xi(\theta, u) \sim \mathcal{N}(m_{\Xi}(\theta, u), \sigma_{\Xi}^2(\theta, u)) \quad (4.45)$$

Using the Gaussian nature of Ξ , we define $\mathbf{q}_1$, $\mathbf{q}_2$ and $\mathbf{q}_3$, as a lower bound, a central estimate and an upper bound respectively, of the value $\min_{\theta} \hat{q}_p(\theta) = \hat{\alpha}_p$, minimum which is attained at $\hat{\theta} = \arg \min_{\theta \in \Theta} \hat{q}_p(\theta)$. Following the notation introduced in the previous chapter, we chose $K_q = 3$. With $k = \Phi^{-1}(1 - 0.025)$ the quantile of order 0.975 of the standard normal distribution, we define

$$\log \mathbf{q}_1 = \left(m_{\Xi}(\tilde{\theta}, u) - k \sigma_{\Xi}(\tilde{\theta}, u) \right)_{([n_u p])} \approx Q_U(m_{\Xi}(\tilde{\theta}, U) - k \sigma_{\Xi}(\tilde{\theta}, U); p) \quad (4.46)$$

$$\log \mathbf{q}_2 = \log \left(\min_{\theta \in \Theta} \hat{q}_p^{\text{PI}}(\theta) \right) = \log \hat{q}_p^{\text{PI}}(\tilde{\theta}) = \log \hat{\alpha}_p \quad (4.47)$$

$$\log \mathbf{q}_3 = \left(m_{\Xi}(\tilde{\theta}, u) + k \sigma_{\Xi}(\tilde{\theta}, u) \right)_{([n_u p])} \approx Q_U(m_{\Xi}(\tilde{\theta}, U) + k \sigma_{\Xi}(\tilde{\theta}, U); p) \quad (4.48)$$

We can then define the margins of uncertainty of the relative-regret of level η , for each of the values $\mathbf{q}_l$:

$$\mathbb{M}_\eta(\mathbf{q}_l) = \left\{ (\theta, u) \mid \frac{\eta}{2} \leq \Phi \left(-\frac{m_\Xi(\theta, u) - \log \mathbf{q}_l}{\sigma_\Xi(\theta, u)} \right) \leq 1 - \frac{\eta}{2} \right\} \quad (4.49)$$

We can then sample points in each of those margins, perform a statistical reduction method in order to get $K_{\mathbb{M}}$ points for each margin, adjust them and finally evaluate the $K_q \cdot K_{\mathbb{M}} = 3 \cdot K_{\mathbb{M}}$ points by the objective function.

Numerically speaking, the margins revealed themselves quite small, so we chose a level $\eta = 0.0005$ in order to increase their volume, and thus facilitate the sampling procedure.

After the sampling of 2000 points in those margins $\mathbb{M}_\eta(\mathbf{q}_l)$ for $1 \leq l \leq 3$, we use the KMeans clustering algorithm in order to select $K_{\mathbb{M}}$ *representative* points among those samples. This number of points selected in each margin $K_{\mathbb{M}}$ changes along the iterations: it starts with a small number of points, $K_{\mathbb{M}} = 3$, and increases until $K_{\mathbb{M}} = 10$.

We can see on Fig. 4.15 the estimated volume of the margin of uncertainty (*i.e.* its measure according to Lebesgue's measure on $\Theta \times \mathbb{U}$), using Monte-Carlo method, along with the 95% confidence intervals associated with the Monte-Carlo estimation (Wilson's score method as introduced Wilson (1927) to account for the small relative volume of the margin of uncertainty). We can notice that first, its volume increases slightly. This can be seen as an exploration phase, where the additional evaluations change the hyperparameters of the GP, which affects the estimation of the conditional minimums, and the candidate quantiles $\mathbf{q}_l$, for $1 \leq l \leq 3$. After enough additional points, the volume of the margin decreases, as points are added that do not change significantly the other estimations.

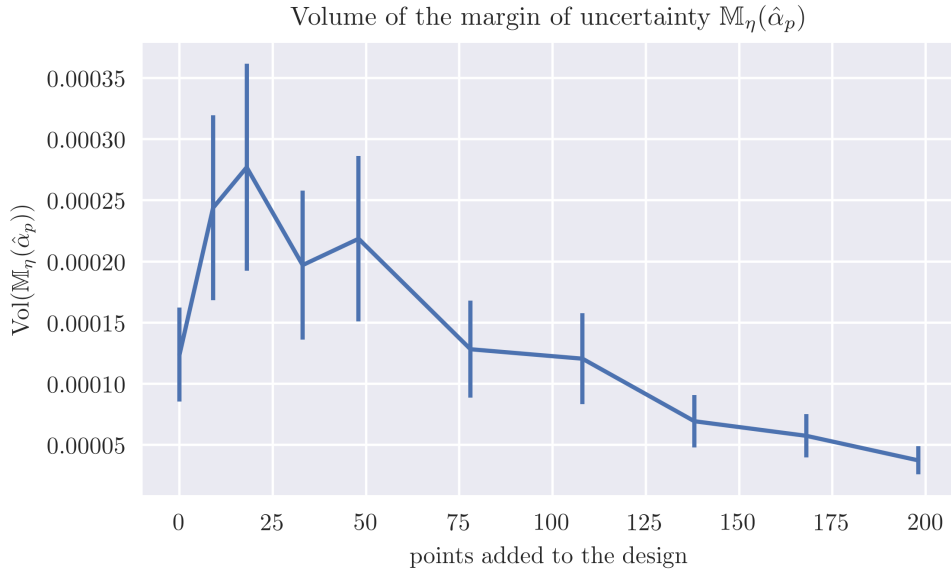


Figure 4.15 – Volume of the margin of uncertainty associated with $\mathbf{q}_2 = \hat{\alpha}_p$, and 95% confidence intervals of its estimation

Using the GP conditioned on the final design, for $p = 0.95$ we can compute

$$\begin{aligned}\hat{\alpha}_p &= 1.3791 \\ [\mathbf{q}_1, \mathbf{q}_3] &= [1.36618, 1.391868]\end{aligned}\tag{4.50}$$

and using the same GP, Fig. 4.16 shows the estimation of the threshold for other values of p . One main difference between the confidence interval given there by $\mathbf{q}_1$ and $\mathbf{q}_3$, and the one given Eq. (4.40) is that the former translates the uncertainty originating from the Gaussian Process, while the latter accounts for the error in the *sample average approximation* of the probability $\hat{\Gamma}_\alpha^{\text{PI}}$.

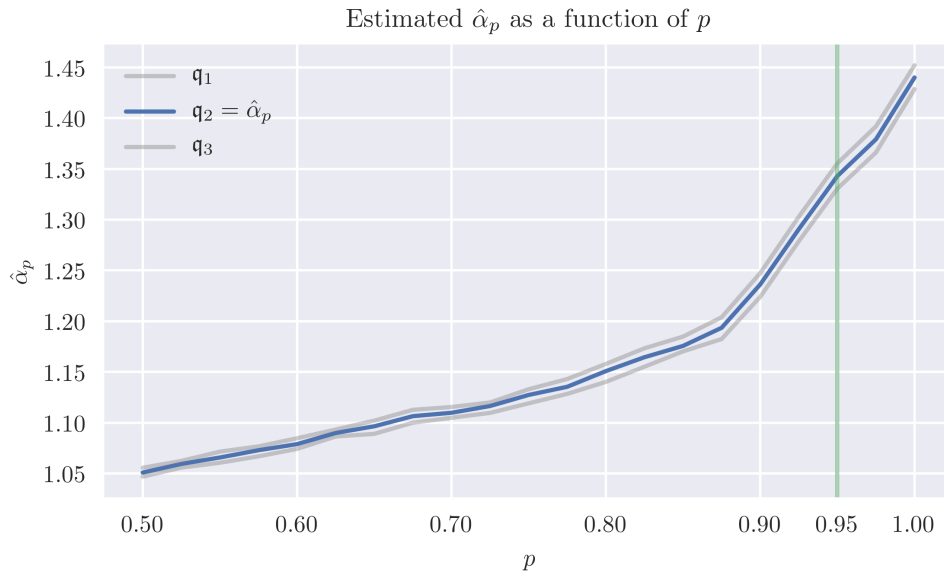
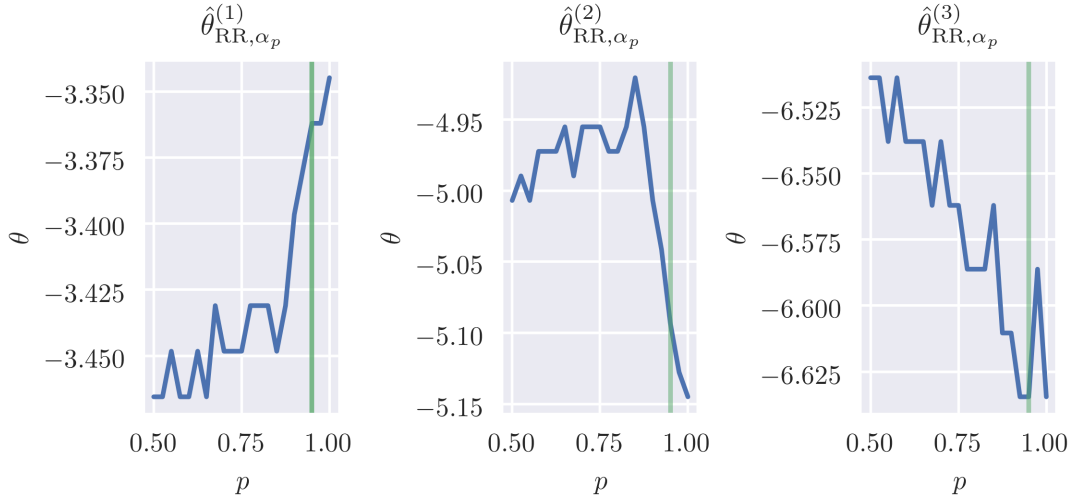


Figure 4.16 – Estimation of the threshold α_p , depending on the level p . The lower and upper bounds $\mathbf{q}_1$ and $\mathbf{q}_3$ are also displayed

Finally, Fig. 4.17 shows the different components of $\hat{\theta}_{\text{RR},\alpha_p}$ for different $p > 0.5$. We can notice that similarly as in Fig. 4.14, the range of values taken by each component is rather small. We can however notice that the first component of $\hat{\theta}_{\text{RR},\alpha_p}$ is again slightly larger than the one the global optimum $\hat{\theta}_{\text{global}}$.

All in all, this study suggests that choosing a value slightly higher than the global optimiser for the first component (which corresponds to Pebbles) would be more robust than choosing the global minimiser. Regarding the other components, their respective influence (given by the SA) and the relative-regret estimates would point toward a value close or equal to the global optimiser.

A comparison of the two methods is summarised in Table 4.4.

Figure 4.17 – Relative-regret based estimates $\hat{\theta}_{RR, \alpha_p}$, depending on the level p

Method	augmented IMSE	Qe-AK MCS
Quantity of interest	$\mathbb{E}_Z[\int_{\Theta \times \mathbb{U}} \sigma_{\Delta_\alpha Z}^2]$	$\mathbb{M}_\eta(\mathbf{q}_l)$
Type	1-step	K -step
Main Bottleneck	Evaluate and optimise integral in $(1 + \dim(\Theta \times \mathbb{U}))$ dimensions	Sampling in unknown regions in $\dim(\Theta \times \mathbb{U})$
Advantage	Optimal criterion of enrichment	K chosen arbitrary
$\hat{\theta}_{RR, \alpha} \alpha = 1.3$	$(-3.43, -5.2, -6.48)$	$(-3.375, -5.05, -6.61)$
$\hat{\theta}_{RR, \alpha_p} p = 0.95$	$(-3.39, -5.28, -6.5)$	$(-3.36, -5.10, -6.63)$
$\hat{\theta}_{\text{global}}$	$(-3.516, -5.078, -6.3459)$	

Table 4.4 – Comparison of methods and numerical results for the robust calibration of CROCO

4.6 Partial conclusion

In this chapter, we addressed the problem of calibration under uncertainties of the numerical model CROCO. After having defined the control and environmental parameters, we performed different studies, in order to have a better understanding of the behaviour of the model at stake.

First, we optimised the objective function on a high-dimensional input space, but without external uncertainties. We chose then to segment the domain according to the type of sediments that can be found at the bottom, and by performing a sensitivity analysis, we could reduce further the dimension of the input space. A similar study has been done in order to reduce the dimension of the environmental parameter as well.

Finally, we applied some of the methods that rely on GP as introduced in [Chapter 3](#), in order to explore further the behaviour of the numerical model under uncertainties.

First, to have a better estimation of the conditional minimum and minimisers we enriched the design using the PEI criterion. Then, in order to get robust estimators of the bottom friction, we used two different approaches to enrich the GP. In the first one, we reduced iteratively the augmented IMSE in order to improve the plug-in approximation of the probability of exceeding a threshold. For the other one, in order to minimise a specific quantile of the relative-regret, we defined and sampled in the margin of uncertainty in order to add points by batches.

This finally leads to relative-regret estimates of the bottom friction, which differ from the global minimiser: according to this study, taking a value slightly larger for the first component of the bottom friction leads seemingly to a more robust solution. The small variations in the values of $\hat{\theta}_{RR}$ for the different levels of confidence also suggest that a compromise between performances and robustness is easily reachable. Indeed, in this configuration, we can see that the calibration problem behaves “nicely” under uncertainties: the environmental parameters seem to only have a limited impact on the estimation of the bottom friction.

From a performance point of view, GP can indeed reduce the computational requirements needed to compute relative-regret estimates. However as mentioned before, GP do not scale particularly well for problems of more than a dozen input variables, rendering a dimension reduction almost mandatory for tractability.

The computational cost of the enrichment process, either using stepwise or sampling procedures, can quickly become non-negligible. Indeed, the estimation of the parameters of the distributions of $\Delta_\alpha(\theta, u)$ and $\Xi(\theta, u)$ require at each point (θ, u) a global optimisation of m_Z , which can amount to a few hundreds of calls to the surrogate m_Z . In itself, this procedure is not that expensive. However, it needs to be performed a large number of times: for a single evaluation of the augmented IMSE for instance, estimation of nested integrals are needed, thus the number of computations of the parameters of Δ_α grows quite large, even more when considering an optimisation of the augmented IMSE on $\Theta \times \mathbb{U}$. Similarly, the volume of the margin of uncertainty can be very small compared of the volume of the whole space $\Theta \times \mathbb{U}$, thus without specific sampling schemes, it can quickly become a computational bottleneck as well, even more so when the prediction variance becomes small.

Some improvements can be considered in order to perform more efficiently those estimations. A two-stage approach could be derived, in order to first identify a value $\tilde{\theta}$ for which we want to reduce the augmented IMSE, and secondly to integrate the augmented IMSE over $\tilde{\theta} \times \mathbb{U}$. By doing so, each iteration would instead require the optimisation of an integral of dimension $1 + \dim \mathbb{U}$, while at the same time focus some computational effort toward the estimation of $\hat{\theta}_{RR,\alpha}$, by correctly choosing a candidate $\tilde{\theta}$. Also, instead of considering the augmented IMSE, we could instead look for the point maximising the prediction variance of Δ_α , and “adjust” its θ component, similarly as proposed for the sampling based scheme, bypassing completely the evaluation of the integral.

A better sampling scheme for Monte-Carlo based methods can also be imagined: the margin of uncertainty comprises points (θ, u) for which $\Xi(\theta, u)$ is “close” to $\log \mathbf{q}_l$ (for $1 \leq l \leq K_q$), thus for importance sampling, we could construct a proposal density which consists in first sampling $u \sim U$, then sampling θ “close” to $\theta_Z^*(u)$, for instance by

choosing $\theta \sim \mathcal{N}(\theta_Z^*(u), \gamma^2 I)$, where I is the identity matrix in dimension $\dim \Theta$, and γ is chosen with respect to the target quantiles $\mathbf{q}_l$.

Both methods introduced in this chapter rely on the ability to compute quickly the conditional minimum m_{Z^*} and the conditional minimiser θ_Z^* in order to get the mean and variance of the processes Δ_α and Ξ . For each of this optimisation, we used a gradient-descent methods, where the starting point is selected randomly in the search space each restart, in order to avoid getting stuck in a local minimiser. This step could be improved as well, by using directly a global optimisation method.

Finally, aside from those technical details of implementation, a similar study could be derived on slightly different configurations, by increasing for instance the assimilation window, in order to increase the observable interactions in the tide constituents, or by considering a different segmentation with more sediment types, or a segmentation based on the depth for instance.

* * *

CONCLUSION AND PERSPECTIVES

Summary

Numerical models are ubiquitous nowadays for the prediction of various phenomena, which have a large impact especially for policy- and decision-making. Due to the increase in computing power and to the various advances in research that allow the quantification of additional effects or interactions that were overlooked until then, realistic models may become increasingly complex, as additional parameters are taken into account. In order to get a meaningful representation of the reality, those parameters have to be chosen accordingly. Furthermore, the numerical model studied may also be affected by some random inputs, making the calibration more difficult. In this thesis, we tackled the estimation problem of parameters under uncertainty. More precisely, we studied the notion of robustness of a calibrated model when a random variable which represents some uncontrollable environmental conditions is taken into account.

In [Chapter 1](#), after having covered common notions of probability and aspects of statistical inference, we described the calibration problem as an optimisation problem by introducing an objective function, that we wish to minimise with respect to the control parameter θ .

However, due to the presence of some random environmental variable U in the study, a plain minimisation of the objective function is not completely relevant, as its result will depend on the value taken by this random variable. Considering the random nature of the environmental parameters, the calibration can be seen as a problem of *optimisation under uncertainties*, and many specific methods and criteria can be defined to treat accordingly this new problem of *robust* calibration.

Some classical criteria are first introduced in [Chapter 2](#). Depending on the framework used to describe the quantities involved, we can define Bayesian or frequentist estimates by keeping a fully probabilistic inference framework. On the other hand, when considering

a variational approach, estimates such as the minimiser of the mean value of the objective function, *i.e.* the minimiser of the *expected loss*, are often used.

In this thesis, we focused on estimates based on the notion of regret: instead of comparing directly the values taken by the objective function for different configurations given by u , the regret allows the modeller to compare the value of the objective with the best attainable performance given this specific environmental variable. This allows to put less emphasis on configurations which lead to bad performances, and to focus more on “salveagable” situations.

Moreover, the user can adjust a parameter in order to reflect either a risk-averse behaviour, by favourising a control of the regret with high probability, or a risk-seeking one, by favourising an estimate that will yield values of the regret closer to its optimum, albeit with lower probability.

In general, criteria of robust optimisation require a global knowledge of the function, since they often involve several evaluations of expectations and probabilities with respect to U . In addition to that, the family of regret-based estimates we introduced depends directly on the conditional minimum and minimiser. In [Chapter 3](#), we proposed to use Gaussian Processes (GP) in order to compute the quantities associated with regret-based estimators. More precisely, we proposed a few methods which aim at improving this estimation by choosing iteratively a new, or a batch of new points to evaluate and to add to the design.

Finally in [Chapter 4](#), we studied an academic problem of calibration of a coastal model based on CROCO. After having reduced the dimension of the input space based on the sediment type at the bottom, we enriched the design in order to improve the estimation of the functions that define the regret-based estimates, which are then optimised.

Limitations and perspectives

Throughout this thesis, we assumed that the forward operator was a deterministic black-box, or *deterministic simulator*, in the sense that the uncertainties in the modelling are “controlled” by the modeller. A *stochastic simulator* in contrast does not take an environmental parameter as input, so each evaluation of the forward operator can be seen as sampling a specific random variable. This may be the case when the environmental variables are not easily parametrised, such as in the presence of functional inputs ([El Amri, 2019](#)). In this case, it is not possible to control the value of u chosen for an enrichment strategies for instance. An alternative strategy would be to consider this uncertainty as *noise* in the output of the simulator, leading to *noisy kriging* and/or noisy optimisation methods ([Picheny and Ginsbourger, 2014](#)).

Based on the assumption that the environmental parameter u was indeed controllable, using the properties of the Gaussian processes we developed criteria which aim at improving the estimations of the functions Γ_α and q_p . In order to get the regret-based estimators, those functions had to be optimised, using a set of samples of U to approximate expectations and quantiles, in a sample average approximation (SAA) fashion. Because of this, for levels of confidence very close to 1, the estimation of quantities in such

high probabilities can be difficult, and usually needs specific methods instead of simple Monte-Carlo sampling in order to get accurate enough results, as in [Razaaly \(2019\)](#).

Also, instead of improving the estimation of Γ_α and of q_p on the whole space, we could develop methods in order to optimise directly those functions, and thus combine the estimation and the optimisation. This could for instance be done in a 2-stage enrichment strategy, as in [Janusevskis and Le Riche \(2010\)](#). First, a value $\tilde{\theta} \in \Theta$ is chosen, with a “high” potential to be the optimiser (quite similarly as the EI criterion), and then the couple (θ_{n+1}, u_{n+1}) to evaluate is chosen in order to reduce a measure of uncertainty associated with the space $\{\tilde{\theta}\} \times \mathbb{U}$ (*e.g.* the IMSE integrated over this space). This would focus the enrichment in regions of interest, and also reduce the dimension of the integral to evaluate.

We introduced also a few sampling methods, that rely on efficient sampling in different margin of uncertainties. Depending on the problem, this task can reveal itself quite challenging: the more accurate the GP is, the less prediction variance it shows, and thus the margin of uncertainty can become very thin. Simple sampling schemes such as the acceptance-rejection method can become inefficient, even more so when the dimension of the problem increases. Refinement such as importance sampling can be considered such as in [Razaaly \(2019\)](#).

Once set aside potential technical improvements in the methods proposed, we can propose a few other possibilities that may warrant further investigations.

In this thesis we focused on the variational formulation of the calibration problem, *i.e.* by defining an objective function, akin to the log-likelihood, or to the log-posterior density that is then optimised. A Bayesian method could be performed in order to estimate the posterior distribution of the parameter θ given the observations. Moreover, this task could be performed using adapted sampling schemes, such as Hamiltonian Monte-Carlo ([Betancourt, 2017](#)), in order to use the gradient that may be available using adjoint method.

This gradient, taken with respect to the control parameter, may be useful at different stages: we can incorporate this additional knowledge in the modelling of the GP, so that the predictions are improved, as done in [Bouhlef and Martins \(2019\)](#); [Laurent et al. \(2019\)](#). However, GP are flexible and useful tools but are not well suited for modelling problems with dimensions higher than about 10: when the designs considered are too large, fitting the GP can also be problematic, as large matrices have to be inverted, and the optimisation of the hyperparameters can become difficult. Reducing the dimension of the input spaces is then often necessary. The segmentation we performed in [Chapter 4](#) for instance is rather coarse and based on external information. A finer dimension reduction method could be done, without prior information, using the gradient for instance, as done in [Benameur et al. \(2002\)](#) or in [Zahm et al. \(2018\)](#).

Finally, in this work, the random variables $J^*(U)$ and $\theta^*(U)$ which represent the conditional minimums and minimisers respectively, play a central role in our definitions of robustness, and further developements could be imagined around those random variables. The distribution of the conditional minimum $J^*(U)$ can be seen as an “ideal” distribution for the objective (*i.e.* the distribution that one could get if the calibrated parameter

had always been chosen optimally for all realisations of the uncertain variable). The minimisation of a measure of misfit between the distribution of $J^*(U)$ and the distribution of $J(\hat{\theta}, U)$ (the objective function of the model calibrated with $\hat{\theta}$) may be worth exploring. Horsetail matching for instance ([Cook et al., 2017](#); [Cook and Jarrett, 2018](#)) could be used so that the a metric between cdf of the two distributions is minimised.

* * *

APPENDIX

Contents

A.1 Lognormal approximation of the ratio of normal random variables	135
A.2 Full sediment repartition in the Bay of Biscay and the English Channel	137
A.3 Misspecified deterministic calibration	137

A.1 Lognormal approximation of the ratio of normal random variables

Let X and Y be two correlated normal random variables:

$$X \sim \mathcal{N}(\mu_X, \sigma_X^2) \tag{A.1}$$

$$Y \sim \mathcal{N}(\mu_Y, \sigma_Y^2) \tag{A.2}$$

$$\rho = \frac{\text{Cov}[X, Y]}{\sigma_X \sigma_Y} \tag{A.3}$$

$$\begin{bmatrix} X \\ Y \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mu_X \\ \mu_Y \end{bmatrix}; \begin{bmatrix} \sigma_X^2 & \rho \sigma_X \sigma_Y \\ \rho \sigma_X \sigma_Y & \sigma_Y^2 \end{bmatrix} \right) \tag{A.4}$$

Let us assume that both random variable are positive with high probability, so that the ratio $T = X/Y$ is defined and positive with high probability. We will first rewrite T using the centered random variables $X_0 = X - \mu_X$ and $Y_0 = Y - \mu_Y$:

$$T = \frac{\mu_X}{\mu_Y} \frac{1 + \frac{X_0}{\mu_X}}{1 + \frac{Y_0}{\mu_Y}} \tag{A.5}$$

and we can then take its logarithm

$$\log T = \log \left(\frac{\mu_X}{\mu_Y} \right) + \log \left(1 + \frac{X_0}{\mu_X} \right) - \log \left(1 + \frac{Y_0}{\mu_Y} \right) \tag{A.6}$$

So if X_0/μ_X and Y_0/μ_Y are small, i.e. if $\mu_X \gg \sigma_X$ and $\mu_Y \gg \sigma_Y$, by Taylor's expansion of the log, we have

$$\log T \approx \log \left(\frac{\mu_X}{\mu_Y} \right) + \frac{X_0}{\mu_X} - \frac{Y_0}{\mu_Y} \quad (\text{A.7})$$

The joint distribution of X_0 and Y_0 is known:

$$\begin{bmatrix} X - \mu_X \\ Y - \mu_Y \end{bmatrix} = \begin{bmatrix} X_0 \\ Y_0 \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}; \begin{bmatrix} \sigma_X^2 & \rho\sigma_X\sigma_Y \\ \rho\sigma_X\sigma_Y & \sigma_Y^2 \end{bmatrix} \right) \quad (\text{A.8})$$

and by scaling by $1/\mu_X$ and $1/\mu_Y$, we have

$$\begin{bmatrix} \frac{1}{\mu_X} & 0 \\ 0 & \frac{1}{\mu_Y} \end{bmatrix} \begin{bmatrix} X_0 \\ Y_0 \end{bmatrix} = \begin{bmatrix} X_0/\mu_X \\ Y_0/\mu_Y \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}; \begin{bmatrix} \frac{\sigma_X^2}{\mu_X^2} & \rho \frac{\sigma_X\sigma_Y}{\mu_X\mu_Y} \\ \rho \frac{\sigma_X\sigma_Y}{\mu_X\mu_Y} & \frac{\sigma_Y^2}{\mu_Y^2} \end{bmatrix} \right) \quad (\text{A.9})$$

Thus

$$\frac{X_0}{\mu_X} - \frac{Y_0}{\mu_Y} \sim \mathcal{N} \left(0; \frac{\sigma_X^2}{\mu_X^2} + \frac{\sigma_Y^2}{\mu_Y^2} - 2\rho \frac{\sigma_X\sigma_Y}{\mu_X\mu_Y} \right) \quad (\text{A.10})$$

A first approximation is then to consider the ratio to be log-normally distributed, and

$$\log T \sim \mathcal{N} \left(\log \left(\frac{\mu_X}{\mu_Y} \right); \frac{\sigma_X^2}{\mu_X^2} + \frac{\sigma_Y^2}{\mu_Y^2} - 2\rho \frac{\sigma_X\sigma_Y}{\mu_X\mu_Y} \right) \quad (\text{A.11})$$

The correlation term can also be possibly dropped from the approximation, as $\rho \frac{\sigma_X\sigma_Y}{\mu_X\mu_Y}$ can be very close to zero. Moreover, as this term is positive, if omitted, this would increase slightly the variance of the log-ratio and thus increase a bit the uncertainty which can be beneficial in some applications.

A.2 Full sediment repartition in the Bay of Biscay and the English Channel

In [Chapter 4](#), we used the data from the SHOM in order to derive a reduced sediment repartition in [Fig. 4.3](#). The original repartition is reproduced here in [Fig. A.1](#), and the correspondence table between the reduced and full repartition is given [Table A.1](#).

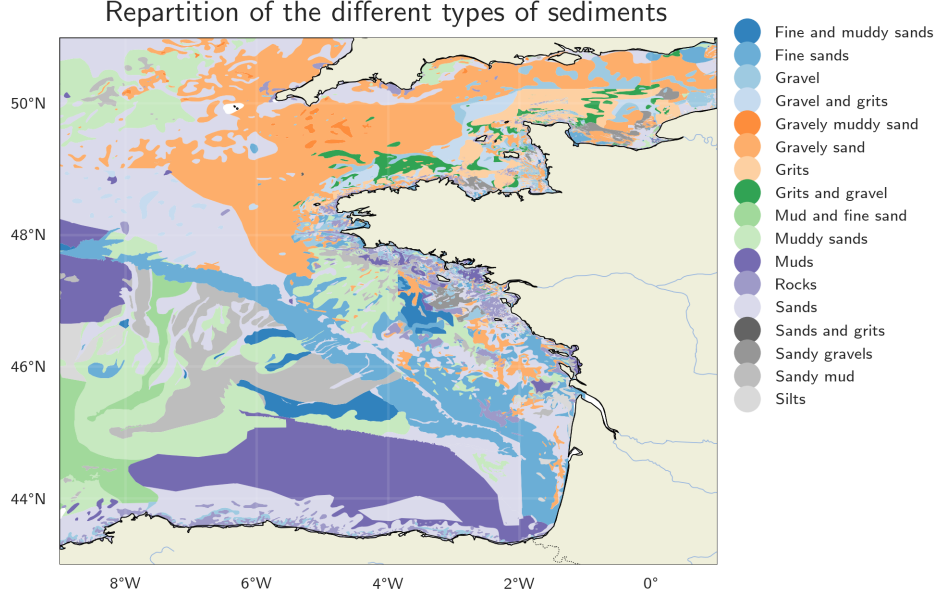


Figure A.1 – Full sediments repartition. Source: SHOM, under CC BY-SA-4.0 license

A.3 Misspecified deterministic calibration

In the next pages are shown the results of different optimisation procedures on the whole domain, similarly as in [Section 4.3.3](#). However, the optimisation will be carried using the objective function defined as

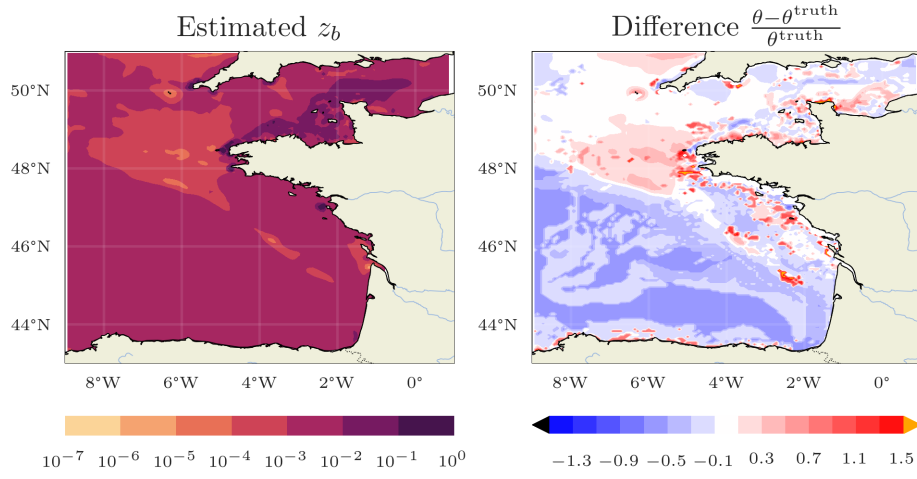
$$\theta \mapsto J(\theta, u^b) = \|\mathcal{M}(\theta, u^b) - y\|_2^2 = \|\mathcal{M}(\theta, u^b) - \mathcal{M}(\theta^{\text{truth}}, u^{\text{truth}})\|_2^2 \quad (\text{A.12})$$

and on 200 iterations. $u^b \in [0, 1]^2$ represents the environmental conditions, which can be misspecified, while $u^{\text{truth}} = (0.5, 0.5)$ is the value used to generate the observations. The top figures show the estimated roughness z_b , and the relative difference of the control variable θ with the truth value θ^{truth} . The bottom plots show the evolution of the objective function and the squared norm of the gradient of the objective function depending on the iteration during the optimisation.

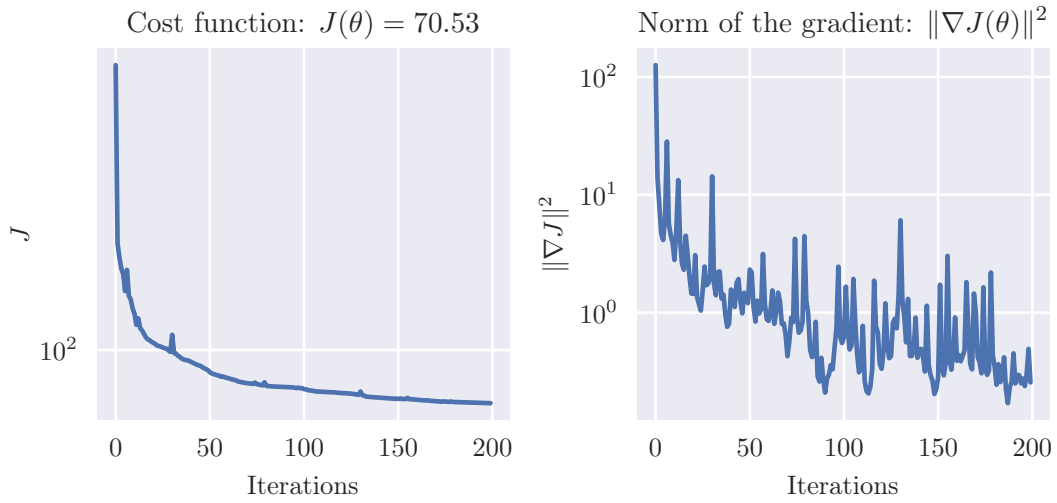
Full repartition	Reduced repartition	Full repartition	Reduced repartition
Rocks	Rocks	Fine sands	Fine Sands
Grits	Pebbles	Fine and muddy sands	
Grits and gravel		Silts	Silts
Gravel	Gravels	Muds	Muds
Gravel and grits		Sandy mud	
Sandy gravels		Mud and fine sand	
Sands	Sands		
Sands and grits			
Gravelly sand			
Gravelly muddy sand			
Muddy sands			

Table A.1 – Correspondence between the full sediments repartition classes and the segmentation used in [Chapter 4](#)

Misspecified environmental parameter: $u^b = (0, 0)$



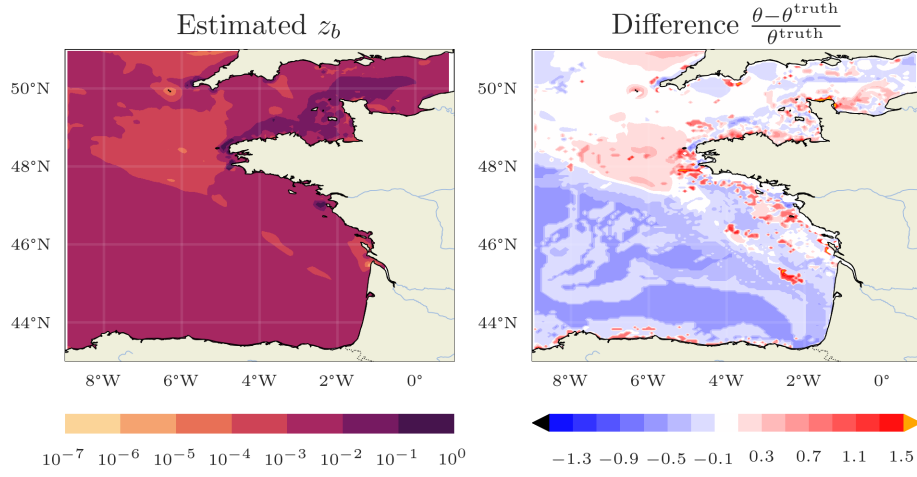
(a) Estimated θ and relative-error



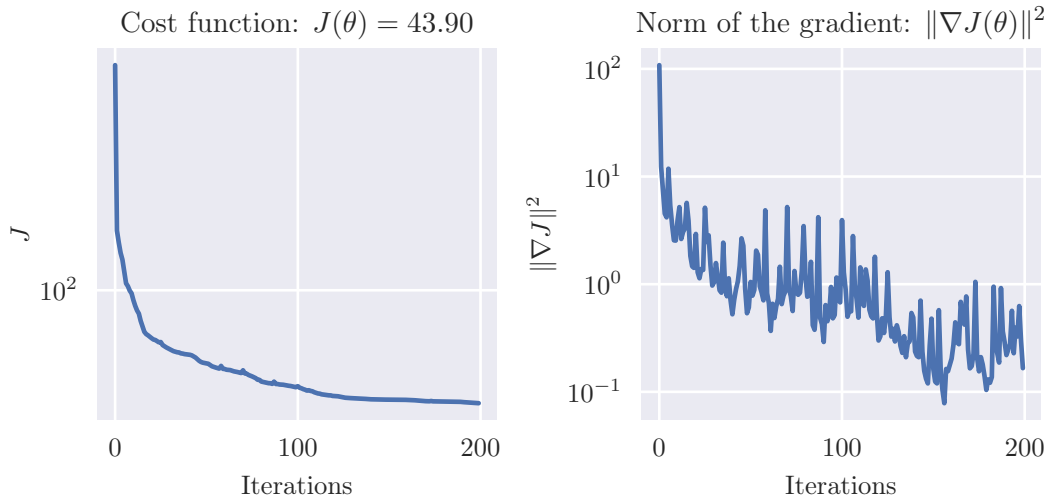
(b) Objective function evolution along the iterations

Figure A.2 – Optimisation on the whole space, $u^b = (0, 0)$

Misspecified environmental parameter: $u^b = (0, 0.5)$



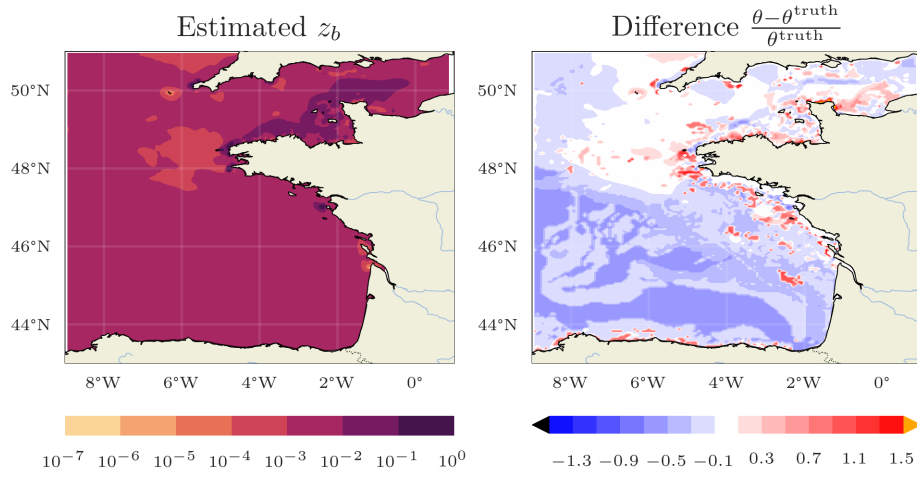
(a) Estimated θ and relative-error



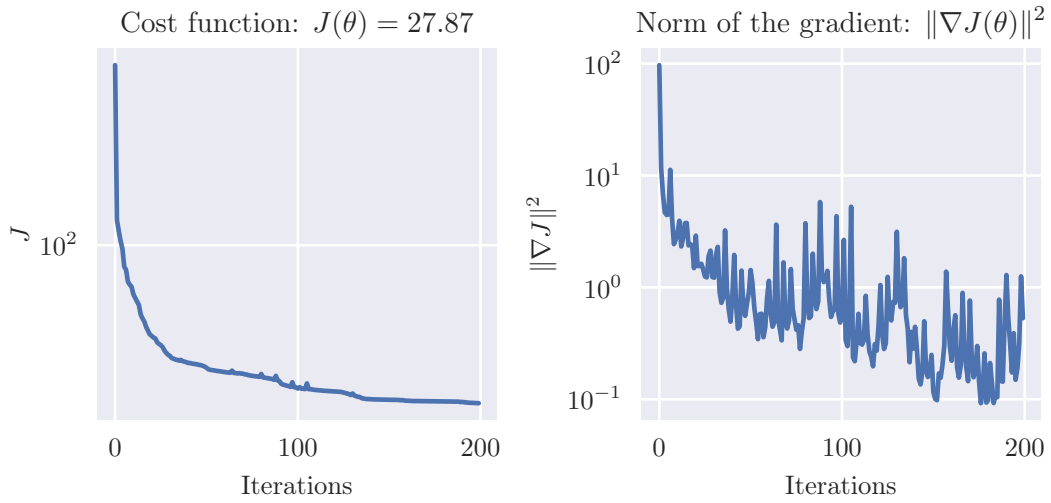
(b) Objective function evolution along the iterations

Figure A.3 – Optimisation on the whole space, $u^b = (0, 0.5)$

Misspecified environmental parameter: $u^b = (0, 1)$



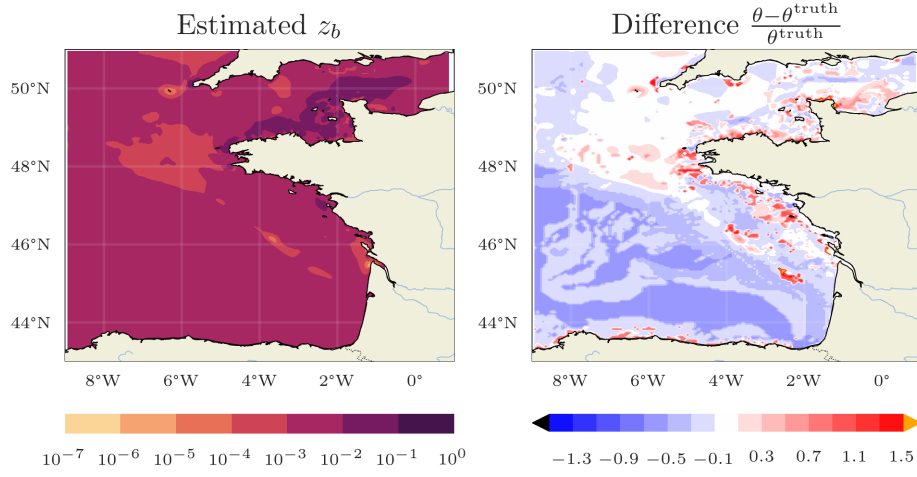
(a) Estimated θ and relative-error



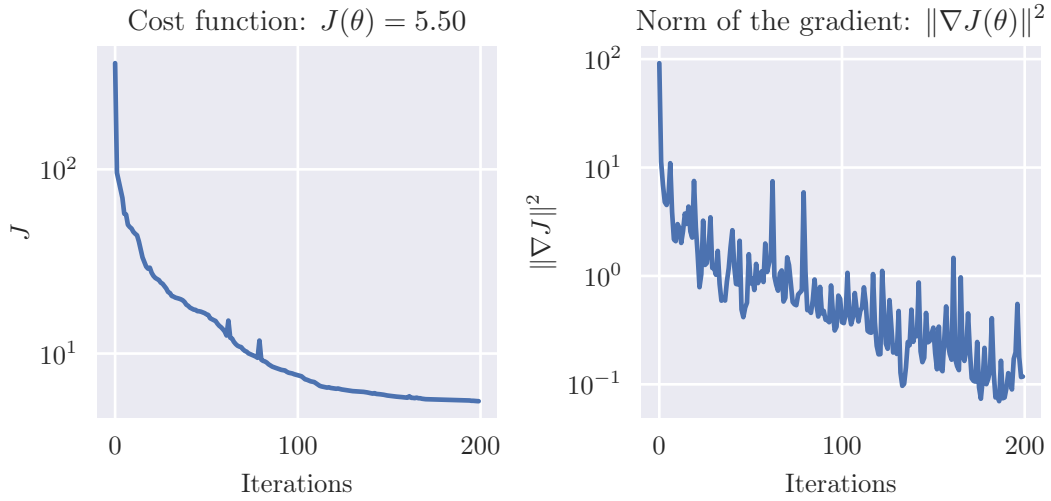
(b) Objective function evolution along the iterations

Figure A.4 – Optimisation on the whole space, $u^b = (0, 1)$

Misspecified environmental parameter: $u^b = (0.5, 0)$



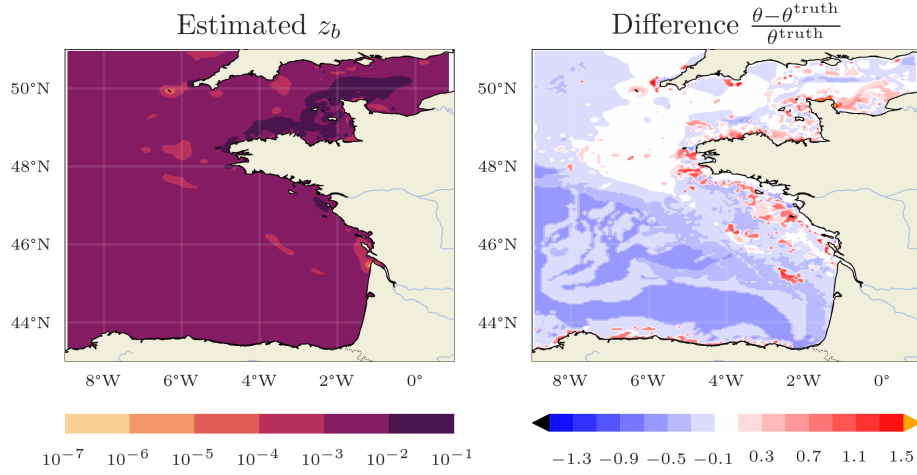
(a) Estimated θ and relative-error



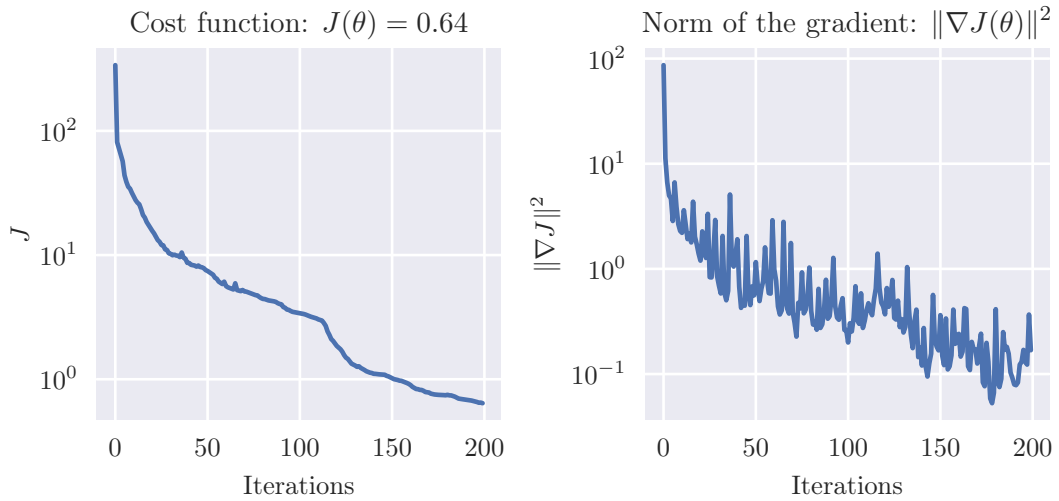
(b) Objective function evolution along the iterations

Figure A.5 – Optimisation on the whole space, $u^b = (0.5, 0)$

Well-specified environmental parameter $u^b = (0.5, 0.5) = u^{\text{truth}}$



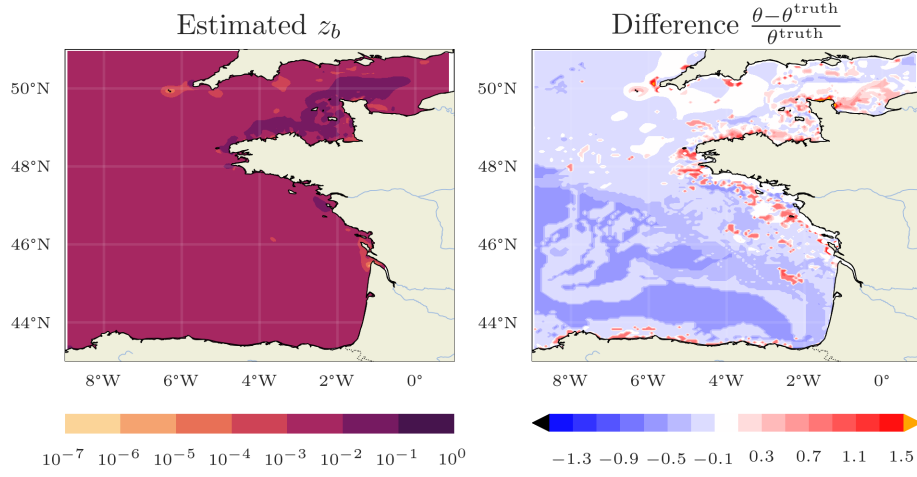
(a) Estimated θ and relative-error



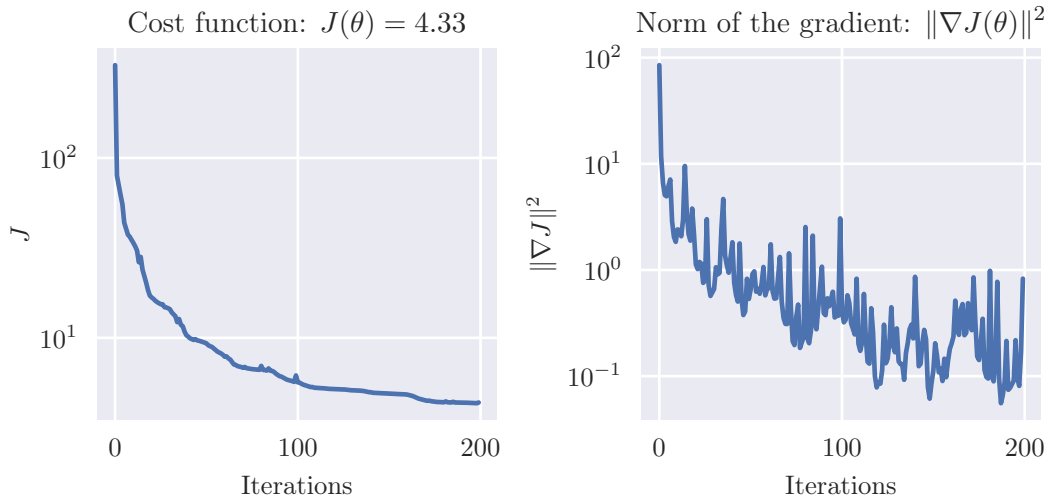
(b) Objective function evolution along the iterations

Figure A.6 – Optimisation on the whole space, $u^b = (0.5, 0.5) = u^{\text{truth}}$

Misspecified environmental parameter: $u^b = (0.5, 1)$



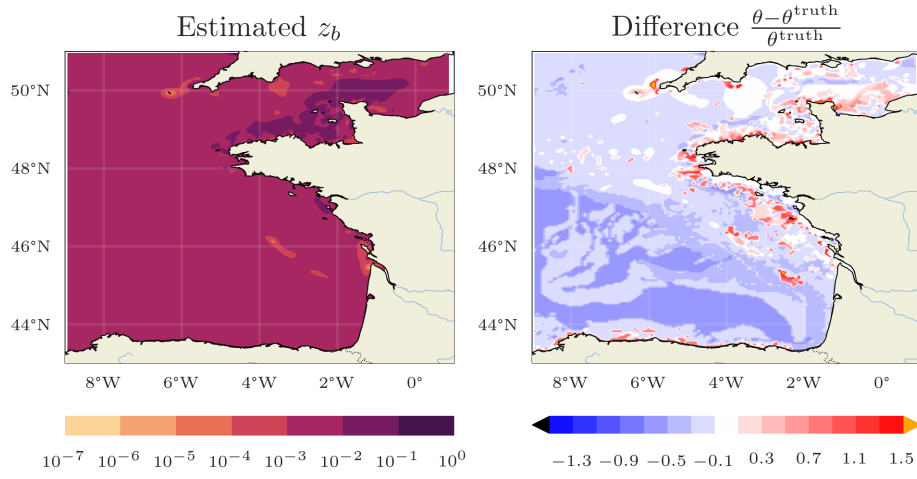
(a) Estimated θ and relative-error



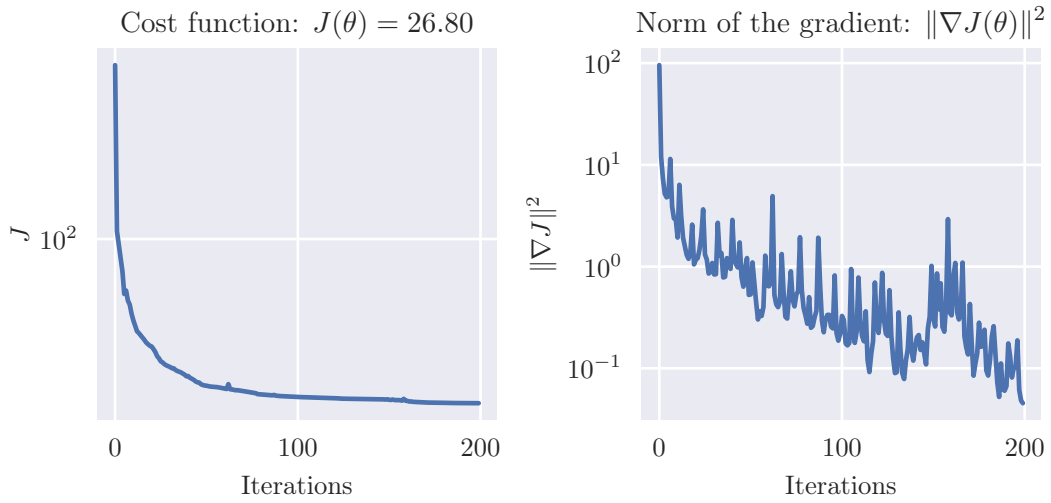
(b) Objective function evolution along the iterations

Figure A.7 – Optimisation on the whole space, $u^b = (0.5, 1)$

Misspecified environmental parameter: $u^b = (1, 0)$



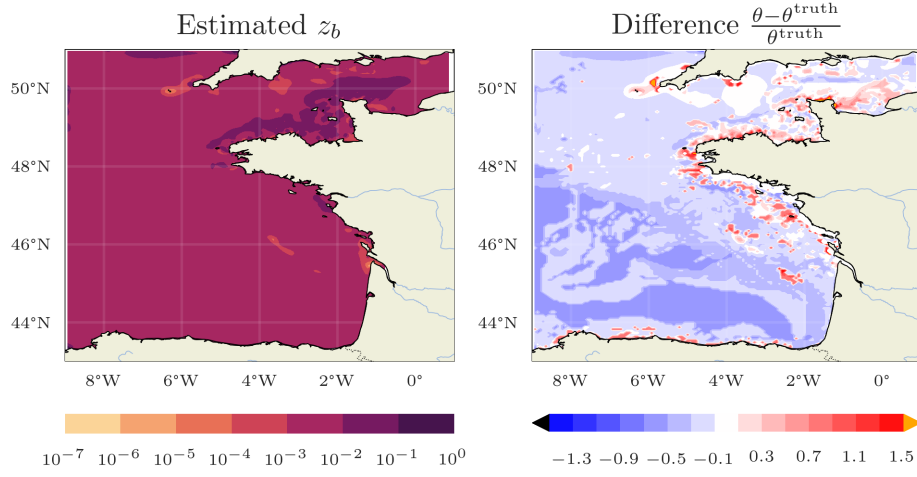
(a) Estimated θ and relative-error



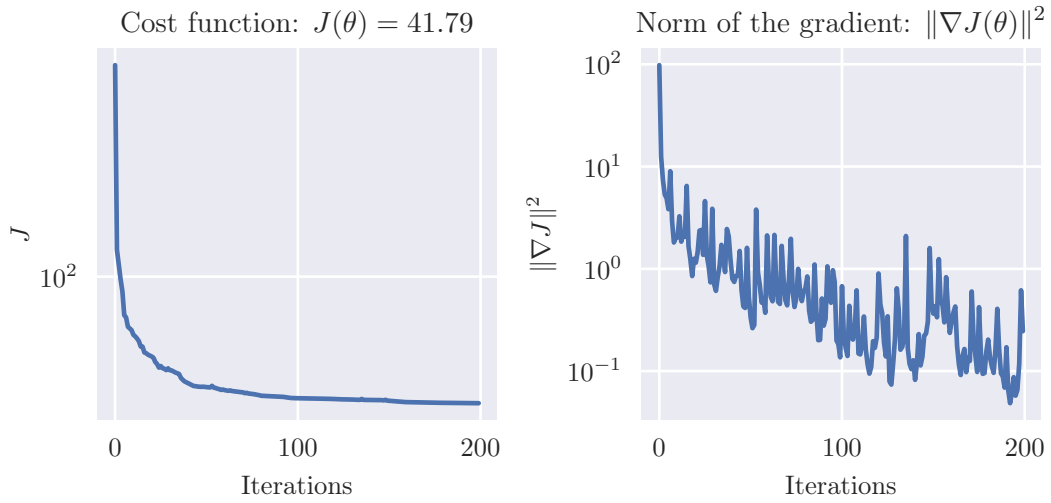
(b) Objective function evolution along the iterations

Figure A.8 – Optimisation on the whole space, $u^b = (1, 0)$

Misspecified environmental parameter: $u^b = (1, 0.5)$



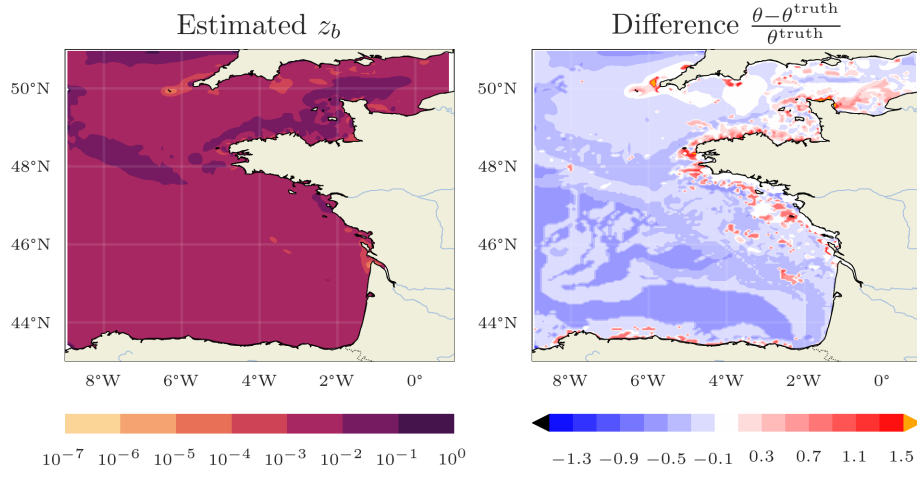
(a) Estimated θ and relative-error



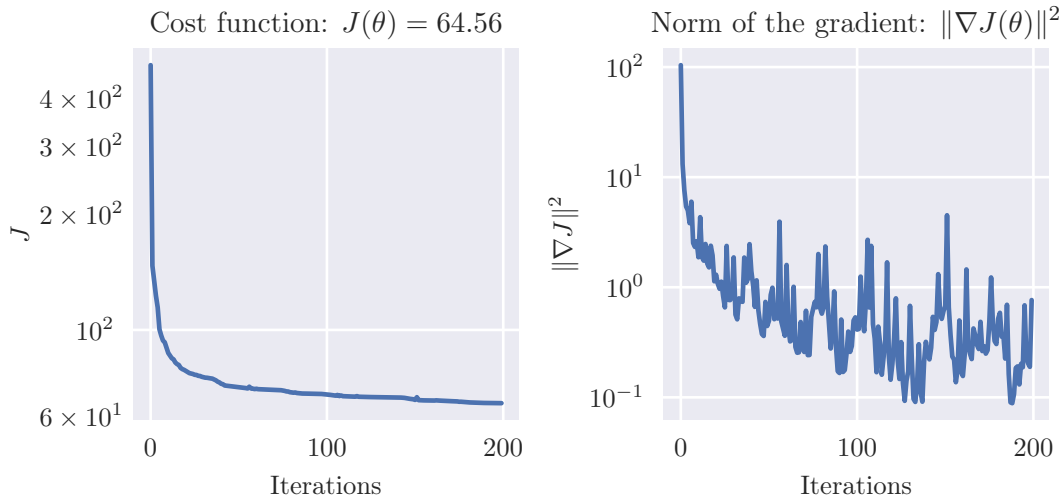
(b) Objective function evolution along the iterations

Figure A.9 – Optimisation on the whole space, $u^b = (1, 0.5)$

Misspecified environmental parameter: $u^b = (1, 1)$



(a) Estimated θ and relative-error



(b) Objective function evolution along the iterations

Figure A.10 – Optimisation on the whole space, $u^b = (1, 1)$

* * *

BIBLIOGRAPHY

- Akaike, H. (1974). A new look at the statistical model identification. *IEEE transactions on automatic control*, 19(6):716–723.
- Andréassian, V., Le Moine, N., Perrin, C., Ramos, M.-H., Oudin, L., Mathevet, T., Lerat, J., and Berthet, L. (2012). All that glitters is not gold: The case of calibrating hydrological models: Invited Commentary. *Hydrological Processes*, 26(14):2206–2210.
- Baudoui, V. (2012). *Optimisation Robuste Multiobjectifs Par Modèles de Substitution*. PhD thesis, Toulouse, ISAE.
- Bect, J., Ginsbourger, D., Li, L., Picheny, V., and Vazquez, E. (2012). Sequential design of computer experiments for the estimation of a probability of failure. *Statistics and Computing*, 22(3):773–793.
- Beirlant, J., Dudewicz, E. J., Györfi, L., and Van der Meulen, E. C. (1997). Nonparametric entropy estimation: An overview. *International Journal of Mathematical and Statistical Sciences*, 6(1):17–39.
- Benameur, H., Chavent, G., and Jaffré, J. (2002). Refinement and Coarsening Indicators for Adaptive Parameterization: Application to the Estimation of Hydraulic Transmissivities. *Inverse Problems*, 18:775.
- Berger, J. O., Liseo, B., and Wolpert, R. L. (1999). Integrated likelihood methods for eliminating nuisance parameters. *Statistical Science*, 14(1):1–28.
- Bernard, L. (2019). *Méthodes Probabilistes Pour l’estimation de Probabilités de Défaillance*. PhD thesis.
- Bertsimas, D., Brown, D. B., and Caramanis, C. (2010). Theory and Applications of Robust Optimization. *arXiv:1010.5445 [cs, math]*.
- Betancourt, M. (2017). A Conceptual Introduction to Hamiltonian Monte Carlo. *arXiv:1701.02434 [stat]*.

- Beyer, H.-G. and Sendhoff, B. (2007). Robust optimization – A comprehensive survey. *Computer Methods in Applied Mechanics and Engineering*, 196(33-34):3190–3218.
- Billingsley, P. (2008). *Probability and Measure*. John Wiley & Sons.
- Bishop, C. M. (2006). *Pattern Recognition and Machine Learning*. springer.
- Bouhlef, M. A. and Martins, J. R. R. A. (2019). Gradient-enhanced kriging for high-dimensional problems. *Engineering with Computers*, 35(1):157–173.
- Boutet, M. (2015). *Estimation Du Frottement Sur Le Fond Pour La Modélisation de La Marée Barotrope*. PhD thesis, Université d’Aix Marseille.
- Briec, W., Kerstens, K., and Jokung, O. (2007). Mean-Variance-Skewness Portfolio Performance Gauging: A General Shortage Function and Dual Approach. *Management Science*, 53(1):135–149.
- Burnham, K. P. and Anderson, D. R. (2004). Multimodel Inference: Understanding AIC and BIC in Model Selection. *Sociological Methods & Research*, 33(2):261–304.
- Cook, L. W. (2018). *Effective Formulations of Optimization Under Uncertainty for Aerospace Design*. Thesis, University of Cambridge.
- Cook, L. W. and Jarrett, J. P. (2018). Horsetail matching: A flexible approach to optimization under uncertainty. *Engineering Optimization*, 50(4):549–567.
- Cook, L. W., Jarrett, J. P., and Willcox, K. E. (2017). Extending Horsetail Matching for Optimization Under Probabilistic, Interval, and Mixed Uncertainties. *AIAA Journal*, 56(2):849–861.
- Couderc, F., Madec, R., Monnier, J., and Vila, J.-P. (2013). *Dassflow-Shallow, Variational Data Assimilation for Shallow-Water Models: Numerical Schemes, User and Developer Guides*. PhD thesis, University of Toulouse, CNRS, IMT, INSA, ANR.
- Das, S. K. and Lardner, R. W. (1991). On the estimation of parameters of hydraulic models by assimilation of periodic tidal data. *Journal of Geophysical Research*, 96(C8):15187.
- Das, S. K. and Lardner, R. W. (1992). Variational parameter estimation for a two-dimensional numerical tidal model. *International Journal for Numerical Methods in Fluids*, 15(3):313–327.
- Debreu, L., Marchesiello, P., Penven, P., and Cambon, G. (2012). Two-way nesting in split-explicit ocean models: Algorithms, implementation and validation. *Ocean Modelling*, 49-50:1–21.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the royal statistical society. Series B (methodological)*, pages 1–38.
- Dobre, S., Bastogne, T., and Richard, A. (2010). Global sensitivity and identifiability implications in systems biology. In *11th IFAC Symposium on Computer Applications in Biotechnology, CAB 2010*, page CDROM, Leuven, Belgium.

- Doucet, A., Godsill, S. J., and Robert, C. P. (2002). Marginal maximum a posteriori estimation using Markov chain Monte Carlo. *Statistics and Computing*, 12(1):77–84.
- Dubourg, V., Sudret, B., and Bourinet, J.-M. (2011). Reliability-based design optimization using kriging surrogates and subset simulation. *Structural and Multidisciplinary Optimization*, 44(5):673–690.
- Dubrule, O. (1983). Cross validation of kriging in a unique neighborhood. *Journal of the International Association for Mathematical Geology*, 15(6):687–699.
- Echard, B., Gayton, N., and Lemaire, M. (2011). AK-MCS: An active learning reliability method combining Kriging and Monte Carlo Simulation. *Structural Safety*, 33(2):145–154.
- Egbert, G. D. and Erofeeva, S. Y. (2002). Efficient Inverse Modeling of Barotropic Ocean Tides. *Journal of Atmospheric and Oceanic Technology*, 19(2):183–204.
- El Amri, M. (2019). *Analyse d’incertitudes et de Robustesse Pour Les Modèles à Entrées et Sorties Fonctionnelles*. Thesis, Grenoble Alpes.
- El Amri, M. R., Helbert, C., Lepreux, O., Zuniga, M. M., Prieur, C., and Sinoquet, D. (2019). Data-driven stochastic inversion via functional quantization. *Statistics and Computing*.
- Friel, N. and Wyse, J. (2011). Estimating the evidence – a review. *arXiv:1111.1957 [stat]*.
- Gilbert, J. C. and Lemaréchal, C. (1989). Some numerical experiments with variable-storage quasi-Newton algorithms. *Mathematical programming*, 45(1-3):407–435.
- Gilquin, L. (2016). *Échantillonnages Monte Carlo et quasi-Monte Carlo pour l’estimation des indices de Sobol’ : application à un modèle transport-urbanisme*. PhD thesis, Université Grenoble Alpes.
- Gilquin, L., Arnaud, E., Prieur, C., and Janon, A. (2019). Making best use of permutations to compute sensitivity indices with replicated orthogonal arrays. *Reliability Engineering and System Safety*, 187:28–39.
- Ginsbourger, D., Baccou, J., Chevalier, C., Perales, F., Garland, N., and Monerie, Y. (2014). Bayesian Adaptive Reconstruction of Profile Optima and Optimizers. *SIAM/ASA Journal on Uncertainty Quantification*, 2(1):490–510.
- Ginsbourger, D., Dupuy, D., Badea, A., Carraro, L., and Roustant, O. (2009). A note on the choice and the estimation of Kriging models for the analysis of deterministic computer experiments. *Applied Stochastic Models in Business and Industry*, 25(2):115–131.
- Ginsbourger, D., Le Riche, R., and Carraro, L. (2010). Kriging is well-suited to parallelize optimization. In Tenne, Y. and Goh, C.-K., editors, *Computational Intelligence in Expensive Optimization Problems*, Springer Series in Evolutionary Learning and Optimization, pages 131–162. springer.

- Hanson, K. (2001). Markov Chain Monte Carlo posterior sampling with the Hamiltonian Method. Technical Report LA-UR-01-1016, Los Alamos National Lab., NM (US).
- Hascoet, L. and Pascual, V. (2013). The Tapenade automatic differentiation tool: Principles, model, and specification. *ACM Transactions on Mathematical Software*, 39(3):1–43.
- Hennig, P. and Schuler, C. J. (2011). Entropy Search for Information-Efficient Global Optimization.
- Higdon, D., Kennedy, M., Cavendish, J. C., Cafeo, J. A., and Ryne, R. D. (2004). Combining field data and computer simulations for calibration and prediction. *SIAM Journal on Scientific Computing*, 26(2):448–466.
- Honnorat, M., Monnier, J., Rivière, N., Huot, É., and Le Dimet, F.-X. (2010). Identification of equivalent topography in an open channel flow using Lagrangian data assimilation. *Computing and Visualization in Science*, 13(3):111–119.
- Huber, P. J. (2011). Robust statistics. In *International Encyclopedia of Statistical Science*, pages 1248–1251. Springer.
- Huyse, L. and Bushnell, D. M. (2001). Free-form airfoil shape optimization under uncertainty using maximum expected value and second-order second-moment strategies.
- Huyse, L., Padula, S. L., Lewis, R. M., and Li, W. (2002). Probabilistic Approach to Free-Form Airfoil Shape Optimization Under Uncertainty. *AIAA Journal*, 40(9):1764–1772.
- Iooss, B. (2011). Revue sur l’analyse de sensibilité globale de modèles numériques. *Journal de la Société Française de Statistique*, 152(1):1–23.
- Janon, A. (2012). *Analyse de Sensibilité et Réduction de Dimension. Application à l’océanographie* *Uncertainties Assessment in Global Sensitivity Indices Estimation from Metamodels Certified Reduced Basis Solutions of Viscous Burgers Equation Parametrized by Initial and Boundary Values*. These de doctorat, Grenoble.
- Janusevskis, J. and Le Riche, R. (2010). Simultaneous kriging-based sampling for optimization and uncertainty propagation. Technical report.
- Jones, D. R., Schonlau, M., and Welch, W. J. (1998). Efficient global optimization of expensive black-box functions. *Journal of Global optimization*, 13(4):455–492.
- Juditsky, A., Nemirovski, A. S., Lan, G., and Shapiro, A. (2009). Stochastic Approximation Approach to Stochastic Programming. In *ISMP 2009 - 20th International Symposium of Mathematical Programming*.
- Kalbfleisch, J. G. (1985). *Probability and Statistical Inference*. Springer Texts in Statistics. Springer New York, New York, NY.
- Kass, R. E. and Raftery, A. E. (1995). Bayes factors. *Journal of the american statistical association*, 90(430):773–795.

- Kennedy, M. C. and O'Hagan, A. (2001). Bayesian calibration of computer models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 63(3):425–464.
- Kent, J. T. (1982). Robust properties of likelihood ratio tests. *Biometrika*, 69(1):19–27.
- Kiureghian, A. D. and Ditlevsen, O. (2009). Aleatory or epistemic? Does it matter? *Structural Safety*, 31(2):105–112.
- Kreitmair, M. J., Draper, S., Borthwick, A. G. L., and van den Bremer, T. S. (2019). The effect of uncertain bottom friction on estimates of tidal current power. *Royal Society Open Science*, 6(1).
- Krige, D. G. (1951). A statistical approach to some basic mine valuation problems on the Witwatersrand. *Journal of the Southern African Institute of Mining and Metallurgy*, 52(6):119–139.
- Kuczera, G., Renard, B., Thyer, M., and Kavetski, D. (2010). There are no hydrological monsters, just models and observations with large uncertainties! *Hydrological Sciences Journal*, 55(6):980–991.
- Kullback, S. and Leibler, R. A. (1951). On Information and Sufficiency. *The Annals of Mathematical Statistics*, 22(1):79–86.
- Lai, K. K., Yu, L., and Wang, S. (2006). Mean-Variance-Skewness-Kurtosis-based Portfolio Optimization. In *First International Multi-Symposiums on Computer and Computational Sciences (IMSCCS'06)*, volume 2, pages 292–297.
- Laurent, L., Le Riche, R., Soulier, B., and Boucard, P.-A. (2019). An Overview of Gradient-Enhanced Metamodels with Applications. *Archives of Computational Methods in Engineering*, 26(1):61–106.
- Le Bars, Y., Lyard, F., Jeandel, C., and Dardengo, L. (2010). The AMANDES tidal model for the Amazon estuary and shelf. *Ocean Modelling*, 31(3):132–149.
- Le Riche, R. (2014). Introduction to kriging.
- Lehman, J. S., Santner, T. J., and Notz, W. I. (2004). Designing computer experiments to determine robust control variables. *Statistica Sinica*, pages 571–590.
- Lehmann, E. L. and Casella, G. (2006). *Theory of Point Estimation*. Springer Science & Business Media.
- Lelièvre, N., Beaurepaire, P., Mattrand, C., Gayton, N., and Otsmane, A. (2016). On the consideration of uncertainty in design: Optimization-reliability-robustness. *Structural and Multidisciplinary Optimization*, 54(6):1423–1437.
- Li, S. Z. (2009). *Markov Random Field Modeling in Image Analysis*. Springer Science & Business Media.

- MacQueen, J. (1967). Some methods for classification and analysis of multivariate observations. In *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Statistics*. The Regents of the University of California.
- Mahalanobis, P. C. (1936). On the generalized distance in statistics. National Institute of Science of India.
- Marler, R. T. and Arora, J. S. (2010). The weighted sum method for multi-objective optimization: New insights. *Structural and Multidisciplinary Optimization*, 41(6):853–862.
- Marmin, S., Chevalier, C., and Ginsbourger, D. (2015). Differentiating the Multipoint Expected Improvement for Optimal Batch Design. In Pardalos, P., Pavone, M., Farinella, G. M., and Cutello, V., editors, *Machine Learning, Optimization, and Big Data*, volume 9432, pages 37–48. Springer International Publishing, Cham.
- Matheron, G. (1962). *Traité de Géostatistique Appliquée. 1 (1962)*, volume 1. Editions Technip.
- McWilliams, J. C. (2007). Irreducible imprecision in atmospheric and oceanic simulations. *Proceedings of the National Academy of Sciences*, 104(21):8709–8713.
- Miranda, J., Kumar, D., and Lacor, C. (2016). Adjoint-based Robust Optimization using Polynomial Chaos Expansions. In *Proceedings of the VII European Congress on Computational Methods in Applied Sciences and Engineering (ECCOMAS Congress 2016)*, pages 8351–8364, Crete Island, Greece. Institute of Structural Analysis and Antiseismic Research School of Civil Engineering National Technical University of Athens (NTUA) Greece.
- Moćkus, J. (1974). On bayesian methods for seeking the extremum. In *Optimization Techniques IFIP Technical Conference Novosibirsk, July 1–7, 1974*, Lecture Notes in Computer Science, pages 400–404. Springer, Berlin, Heidelberg.
- Moustapha, M., Sudret, B., Bourinet, J.-M., and Guillaume, B. (2016). Quantile-based optimization under uncertainties using adaptive Kriging surrogate models. *Structural and Multidisciplinary Optimization*, 54(6):1403–1421.
- Nielsen, F. (2016). Hierarchical Clustering. pages 195–211.
- Ogryczak, W. and Ruszczyński, A. (1997). On stochastic dominance and mean-semideviation models.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M., and Duchesnay, É. (2011). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*, 12:2825–2830.
- Picheny, V. and Ginsbourger, D. (2014). Noisy kriging-based optimization methods: A unified implementation within the DiceOptim package. *Computational Statistics & Data Analysis*, 71:1035–1053.

- Rao, V., Sandu, A., Ng, M., and Nino-Ruiz, E. (2015). Robust data assimilation using L_1 and Huber norms. *SciRate*.
- Rasmussen, C. E. and Williams, C. K. I. (2006). *Gaussian Processes for Machine Learning*. Adaptive Computation and Machine Learning. MIT Press, Cambridge, Mass.
- Razaaly, N. (2019). *Rare Event Estimation and Robust Optimization Methods with Application to ORC Turbine Cascade*. These de doctorat, Université Paris-Saclay (ComUE).
- Razaaly, N. and Congedo, P. M. (2020). Extension of AK-MCS for the efficient computation of very small failure probabilities. *Reliability Engineering & System Safety*, 203:107084.
- Razaaly, N., Persico, G., Gori, G., and Congedo, P. M. (2020). Quantile-based robust optimization of a supersonic nozzle for organic rankine cycle turbines. *Applied Mathematical Modelling*, 82:802–824.
- Reid, N. (2013). Aspects of likelihood inference. *Bernoulli*, 19(4):1404–1418.
- Ribaud, M. (2018). *Krigeage Pour La Conception de Turbomachines : Grande Dimension et Optimisation Multi-Objectif Robuste*. Thesis, Lyon.
- Ribaud, M., Blanchet-Scalliet, C., Gillot, F., and Helbert, C. (2019). Robustness kriging-based optimization.
- Sacks, J., Schiller, S. B., and Welch, W. J. (1989). Designs for Computer Experiments. *Technometrics*, 31(1):41–47.
- Savage, L. J. (1951). The Theory of Statistical Decision. *Journal of the American Statistical Association*, 46(253):55–67.
- Schöbi, R., Sudret, B., and Marelli, S. (2017). Rare Event Estimation Using Polynomial-Chaos Kriging. *ASCE-ASME Journal of Risk and Uncertainty in Engineering Systems, Part A: Civil Engineering*, 3(2).
- Schwarz, G. (1978). Estimating the Dimension of a Model. *Annals of Statistics*, 6(2):461–464.
- Scott, D. W. (1979). On Optimal and Data-Based Histograms. *Biometrika*, 66(3):605.
- Seshadri, P., Constantine, P., Iaccarino, G., and Parks, G. (2014). A density-matching approach for optimization under uncertainty. *arXiv:1409.7089 [math, stat]*.
- Shah, A., Wilson, A. G., and Ghahramani, Z. (2014). Student-t Processes as Alternatives to Gaussian Processes. *arXiv:1402.4306 [cs, stat]*.
- Sinha, B. and Pingree, R. D. (1997). The principal lunar semidiurnal tide and its harmonics: Baseline solutions for M2 and M4 constituents on the North-West European Continental Shelf. *Continental Shelf Research*, 17(11):1321–1365.

- Sobol, I. M. (1993). Sensitivity analysis for non-linear mathematical models. *Mathematical modelling and computational experiment*, 1:407–414.
- Sobol, I. M. (2001). Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Mathematics and computers in simulation*, 55(1-3):271–280.
- Sudret, B. (2015). Polynomial chaos expansions and stochastic finite element methods. In Kok-Kwang Phoon, J. C., editor, *Risk and Reliability in Geotechnical Engineering*, pages 265–300. CRC Press.
- Tarantola, A. (2005). *Inverse Problem Theory and Methods for Model Parameter Estimation*. SIAM, Society for Industrial and Applied Mathematics, Philadelphia, Pa.
- Tibshirani, R. (2011). Regression shrinkage and selection via the lasso: A retrospective. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 73(3):273–282.
- Tikhonov, A. and Arsenin, V. (1977). *Solutions of Ill-Posed Problems*, volume 14.
- Trappler, V., Arnaud, É., Vidard, A., and Debreu, L. (2020). Robust calibration of numerical models based on relative regret. *Journal of Computational Physics*, page 109952.
- Vapnik, V. (1992). Principles of risk minimization for learning theory. In *Advances in Neural Information Processing Systems*, pages 831–838.
- Villemonteix, J., Vazquez, E., and Walter, E. (2006). An informational approach to the global optimization of expensive-to-evaluate functions. *arXiv:cs/0611143*.
- Vorobyev, O. and Vorobyev, A. (2003). On the New Notion of the Set-Expectation for a Random Set of Events. *University Library of Munich, Germany, MPRA Paper*.
- Wald, A. (1945). Statistical Decision Functions Which Minimize the Maximum Risk. *Annals of Mathematics*, 46(2):265–280.
- Walker, W. E., Harremoës, P., Rotmans, J., van der Sluijs, J. P., van Asselt, M. B., Janssen, P., and Kreyer von Krauss, M. P. (2003). Defining uncertainty: A conceptual basis for uncertainty management in model-based decision support. *Integrated assessment*, 4(1):5–17.
- White, H. (1982). Maximum Likelihood Estimation of Misspecified Models. *Econometrica*, 50(1):1–25.
- Wiener, N. (1938). The Homogeneous Chaos. *American Journal of Mathematics*, 60(4):897–936.
- Wilks, S. S. (1938). The Large-Sample Distribution of the Likelihood Ratio for Testing Composite Hypotheses. *Annals of Mathematical Statistics*, 9(1):60–62.
- Wilson, E. B. (1927). Probable inference, the law of succession, and statistical inference. *Journal of the American Statistical Association*, 22(158):209–212.

- Xiu, D. and Karniadakis, G. (2002). The Wiener–Askey Polynomial Chaos for Stochastic Differential Equations. *SIAM Journal on Scientific Computing*, 24(2):619–644.
- Yizong Cheng (Aug./1995). Mean shift, mode seeking, and clustering. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 17(8):790–799.
- Zahm, O., Cui, T., Law, K., Spantini, A., and Marzouk, Y. (2018). Certified dimension reduction in nonlinear Bayesian inverse problems.
- Zanna, L. (2011). Ocean Model Uncertainty in Climate Prediction. page 8.

Abstract

To understand and to be able to forecast natural phenomena is increasingly important nowadays, as those predictions are often the basis of many decisions, whether economical or ecological. In order to do so, mathematical models are introduced to represent the reality at a specific scale, and are then implemented numerically. However in this process of modelling, many complex phenomena occurring at a smaller scale than the one studied have to be simplified and quantified. This often leads to the introduction of additional parameters, which then need to be properly estimated. Classical methods of estimation usually involve an objective function, that measures the distance between the simulations and some observations, which is then optimised. Such an optimisation require many runs of the numerical model and possibly the computation of its gradient, thus can be expensive to evaluate computational-wise.

However, some other uncertainties can also be present, which represent some uncontrollable and external factors that affect the modelling. Those variables will be qualified as *environmental*. By modelling them with a random variable, the objective function is then a random variable as well, that we wish to minimise in some sense. Omitting the random nature of the environmental variable can lead to localised optimisation, and thus a value of the parameters that is optimal only for the fixed nominal value. To overcome this, the minimisation of the expected value of the objective function is often considered in the field of optimisation under uncertainty for instance.

In this thesis, we focus instead on the notion of regret, that measures the deviation of the objective function from its optimal value given a realisation of the environmental variable. This regret (either additive or relative) translates a notion of robustness through its probability of exceeding a specified threshold. So, by either controlling the threshold or the probability, we can define a family of estimators based on this regret.

The regret can quickly become expensive to evaluate since it requires an optimisation of the objective for every realisation of the environmental variable. We then propose to use Gaussian Processes (GP) in order to reduce the computational burden of this evaluation. In addition to that, we propose a few adaptive methods in order to improve the estimation : the next points to evaluate are chosen sequentially according to a specific criterion, in a Stepwise Uncertainty Reduction (SUR) strategy.

Finally, we will apply some of the methods introduced in this thesis on an academic problem of parameter estimation. We will study the calibration of the bottom friction of a model of the Atlantic ocean near the French coasts, while introducing some uncertainties in the forcing of the tide, and get a robust estimation of this friction parameter in a twin experiment setting.

Keywords : Optimisation under uncertainties ; Robust calibration ; Gaussian Processes ; Ocean modelling ; Regret

* * *

Résumé

De nombreux phénomènes physiques sont modélisés afin d'en mieux connaître les comportements ou de pouvoir les prévoir. Cependant pour représenter la réalité, de nombreux processus doivent être simplifiés, car ils sont souvent trop complexes, ou apparaissent à une échelle bien inférieure à celle de l'étude du phénomène. Au lieu de complètement les omettre, les effets de ces processus sont souvent retranscrits dans les modèles à l'aide de paramétrisations, c'est-à-dire en introduisant des termes les quantifiant, et qui doivent être ensuite estimées. Les méthodes classiques d'estimation se basent sur la définition d'une fonction objectif qui mesure l'écart entre le modèle numérique et la réalité, qui est ensuite optimisée.

Cependant, au delà de l'incertitude sur la valeur du paramètre à estimer, un autre type d'incertitude peut aussi être présent. Cela permet de représenter la variabilité intrinsèque de certains processus externes, qui vont avoir un effet sur la modélisation. Ces variables vont être qualifiées d'*environnementales*. En les modélisant à l'aide d'une variable aléatoire, la fonction objectif devient à son tour une variable aléatoire, que l'on va chercher à minimiser dans un certain sens. Si on omet ce caractère aléatoire, on peut se retrouver avec un paramètre optimal uniquement pour la valeur nominale du paramètre environnemental, et le modèle peut s'éloigner de la réalité pour d'autres réalisations. Ce problème d'optimisation sous incertitudes est souvent abordé en optimisant les premiers moments de la variable aléatoire, l'espérance en particulier.

Dans cette thèse, nous nous intéressons plutôt à la notion de regret, qui mesure l'écart entre la fonction objectif et la valeur optimale qu'elle peut atteindre, pour la réalisation de la variable environnementale donnée. Cette idée de regret (additif ou bien relatif) nous permet de proposer une notion de robustesse à travers l'étude de sa probabilité de dépasser un certain seuil, ou inversement à travers le calcul de ses quantiles. À l'aide de ce seuil, ou de l'ordre du quantile choisi, on peut donc définir une famille d'estimateurs basés sur le regret.

Néanmoins, le calcul du regret, et donc des quantités dérivées peut vite devenir très coûteux, car il nécessite une optimisation par rapport au paramètre de contrôle. Nous proposons donc d'utiliser des processus Gaussiens (GP) afin de construire un modèle de substitution, et donc de réduire cette contrainte en pratique. Nous proposons aussi des méthodes itératives basées notamment sur la stratégie SUR (Stepwise Uncertainty Reduction, Réduction d'incertitudes séquentielle) : le point à évaluer ensuite est choisi selon un critère permettant d'améliorer au mieux des quantités associées au regret-relatif.

Enfin, nous appliquons les outils présentés dans cette thèse à un problème académique d'estimation de paramètre. Nous étudions ainsi la calibration sous incertitudes du paramètre de friction de fond d'un modèle océanique, représentant la façade atlantique des côtes françaises, ainsi que la Manche dans un cadre d'expériences jumelles.

Mots-Clés : Optimisation sous incertitudes ; Calibration robuste ; Processus Gaussiens ; Modélisation de l'océan ; Regret