

DISSERTATION

SURROGATE MODELING FOR EFFICIENT ANALYSIS AND DESIGN OF ENGINEERING  
SYSTEMS

Submitted by

Min Li

Department of Civil and Environmental Engineering

In partial fulfillment of the requirements

For the Degree of Doctor of Philosophy

Colorado State University

Fort Collins, Colorado

Fall 2021

Doctoral Committee:

Advisor: Gaofeng Jia

Bruce Ellingwood

John van de Lindt

Xinfeng Gao

Copyright by Min Li 2021

All Rights Reserved

## ABSTRACT

### SURROGATE MODELING FOR EFFICIENT ANALYSIS AND DESIGN OF ENGINEERING SYSTEMS

Surrogate models, trained using a data-driven approach, have been extensively used to approximate the input/output relationship for expensive high-fidelity models (e.g., large-scale physical experiments and high-resolution computationally expensive numerical simulations). The computational efficiency of surrogate models is greatly increased compared with the high-fidelity models. Once trained, the original high-fidelity models can be replaced by the surrogate models to facilitate efficient subsequent analysis and design of engineering systems. The quality of surrogate based analysis and design of engineering systems relies largely on the prediction accuracy of the constructed surrogate model. To ensure the prediction accuracy, the training data should be adequate in terms of the size of the training data and their sampling. Unfortunately, constrained by limited computational budgets, typically it is challenging to obtain a lot of training data by running high-fidelity models. Furthermore, significant challenge arises in obtaining sufficient training data for problems with high-dimensional model inputs due to the well-known curse of dimensionality.

In order to build surrogate models with high prediction accuracy and generalization performance while using as less computational resources as possible, this dissertation proposes several advanced strategies and examines their performances within several practical engineering applications. The fundamental idea of the proposed strategies is to embed extra knowledge about the high-fidelity models in the surrogate model by enriching the training data (e.g., leverage additional

low-fidelity data, or censored/bounded data) and enhancing model assumption (e.g., explicitly incorporate prior knowledge about the physics of the problem, or explore low-dimensional latent structures/features), which reduces the required size of high-fidelity training data and meanwhile effectively boosts the prediction accuracy of the established surrogate model. Among different surrogate models, Gaussian process models have been gaining popularity due to its flexibility in modeling complex functions and ability to provide closed-form predictive distributions. Therefore, the strategies are developed in the context of Gaussian process model, but the ideas are expected to be applicable to other types of surrogate models. In particular, this dissertation (i) develops a physics-constrained Gaussian process model to efficiently incorporate our prior knowledge about physical constraints/characteristics of the input/output relationship by designing specific kernels, (ii) proposes a general multi-fidelity Gaussian process model capable of integrating training data with different level of accuracy (i.e., both high-fidelity data and low-fidelity data) and completeness (i.e., both accurate data and censored data), and (iii) develops an efficient surrogate modeling approach for problems with high-dimensional binary model inputs by integrating dimension reduction technique and Gaussian process model, and investigates its application in design optimization problems. The excellent performance of the proposed strategies are then validated through analysis and design of several different engineering systems, including (i) calculating hydrodynamic characteristics of wave energy converters (WECs) in an array, (ii) predicting the deformation capacity of reinforced concrete columns under cyclic loading, and (iii) optimizing topology of periodic structures.

## ACKNOWLEDGEMENTS

I would first like to convey my deepest gratitude to my advisor, Prof. Gaofeng Jia, for giving me the opportunity to work on so many interesting research topics of surrogate modeling, machine learning, and uncertainty quantification. His enthusiasm in conducting research has significantly motivated me. Without his patience, unwavering support, and valuable suggestions, I could not imagine the completion of my dissertation. I am also so grateful for his encouragements and advice when I felt uncertain and anxious about my future career and life. It has been truly lucky to work with such an admirable researcher.

I would also like to express my appreciation to my committee members, Prof. Bruce Ellingwood, Prof. John van de Lindt, and Prof. Xinfeng Gao, for their invaluable comments, insights, and suggestions to this dissertation. I am also thankful to Prof. Bruce Ellingwood for his inspiring lectures, which have laid the foundation for my research during these years.

In addition, I would like to recognize the invaluable assistance that my friends at Colorado State University, Zhenqiang Wang, Abdelrahman Abdallah, Chao Jiang, Jiate Li, Yue Dong, et al., have provided during my study.

Last but not least, I would like to express my sincere gratitude to my parents for their consistent support, faith, and encouragement throughout my life. Also, I want to convey my special thanks to Qiling Zou for his love and support. This journey would not have been possible without their support.

## DEDICATION

*To my parents.*

## TABLE OF CONTENTS

|                                                                                                         |    |
|---------------------------------------------------------------------------------------------------------|----|
| ABSTRACT . . . . .                                                                                      | ii |
| ACKNOWLEDGEMENTS . . . . .                                                                              | iv |
| DEDICATION . . . . .                                                                                    | v  |
| LIST OF TABLES . . . . .                                                                                | ix |
| LIST OF FIGURES . . . . .                                                                               | x  |
| <br>Chapter 1                                                                                           |    |
| Introduction . . . . .                                                                                  | 1  |
| 1.1 Background . . . . .                                                                                | 1  |
| 1.2 Literature Review . . . . .                                                                         | 4  |
| 1.2.1 Physics-constrained surrogate model . . . . .                                                     | 4  |
| 1.2.2 Multi-fidelity surrogate model . . . . .                                                          | 7  |
| 1.2.3 Surrogate model considering data censoring . . . . .                                              | 9  |
| 1.2.4 Dimension reduction assisted surrogate model . . . . .                                            | 11 |
| 1.2.5 Surrogate model based optimization . . . . .                                                      | 13 |
| 1.3 Research Gaps and Motivation . . . . .                                                              | 14 |
| 1.3.1 Physics-constrained Gaussian process model . . . . .                                              | 15 |
| 1.3.2 Multi-fidelity Gaussian process model considering censored data . . . . .                         | 16 |
| 1.3.3 Dimension reduction assisted Gaussian process model and its application in optimization . . . . . | 16 |
| 1.4 Summary of contributions and objectives . . . . .                                                   | 17 |
| 1.5 Outline of Dissertation . . . . .                                                                   | 19 |
| <br>Chapter 2                                                                                           |    |
| Review of Gaussian Process Model . . . . .                                                              | 21 |
| 2.1 Introduction . . . . .                                                                              | 21 |
| 2.2 Gaussian Process Model . . . . .                                                                    | 21 |
| 2.3 Model Parameter Calibration . . . . .                                                               | 24 |
| 2.4 Prediction . . . . .                                                                                | 25 |
| 2.5 Design of Experiment . . . . .                                                                      | 26 |
| 2.6 Multidimensional outputs . . . . .                                                                  | 28 |
| 2.7 Derivatives of model outputs . . . . .                                                              | 29 |
| <br>Chapter 3                                                                                           |    |
| Physics-Constrained Gaussian Process Model Through Kernel Design . . . . .                              | 31 |
| 3.1 Introduction . . . . .                                                                              | 31 |
| 3.2 Physical Constraints: Invariance, Symmetry, and Additivity . . . . .                                | 32 |
| 3.2.1 Invariance and symmetry . . . . .                                                                 | 32 |
| 3.2.2 Additivity . . . . .                                                                              | 33 |
| 3.3 Physics-constrained Gaussian Process Model through Kernel Design . . . . .                          | 35 |
| 3.3.1 Basics of kernel function . . . . .                                                               | 36 |
| 3.3.2 Invariant kernel . . . . .                                                                        | 37 |
| 3.3.3 Additive kernel . . . . .                                                                         | 38 |
| 3.3.4 Combined invariant and additive kernel . . . . .                                                  | 40 |

|           |                                                                                                                               |     |
|-----------|-------------------------------------------------------------------------------------------------------------------------------|-----|
| 3.4       | Illustrative Example: Hydrodynamic Characteristics of WECs in an Array                                                        | 41  |
| 3.4.1     | Motivation . . . . .                                                                                                          | 41  |
| 3.4.2     | Problem formulation: hydrodynamic interaction between WECs . . . . .                                                          | 45  |
| 3.4.3     | Physics-constrained Gaussian process model for predicting hydrodynamic characteristics . . . . .                              | 50  |
| 3.4.4     | Implementation details . . . . .                                                                                              | 58  |
| 3.4.5     | Prediction accuracy . . . . .                                                                                                 | 60  |
| 3.4.6     | Performance of physics-constrained kernel . . . . .                                                                           | 63  |
| 3.4.7     | Computational efficiency . . . . .                                                                                            | 70  |
| 3.5       | Conclusions . . . . .                                                                                                         | 71  |
| Chapter 4 | Multi-fidelity Gaussian Process Model Integrating Low- and High-Fidelity Data Considering Censoring . . . . .                 | 73  |
| 4.1       | Introduction . . . . .                                                                                                        | 73  |
| 4.2       | Multi-Fidelity Gaussian Process Prediction Model . . . . .                                                                    | 75  |
| 4.2.1     | Model formulation . . . . .                                                                                                   | 75  |
| 4.2.2     | Posterior predictions at new inputs . . . . .                                                                                 | 77  |
| 4.3       | Posterior Distributions of Model Parameters . . . . .                                                                         | 80  |
| 4.3.1     | Prior distributions of model parameters . . . . .                                                                             | 80  |
| 4.3.2     | Likelihood functions of model parameters . . . . .                                                                            | 82  |
| 4.4       | Sampling from Posterior Distributions of Model Parameters . . . . .                                                           | 84  |
| 4.4.1     | Posterior distributions of model parameters . . . . .                                                                         | 84  |
| 4.4.2     | Sampling from posterior distributions of model parameters . . . . .                                                           | 87  |
| 4.5       | Illustrative Example: Probabilistic Deformation Capacity of RC Columns . . . . .                                              | 90  |
| 4.5.1     | High-fidelity database with censored data . . . . .                                                                           | 91  |
| 4.5.2     | Low-fidelity database . . . . .                                                                                               | 92  |
| 4.5.3     | Implementation details . . . . .                                                                                              | 94  |
| 4.5.4     | Posterior statistics of model parameters . . . . .                                                                            | 96  |
| 4.5.5     | Posterior predictive distributions for the outputs . . . . .                                                                  | 100 |
| 4.6       | Conclusions . . . . .                                                                                                         | 103 |
| Chapter 5 | Dimension Reduction Assisted Gaussian Process Model for High-Dimensional Inputs and Its Application in Optimization . . . . . | 105 |
| 5.1       | Introduction . . . . .                                                                                                        | 105 |
| 5.2       | Problem Formulation: Optimization for Problems with High-dimensional Inputs . . . . .                                         | 106 |
| 5.3       | Low-dimensional Representation of High-dimensional Inputs . . . . .                                                           | 107 |
| 5.3.1     | Set of reference inputs . . . . .                                                                                             | 107 |
| 5.3.2     | Dimension reduction by logistic PCA . . . . .                                                                                 | 108 |
| 5.4       | Optimization in Low-dimensional Continuous Design Space . . . . .                                                             | 110 |
| 5.5       | Surrogate based Optimization with Adaptive Design of Experiment . . . . .                                                     | 111 |
| 5.6       | Illustrative Example: Topology Optimization of Periodic Structures . . . . .                                                  | 113 |
| 5.6.1     | Motivation . . . . .                                                                                                          | 113 |
| 5.6.2     | Frequency bandgap of periodic structures . . . . .                                                                            | 116 |
| 5.6.3     | Topology optimization of unit cell of periodic structures . . . . .                                                           | 118 |

|              |                                                                                                                                                      |     |
|--------------|------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 5.6.4        | DRESTO: overview . . . . .                                                                                                                           | 119 |
| 5.6.5        | DRESTO: implementation details . . . . .                                                                                                             | 121 |
| 5.6.6        | Optimal topologies . . . . .                                                                                                                         | 124 |
| 5.6.7        | Computational efficiency . . . . .                                                                                                                   | 127 |
| 5.7          | Conclusions . . . . .                                                                                                                                | 131 |
| Chapter 6    | Summary of Dissertation and Future Studies . . . . .                                                                                                 | 133 |
| 6.1          | Conclusions . . . . .                                                                                                                                | 133 |
| 6.2          | Future directions . . . . .                                                                                                                          | 137 |
| Bibliography | . . . . .                                                                                                                                            | 141 |
| Appendix A   | Derivation of Marginal Posterior Distribution for $\theta$ in Multi-fidelity Gaussian Process Model . . . . .                                        | 169 |
| Appendix B   | High-fidelity Database for Developing Multi-fidelity Gaussian Process Model to Predict Deformation Capacity of Reinforced Concrete Columns . . . . . | 174 |

## LIST OF TABLES

|     |                                                                                                                           |     |
|-----|---------------------------------------------------------------------------------------------------------------------------|-----|
| 3.1 | Prediction accuracy metrics of physics-constrained Gaussian process model for $\omega = 1.4$ rad/s . . . . .              | 62  |
| 3.2 | Prediction accuracy metrics from physics-constrained Gaussian process model and standard Gaussian process model . . . . . | 70  |
| 4.1 | MAP estimates for $\theta_l$ and $\theta_\delta$ . . . . .                                                                | 97  |
| 4.2 | Posterior statistics of $\beta_l$ . . . . .                                                                               | 98  |
| 4.3 | Posterior statistics of $\beta_\rho$ . . . . .                                                                            | 98  |
| 4.4 | Posterior statistics of $\beta_\delta$ . . . . .                                                                          | 99  |
| 4.5 | Comparison of the mean errors . . . . .                                                                                   | 101 |
| B.1 | Accurate high-fidelity data . . . . .                                                                                     | 175 |
| B.2 | Censored high-fidelity data. . . . .                                                                                      | 176 |
| B.3 | Description and unit of notations . . . . .                                                                               | 177 |

## LIST OF FIGURES

|      |                                                                                                                                                                                                                         |     |
|------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 2.1  | Illustrative example of Gaussian process model (GP in the figure stands for Gaussian process model). . . . .                                                                                                            | 26  |
| 3.1  | Permutation invariance property of interatomic potential $E$ . . . . .                                                                                                                                                  | 33  |
| 3.2  | Illustration of additivity property of a three-body example. . . . .                                                                                                                                                    | 35  |
| 3.3  | Comparison between true response and predictions by Gaussian process models using standard kernel and invariant kernel (black dots correspond to the training samples). . .                                             | 38  |
| 3.4  | Additive structure in kernel and Gaussian process model (adapted from Duvenaud (2014)). . . . .                                                                                                                         | 40  |
| 3.5  | (a) Wave energy converter (WEC) (adapted from (Prakash et al. 2016)), (b) wave farm (adapted from (Drew et al. 2009)), and (c) Cartesian representation of an array of $N$ WECs. . . . .                                | 46  |
| 3.6  | Illustration of hydrodynamic interactions between WECs and concept of MS approach. .                                                                                                                                    | 49  |
| 3.7  | Illustration of symmetry property of hydrodynamic characteristics. . . . .                                                                                                                                              | 53  |
| 3.8  | Illustration of decomposition of many-body systems (i.e., array of $N$ WECs). . . . .                                                                                                                                   | 54  |
| 3.9  | Flowchart of the proposed physics-constrained Gaussian process model for prediction of hydrodynamic characteristics. . . . .                                                                                            | 58  |
| 3.10 | Comparison of the hydrodynamic characteristics calculated by numerical model (black lines) and predicted by physics-constrained Gaussian process model (red dots). . . . .                                              | 64  |
| 3.11 | Covariance matrix calculated by physics-constrained kernel and standard kernel. . . .                                                                                                                                   | 66  |
| 3.12 | Permutation invariance of $F^{(1)}$ from true function, physics-constrained Gaussian process model and standard Gaussian process model (GP in the figure stands for Gaussian process model). . . . .                    | 68  |
| 3.13 | $x$ -axis symmetry of $F^{(1)}$ from true function, physics-constrained Gaussian process model and standard Gaussian process model (GP in the figure stands for Gaussian process model). . . . .                        | 68  |
| 4.1  | Illustrative example of multi-fidelity Gaussian process model (MFGP in the figure stands for multi-fidelity Gaussian process model). . . . .                                                                            | 78  |
| 4.2  | Prior and posterior distributions for $\sigma_l^2$ , $\sigma_\delta^2$ and $\sigma_\epsilon^2$ . . . . .                                                                                                                | 99  |
| 4.3  | Comparison between the observations and the mean predictions by the multi-fidelity Gaussian process model established using (a) “complete data”, (b) “accurate data”. . .                                               | 102 |
| 4.4  | Comparison between the observations and the predictive distributions (showing $\mu \pm \sigma$ ) by the multi-fidelity Gaussian process model established using (a) “complete data”, (b) “accurate data”. . . . .       | 102 |
| 4.5  | Probabilities that predictions by the multi-fidelity Gaussian process models will exceed the observations for the censored data points. . . . .                                                                         | 103 |
| 5.1  | Periodic structures for (a) sound and vibration reduction (adapted from (Bilal et al. 2018)), (b) seismic isolation (adapted from (Kim and Das 2012)), and (c) wave guide (adapted from (Vasseur et al. 2011)). . . . . | 114 |

|      |                                                                                                                                                                                                                                                               |     |
|------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 5.2  | Illustration of the (a) two dimensional periodic structure, (b) the unit cell, (c) the first irreducible Brillouin zone, (d) frequency bandgap. . . . .                                                                                                       | 117 |
| 5.3  | Flowchart for the proposed DRESTO for topology optimization of periodic structures. . . . .                                                                                                                                                                   | 120 |
| 5.4  | Several example topologies in the set of basic topologies. . . . .                                                                                                                                                                                            | 122 |
| 5.5  | Illustration of latent space representation of different topologies $\Sigma$ by three-dimensional latent inputs $\mathbf{x}$ (i.e., $n_x = 3$ ), and four example topologies corresponding to the four points in the latent space (red colored ones). . . . . | 122 |
| 5.6  | Optimal topology and corresponding band structures for the third bandgap of the in-plane mode. . . . .                                                                                                                                                        | 125 |
| 5.7  | Optimal topology and corresponding band structures for the seventh bandgap of the in-plane mode. . . . .                                                                                                                                                      | 125 |
| 5.8  | Optimal topologies for the third (left) and seventh (right) bandgap of the in-plane mode in Dong et al. (2014). . . . .                                                                                                                                       | 126 |
| 5.9  | Topologies identified by DRESTO over the iterations for (a) the third bandgap and (b) the seventh bandgap. . . . .                                                                                                                                            | 128 |
| 5.10 | Optimal topology and corresponding band structures for the in-plane mode with the sum of the bandgap widths as objective function. . . . .                                                                                                                    | 129 |
| 5.11 | Topologies identified by DRESTO over the iterations for the sum of normalized bandgaps. . . . .                                                                                                                                                               | 130 |

# CHAPTER 1:

## INTRODUCTION

### 1.1 Background

Physical experiments or numerical simulations are usually carried out to explore the behavior of engineering systems. However, physical experiments can be extremely expensive in many situations, such as destructive tests and prototyping (Durrande and Le Riche 2017). Numerical simulations, as an alternative to physical experiments, are less expensive to perform. Nevertheless, to capture the behavior of engineering systems with high accuracy, many detailed simulations are extremely time-consuming to run. Moreover, in situations where a large number of system model evaluations are needed (e.g., parametric study, design optimization, uncertainty quantification), direct adoption of physical experiments or detailed numerical simulations (hereinafter referred to as *high-fidelity models*) becomes prohibitive from the cost or time perspective (Huang et al. 2011).

One way to circumvent the challenges associated with running the high-fidelity models is to construct surrogate models to approximate the input-output relationship for the high-fidelity models. As a special case of supervised machine learning, surrogate models have been applied in the engineering field for efficient analysis and design. Surrogate models are built based on a database obtained by evaluating the high-fidelity models, formally denoted as *training data*. The computational efficiency of surrogate models is greatly increased compared with the high-fidelity models. Therefore, once trained, the original high-fidelity models can be replaced by the surrogate models for efficient analysis and design of engineering systems (Jin et al. 2002, 2003; Kaymaz 2005; Xiong et al. 2007; Sudret 2008; Huang et al. 2011; Taflanidis et al. 2013; Tao et al. 2017). Various surrogate models have been used in the literature, for example, polynomial response surfaces

(Bucher and Bourgund 1990; Wang 2003), artificial neural networks (Hassoun et al. 1995; Jain et al. 1996; Hurtado and Alvarez 2001; Deng et al. 2005), support vector machines (Hearst et al. 1998; Gunn et al. 1998; Hurtado 2004; Steinwart and Christmann 2008), polynomial chaos expansions (Sudret and Der Kiureghian 2002; Choi et al. 2004; Sudret 2008; Crestaux et al. 2009), and Gaussian process model (Sacks et al. 1989; Beers and Kleijnen 2003; Rasmussen 2004; Forrester et al. 2008). Comparative studies have been done by some researchers to investigate the performance of different surrogate models (Simpson et al. 2001b; Willmes et al. 2003; Forrester et al. 2008; Paiva et al. 2010; Luo and Lu 2014). Among different surrogate models, the Gaussian process model has been gaining popularity due to its high flexibility in approximating complex functions. More importantly, it not only provides the prediction but also the associated predictive uncertainty, namely the local variance, which can be used to guide adaptive sampling and training (Picheny et al. 2010), to guide optimization (Jones et al. 1998), and be explicitly incorporated when needed (e.g., in the context of risk assessment (Jia and Taflanidis 2013), and sensitivity analysis (Le Gratiet et al. 2017)).

The quality of surrogate based analysis and design of engineering systems relies largely on the prediction accuracy of the constructed surrogate model, and therefore a high approximation accuracy is usually desired for surrogate modeling. Unlike typical machine learning methods involving big data, researchers are interested in using small number of training data to obtain desired prediction accuracy for surrogate models (Viana et al. 2010; Sudret et al. 2019). However, for complex systems with highly nonlinear input-output relationship, surrogate modeling with limited training data usually struggles to achieve the goal. To increase the prediction accuracy, a common practice is to increase the training data size so that more information about the high-fidelity model can be incorporated. Unfortunately, limited computational budgets often makes it difficult for us to

run more high-fidelity models. For high-dimensional model inputs, we face particularly significant challenge in obtaining sufficient training data, since the number of required training data to achieve desired accuracy is typically exponential with the input dimensionality (i.e., the so-called *curse of dimensionality*) (Viana et al. 2010). On the other hand, although sometimes we might be able to obtain a large number of high-fidelity data, the training cost of the surrogate model could also be beyond our affordability. This is especially the case for kernel-based surrogate models such as Gaussian process models and support vector machines, since their training time scales at least quadratically to the number of training data (DeCoste and Schölkopf 2002). As a results, how to use as less computational cost as possible to efficiently inform an accurate surrogate model has attracted extensive research interests.

In light of the high efficiency of using surrogate models for analysis and design of complex engineering systems, and the significant challenge in developing accurate surrogate models using limited computational budget, this dissertation aims at proposing strategies from different perspectives in order to effectively reduce the computational cost of surrogate modeling and improve the prediction performance. The fundamental idea of the proposed strategies is to embed extra knowledge about the high-fidelity models in the surrogate model by enriching the training data (e.g., leverage additional low-fidelity data, or censored/bounded data) and enhancing model assumption (e.g., explicitly incorporate prior knowledge about the physics of the problem, or explore low-dimensional latent structures/features), which reduces the required size of high-fidelity training data and meanwhile effectively boosts the prediction accuracy of the established surrogate model. Among various surrogate models proposed in the literature, Gaussian process models (Sacks et al. 1989; Rasmussen 2004) have been gaining popularity due to its flexibility in modeling complex functions and ability to provide closed-form predictive distributions. The predictive mean of Gaus-

sian process model is known as the Best Linear Unbiased Predictor (BLUP), and for prediction it is based on only matrix manipulations. More importantly, the Gaussian process model not only gives the mean prediction but also the associated uncertainty. This information can be further explicitly incorporated to adaptively improve the prediction accuracy and facilitate effective analysis or design (Jones et al. 1998; Jia and Taflanidis 2013; Le Gratiet et al. 2017). Therefore, the specific strategies are developed in the context of Gaussian process model, but the ideas are expected to be applicable to other types of surrogate models.

## 1.2 Literature Review

This section reviews the past research dedicated to improving the computational efficiency in building accurate surrogate models.

### 1.2.1 *Physics-constrained surrogate model*

In many practical engineering problems, we have some prior knowledge about the behavior of the interested system, and one example is that some transformations on the model input do not change the model output (i.e., symmetry or invariance). For example, in a chemical environment, the interatomic potential of a molecule or crystal is permutation invariant with respect to the ordering of the atoms in the same species (Bartók et al. 2013). In disaster assessment, the model used for super-resolving remote sensing imagery is invariant to temporal ordering of the collected low-resolution images (Valsesia and Magli 2021). In a group pile foundation system where all piles have the same design parameters, the seismic performance of the whole system will not be impacted if the ordering of the piles is switched. Additionally, power generation of a wave farm is permutation invariant to the ordering of the wave energy converters in the array. Another example of the physical constraints is the additivity feature. Mathematically, a model function can

be decomposed/expanded as the sum of the contributions from all orders of interactions between different input dimensions, and this concept is similar to the well-known high-dimensional model representation (HDMR) (Sobol' 1993; Li et al. 2001; Sobol' 2003). The function decomposition can be physically interpretable for many engineering problems. A good example is the classical many-body interaction problems in the field of quantum mechanics and molecular dynamics, where the system model output (i.e., the total potential) can be decomposed to the sum of the output from sub-systems, represented by many-body terms (i.e., individual terms and interaction terms) (Yao et al. 2017; Zhang et al. 2020a).

Standard way of building surrogate model (i.e., simply relying on input-output data, or so-called data-driven model) ignores prior knowledge on physical characteristics of the problem or input-output relationship (e.g., the aforementioned symmetry, invariance and additivity), which may increase the training cost and lead to surrogate models with lower prediction accuracy and generalization ability (Karniadakis et al. 2021). The available prior knowledge about the physical characteristics of the input-output relationship should be incorporated into surrogate modeling, since it may potentially reduce the required training cost, and improve the prediction accuracy and generalization ability of the surrogate model under the same number of training data (Rasmussen 2004; Haasdonk and Burkhardt 2007; Zhu et al. 2019).

Data augmentation is a common way to encode the prior knowledge about the symmetry and invariance features. By transforming the original training data based on the invariance information while keeping the corresponding response unchanged, new training data can be generated and added to the training set (Beymer and Poggio 1995; Niyogi et al. 1998; Dao et al. 2019). The augmented training set then constrains the surrogate model to produce predictions which are physically plausible (van der Wilk et al. 2018). However, the direct data augmentation may result in

additional computational cost to the training of surrogate models, since completely encoding some prior knowledge can easily increase the data size to the hundreds or thousands of times (DeCoste and Schölkopf 2002; Rasmussen 2004). On the other hand, for system models which demonstrate additivity features, surrogate models can be separately constructed for sub-systems and then added together to obtain the total response (Zhang et al. 2020a). However, to apply the approach, information on the contributions from subsets of the multiple-body system is needed; however, such information is not always available, and its availability (i.e., whether it can be calculated) depends on the (numerical) model used.

In recent years, physics-informed or physics-constrained techniques have emerged in the surrogate modeling community. Physics-constrained surrogate modeling aims to introduce physical laws and constraints (e.g., partial differential equations, boundary conditions) to the model construction and thus guide the model training to produce physically consistent predictive response (Frankel et al. 2020; Karniadakis et al. 2021). Such surrogate models are able to integrate additional information about the physical laws/constraints to enhance the efficiency of the pure data-driven learning, and this potentially reduce the demand for a large training set (Zhang et al. 2020b). In general, a surrogate model can be made physics-aware by introducing observational biases, inductive biases or learning biases (Karniadakis et al. 2021). The aforementioned data augmentation is often used to introduce observational biases, since the underlying laws of physics are encoded in the augmented data and the surrogate model trained based on adequate data is able to generate predictions following the corresponding physical principles. Inductive biases (i.e., model assumptions (Marino and Manic 2019)) describes the prior knowledge about the physical laws which can be incorporated by designing specific model architectures. Such prior knowledge is usually related to some symmetry/invariance features, i.e., transformations (such as reflection, permutation,

translation, and rotation) on model input do not change the model output. An example of surrogate models designed to implicitly learn symmetric/invariant features is the well-known convolutional neural network. Its architecture is designed based on convolution operations which enables the network to be invariant to some degrees of object/pixel translation, scaling and local deformation (Lawrence et al. 1997; Albawi et al. 2017). Recently, a new neural network architecture, named *Set Transformer* (Lee et al. 2019), was proposed to preserve the permutation invariance of problems involving multiple instance learning (e.g., three-dimensional shape recognition). Such problems involve set inputs rather than the conventional vector inputs, and the solutions do not rely on the ordering of the elements in a set. In addition to these neural networks, there are many examples of network architectures designed for the purpose of enforcing symmetry/invariance properties (see Bronstein et al. (2017); Zaheer et al. (2017); Tai et al. (2019); Cohen et al. (2019)). Finally, learning biases correspond to the prior knowledge about the physical constraints that can be enforced by penalizing loss functions. Many of currently well-known physics-informed neural networks (Raissi 2018; Raissi et al. 2019; Zhu et al. 2019; Sun et al. 2020) are in this category. These networks are mainly designed for problems governed by partial differential equations (PDEs), and by encoding PDEs into the loss functions, the networks are constrained to satisfy the underlying physical laws.

### 1.2.2 *Multi-fidelity surrogate model*

To build surrogate models, training data obtained by evaluating high-fidelity models is needed. For some extremely expensive high-fidelity models such as large-scale experiments and detailed numerical simulations, the number of high-fidelity runs is usually limited (e.g., due to the insufficient computational resources). With a small amount of high-fidelity data, the accuracy of established surrogate model might be compromised (Xia and Shi 2018). The aforementioned training

data is typically accurate (i.e., reflect the behavior of a system model with acceptable accuracy (Fernández-Godino et al. 2016)) but costly to obtain. By contrast, low-fidelity data is less computationally demanding to establish, such as data from coarse-resolution numerical simulations or reduced order models (Willcox 2000; Lucia et al. 2004; Willcox and Megretski 2005). Although lacking accuracy, the low-fidelity data also contains useful information about the system model (e.g., the general trend of the output from the high-fidelity model) (Liu et al. 2018a) and can be used for surrogate model training, especially considering that the low-fidelity data can be efficiently collected. However, only using low-fidelity data may also result in low prediction accuracy or bias of the surrogate model (Cheng et al. 2021).

To address the constraint, multi-fidelity surrogate models (Kennedy and O’Hagan 2000; Forrester et al. 2007; Qian and Wu 2008; Le Gratiet 2013a), which can leverage the accuracy of high-fidelity data and the efficiency of low-fidelity data, have been proposed. By integrating a limited number of high-fidelity training data and a large quantity of low-fidelity training data (Tao et al. 2019; Hao et al. 2020), the prediction accuracy and computational efficiency can be enhanced compared to the surrogate models using only single-fidelity data. This approach is helpful in the context of building predictive model using multi-fidelity data, since for many problems the simulations can be run at different fidelity level with different costs. Kennedy and O’Hagan (2000) suggested that multi-level simulations can be used to make inference about the results from the most complex simulation, and the computational efficiency can be greatly improved. Qian and Wu (2008) developed a multi-fidelity Gaussian process model that can incorporate uncertainties in the model parameters using a Bayesian approach. However, the posterior distributions of the model parameters do not have a closed form expression. To address this problem, Le Gratiet (2013a,b) proposed a new multi-fidelity Gaussian process model which extends the scaling parameter to a more practical condition

and can give the analytical expressions of the posteriors of the model parameters. These research works mainly focus on developing various model forms (e.g., using different scaling/additive terms to describe the difference between the low-fidelity model and the high-fidelity model), and then establish a Bayesian framework to calibrate the introduced model parameters. Another research focus of the multi-fidelity surrogate models is to investigate their applications in design optimization (Gano et al. 2004; Forrester et al. 2007; Keane 2012; Liu et al. 2018b). For example, Gano et al. (2004) have studied Variable Fidelity Optimization (VFO) where the high-fidelity models are approximated by the low-fidelity models and a scaling function generated by Gaussian process model. Compared to the single-fidelity surrogate model based optimization approach, the prediction accuracy of the surrogate model can be enhanced and the optimal solutions can be found more quickly when using multi-fidelity surrogate model based optimization (Forrester et al. 2007). Although there are also different types of surrogate models (e.g., radial basis function and polynomial chaos expansion) established based on data of different levels of accuracy and cost (Ng and Eldred 2012; Cai et al. 2017; Durantin et al. 2017; Palar et al. 2018), existing multi-fidelity surrogate models are mainly developed based on Gaussian process models.

### 1.2.3 *Surrogate model considering data censoring*

Standard surrogate models usually deal with high-fidelity data that are “accurate”. Here “accurate” means the high-fidelity model provides exact values for the interested output. However, in many engineering applications, for various reasons the outputs of physical experiments or computational simulations are not always the exact values of interested output. Frequently, instead of *accurate data*, we have the so-called *censored data* (Quesenberry Jr et al. 1989; Lin et al. 1997). For example, when doing experiments to investigate the deformation capacity of a reinforced con-

crete (RC) column, we define the deformation capacity as the displacement of the column when the lateral resistance drops to 20% of its peak value (Gardoni et al. 2002; Wu et al. 2004). However, for some columns we can often observe that even when the lateral resistance has dropped below 80% of its peak value, the column still does not reach its ultimate deformation. In this case, we can only obtain a certain range of the deformation capacity. Such type of practical problems are called *censoring*, and the corresponding inputs/outputs of the experiments or simulation are called censored data. Another common example for censored data is the data from partially converged numerical simulations, which instead of providing the the exact values for the interested output provides an upper bound or lower bound or both for the interested output. Censored data also contains useful information about the system model (i.e., interval of the responses), and this type of information can be fully exploited to establish better surrogate models.

Regression models have been developed to fit censored data (i.e., censored regression models), and most of these regression models are of parametric forms (e.g., Weibull or Gamma distributions) (DeMaris 2004; Hashimoto et al. 2010; Pratavia et al. 2019) or semi-parametric forms (e.g., proportional hazard models) (Fox and Weisberg 2002; Harrell 2015). Many research efforts have also been devoted to modeling censored data using surrogate models such as support vector machines (Shivaswamy et al. 2007; Khan and Zubek 2008; Goldberg and Kosorok 2017) and neural networks (Zhu et al. 2016; Luck et al. 2017; Gensheimer and Narasimhan 2019). These surrogate models mainly focus on designing specific loss/cost functions to deal with censored data (Zhao and Feng 2020). In the field of Gaussian process model, in order to address censored data, many approaches have been proposed (Dubrule and Kostov 1986; Kostov and Dubrule 1986; Stein 1992; Militino and Ugarte 1999; Abrahamsen and Benth 2001). Nevertheless, these methods all suffer from some drawbacks, e.g., the amount of information contained in accurate data and censored data is not

differentiated (De Oliveira 2005). To address the issues, De Oliveira (2005) proposed to use *data augmentation* (note: different from the aforementioned one for incorporating prior knowledge) to deal with the censored data. This approach treats the censored data as unknown model parameters and the extended model with these parameters is then built. Later, in Groot's study (2012), the likelihood of the model parameters based on data subject to censoring is presented. The likelihood results in no analytical expressions for the posterior distributions of the model parameters, so expectation propagation (Minka 2001, 2013) is applied to approximately make inference for the model. Similarly, Gammelli et al. (2020) developed a censored likelihood function based on the Tobit censored regression model (Chib 1992), and then approximately estimated the posterior predictive distribution via expectation propagation. Essentially, the aforementioned methods utilize Bayesian approaches for inference purpose and the censored data are taken into account by modifying the likelihood functions.

#### *1.2.4 Dimension reduction assisted surrogate model*

Parametrization and training of many surrogate models typically face difficulties when the model inputs become high-dimensional (Lataniotis et al. 2020). This is because the surrogate models suffer from the curse of dimensionality (Verleysen and François 2005), and the number of training data required for model construction (to achieve desired prediction accuracy) grows exponentially with the input dimension (Tripathy et al. 2016). Many engineering problems involve high-dimensional inputs, and directly building surrogate models using the high-dimensional inputs may result in low approximation quality (Palar and Shimoyama 2018).

To address the challenges in building surrogate models for problems with high-dimensional inputs, many approaches have been proposed in previous literature. One common case is to build

surrogate models for models with high-dimensional structured data (Ramsay and Dalzell 1991; Tripathy et al. 2016), which shows strong correlations between different dimensions (i.e., time-series or spatial quantities) (Lataniotis et al. 2020). In such problems, some general dimension reduction techniques can be applied to the model input before surrogate modeling, which involves projecting the original high-dimensional input to a lower-dimensional latent space (Lataniotis et al. 2018). Typically, the dimension reduction techniques can be categorized to two classes: unsupervised dimension reduction and supervised dimension reduction (Joy et al. 2019). One of the supervised dimension reduction techniques having attracted extensive attention is the active subspace (Constantine et al. 2014; Constantine 2015), which is a low-dimensional representative subspace of the model input obtained by detecting and exploiting the most important directions of the model output, i.e., the directions of strongest variability of the model. An active subspace involves finding an orthogonal projection matrix to project the model input to a low-dimensional subspace, and one needs to evaluate the gradients of the model at a number of input sample points to obtain the projection matrix (Constantine 2015). This method has been successfully used to reduce input dimension in surrogate model based aerodynamic shape optimization (Grey and Constantine 2018; Li et al. 2019a). However, one drawback of this method is that for those models with a high-dimensional input which do not have explicit mathematical expressions or computer functions, the computation of the gradients typically involves significant or sometimes even prohibitive computational efforts (Palar and Shimoyama 2018). An alternative approach is the partial least-squares algorithm (Bouhlel et al. 2016a,b; Papaioannou et al. 2019). This method identifies the directions with largest significance by analyzing the covariance between the model input and output, and does not require the calculation of the gradients.

For unsupervised dimension reduction, commonly used techniques include Karhunen–Loève (KL) expansion (Ghanem and Spanos 2003), principal component analysis (PCA) (Jolliffe 2002), kernel PCA (Wu et al. 1997), and locally linear embedding method (Roweis and Saul 2000). These dimension reduction methods have been successfully applied to address the issues with surrogate modeling for high-dimensional inputs. For example, KL expansion has been combined with polynomial chaos expansion to explore a high-dimensional groundwater model (Zhang et al. 2015; Dai et al. 2016). Kernel PCA was also used with radial basis function networks to accelerate the high-dimensional evolutionary optimization (Kapsoulis et al. 2016). Raponi et al. (2020) addressed the scalability of surrogate-assisted global optimization by integrating it with PCA. In addition to optimization tasks, unsupervised dimension reduction techniques have also been used to high-dimensional reliability analysis (Zhou and Peng 2020). Among these techniques, PCA discovers the linear projections of the data with maximum variance, or equivalently, the lower dimensional subspace that yields the minimum squared reconstruction error (Tipping and Bishop 1999). Compared to other techniques, PCA has the advantage of simplicity in implementation, and more importantly, PCA provides a bidirectional transformation between the original space and the latent space (Jolliffe 2002). This especially is of great importance to optimization problems since the optimal design established in the latent input space needs to be transformed back to the original input space.

### *1.2.5 Surrogate model based optimization*

Considering the high efficiency of surrogate models, any global optimization algorithm can be used to efficiently find the optimum by replacing the high-fidelity models with the surrogate models. Early surrogate based optimization methods mainly focused on response surface surrogate

models (Jones et al. 1998) where response surfaces were used to fit a small number of high-fidelity data, and then used in global optimization instead of the expensive high-fidelity models. Later, optimization approaches based on other alternative surrogate models were investigated as well, e.g., Gaussian process models (Simpson et al. 2001a; Forrester and Keane 2009), radial basis functions (Regis and Shoemaker 2005; Jakobsson et al. 2010), and neural networks (Muralitharan et al. 2018; Zou and Chen 2021). Different ways of constructing surrogate models for optimization and their advantages have been summarized (Jin et al. 2001; Forrester et al. 2008; Forrester and Keane 2009). However, directly replacing the high-fidelity model with the surrogate model for optimization could easily lead to a local optimal solution. This is because this way of search essentially assumes that the predictor is an accurate prediction of the actual output values, i.e., it puts too much emphasis on exploiting the predictor but puts no emphasis on exploring points where we are uncertain (Jones et al. 1998). To address this, we need to sample more at points where the uncertainty is high. Therefore, another focus of the research on surrogate based optimization is how to adaptively select the infill points so that the optimum can be obtained quickly and accurately. In order to trade-off between exploiting the predictor and exploring the uncertain regions, some infill criteria have been developed to guide the adaptive design of experiment. Among those infill criteria which balance exploitation and exploration (see Forrester et al. (2008)), expected improvement is commonly used, where infill points are adaptively added at locations that give the largest expected improvement in the objective function based on the surrogate model (Forrester and Keane 2009; Jones 2001).

### **1.3 Research Gaps and Motivation**

This dissertation focuses on proposing multiple strategies to efficiently train accurate Gaussian process models (i.e., obtain high prediction accuracy and generalization ability with as less compu-

tational cost as possible) from different perspectives. More specifically, we enhance the Gaussian process models by (i) embedding physical constraints, (ii) integrating high-fidelity training data and low-fidelity training data as well as incorporating both accurate and censored training data, and (iii) applying dimension reduction to handle high-dimensional inputs. Section 1.2 has already reviewed the existing research on these topics. In this section, the research gaps in these topics will be identified and also the research motivation of this dissertation will be stated.

### *1.3.1 Physics-constrained Gaussian process model*

Existing research on physics-constrained surrogate models is mainly related to neural networks, including proposing specialized network architectures to include symmetry/invariance features of the problems, or imposing loss function penalty constraints to enable the predictions to satisfy the governing partial differential equations (PDEs) of the problems. However, currently there is only a few literature on physics-informed or physics-constrained Gaussian process models, and the existing limited research works are focusing on physical knowledge represented in the form of PDEs (Yang et al. 2018; Wu et al. 2019; Albert 2019; Hanuka et al. 2019; Tong et al. 2020). In many engineering problems, we frequently have prior knowledge about the symmetry/invariance features of the input-output relationship, and current physics-constrained Gaussian process models has little research on how to incorporate such physical constraints. In addition, the prior knowledge about the additivity feature (i.e., the many-body expansion principle) has not exploited under the topic of physics-constrained Gaussian process models. Overall, there is strong need to develop Gaussian process models that can explicitly incorporate symmetry, invariance, and additivity features.

### *1.3.2 Multi-fidelity Gaussian process model considering censored data*

There have been a lot of research works on multi-fidelity Gaussian process models, but existing multi-fidelity Gaussian process models usually use a constant scaling factor to match the low-fidelity model to the high-fidelity model (Forrester et al. 2008; Kennedy and O'Hagan 2000; Kuya et al. 2011). Also, measurement error is usually not explicitly modeled. In order to accommodate more complex relationships between models of different fidelity, a more general scaling factor (e.g., non-constant scaling factor that varies with the input) should be used, and also measurement error needs to be considered. On the other hand, the focus of existing research on incorporating censored data has been in the development of Gaussian process models rather than multi-fidelity Gaussian process models. To the author's best knowledge, multi-fidelity Gaussian process models that can incorporate both accurate data and censored data have not been studied. In this setting, the inclusion of censored data entails significant computational challenges, i.e., in establishing the likelihood function needed for calculating the posterior distribution of the model parameters. To facilitate the application of multi-fidelity Gaussian process models for various engineering problems where multi-fidelity data and censored data are available, the above research gaps need to be closed and associated challenges need to be addressed.

### *1.3.3 Dimension reduction assisted Gaussian process model and its application in optimization*

Though effective, surrogate based optimization (also referred as efficient global optimization) still faces challenges for problems with high-dimensional design variables stemming from the difficulty in building accurate surrogate models for problems with high-dimensional inputs (i.e., design variables). Therefore, direct application of Gaussian process model based optimization to high-dimensional optimization faces challenges. As a popular dimension reduction technique, PCA has

been mostly used to reduce output dimensions to facilitate efficient Gaussian process models for high-dimensional outputs (Jia and Taflanidis 2013; Jia et al. 2016; Aversano et al. 2019; Li et al. 2020). However, it has not been thoroughly investigated in assisting input dimension reduction for Gaussian process modeling, especially in the context of design optimization. To the author's best knowledge, the only literature on PCA assisted Gaussian process model for high-dimensional optimization is in Raponi et al. (2020), which focused on optimization problems with high-dimensional continuous design variables. However, for high-dimensional discrete optimization (i.e., design variables take discrete values) or binary optimization (i.e., design variables take 0 or 1), which is also common in engineering field (e.g., topology optimization) and typically more challenging, there is no research yet on integrating PCA with Gaussian process model. Furthermore, PCA typically performs well when the data exhibits some latent/low-dimensional features. Consequently, if the model inputs only demonstrate such features over part(s) of the design space rather than the entire design space, PCA cannot be applied directly and some data processing procedures might be needed. This issue has not been addressed in the existing literature. With the research gaps identified, this dissertation aims to provide a deeper insight into the PCA assisted Gaussian process model for high-dimensional optimization, with a focusing on problems with high-dimensional binary design variables.

#### **1.4 Summary of contributions and objectives**

Overall, the proposed research seeks to facilitate efficient Gaussian process model based engineering system analysis and design. The key novel contribution of this dissertation is on proposing strategies from different perspectives in order to effectively reduce the computational cost of Gaussian process modeling and improve the prediction performance. The fundamental idea of the pro-

posed strategies is to embed extra knowledge about the high-fidelity models in the surrogate model by enriching the training data (e.g., leverage additional low-fidelity data, or censored/bounded data) and enhancing model assumption (e.g., explicitly incorporate prior knowledge about the physics of the problem, or explore low-dimensional latent structures/features), which can effectively reduce the required size of high-fidelity training data and meanwhile effectively boosts the prediction accuracy of the established Gaussian process model. The strategies are further examined in the context of several engineering applications. In particular, the main research objectives of this dissertation are:

- Develop a physics-constrained Gaussian process model to efficiently incorporate our prior knowledge about physical constraints/characteristics of the input-output relationship by designing specific kernels.
- Establish a more general multi-fidelity Gaussian process model integrating a small number of expensive high-fidelity data and a large number of cheap low-fidelity data by developing a more general model form to enhance the existing multi-fidelity Gaussian process model and by developing a corresponding Bayesian calibration framework.
- Propose a multi-fidelity Gaussian process model capable of integrating training data with different level of accuracy (high-fidelity data and low-fidelity data) and completeness (i.e., accurate data and censored data).
- Develop an efficient surrogate modeling approach for design optimization problems with high-dimensional binary model inputs (or design variables) by integrating dimension reduction technique and Gaussian process model and by using adaptive sampling scheme.

## 1.5 Outline of Dissertation

The remainder of this dissertation is organized as follows. Chapter 1 provides the research background for this dissertation and presents the literature review on four main research topics. Then, the research gap for each research topic are identified and the research objectives are described. Chapter 2 reviews the construction of a standard Gaussian process model, along with some related important topics such as design of experiment and multidimensional outputs. This sets the foundation for the proposed strategies for improving the performance of the Gaussian process model. Chapter 3 proposes a physics-constrained Gaussian process model which can accurately predict response exhibiting symmetry, invariance, and additivity features. These physical constraints are first mathematically formulated, based on which physics-constrained kernels are then designed and used for the Gaussian process model to effectively and efficiently incorporate the physical constraints. An illustrative example on predicting hydrodynamic characteristics of WECs in an array using the proposed algorithm is presented, and the corresponding research findings are summarized. Chapter 4 proposes a general multi-fidelity Gaussian process model integrating low-fidelity data and high-fidelity data considering censoring in high-fidelity data. The uncertain model parameters are calibrated using a Bayesian approach. The likelihood function considering censored data is first estimated by adopting data augmentation algorithm, and closed form conditional posteriors are derived for the augmented model parameters. Then Gibbs sampling is used to efficiently generate samples from the posterior distributions for the model parameters, and the generated posterior samples are used to establish the posterior statistics for the output predictions at new inputs. As an illustrative example, the proposed model is applied to establish predictive model for the deformation capacity of RC columns, and the results are presented and discussed. Chapter 5 proposes a

dimension reduction assisted Gaussian process model within the context of optimization to address the challenges stemming from high dimensionality of binary inputs (i.e., design variables take 0 or 1). An optimization problem with high-dimensional discrete design variables is first formulated and the associated computational challenges are discussed. The proposed algorithm is then presented, focusing on dimension reduction of design space and adaptive surrogate based optimization. Finally, an illustrative example on topology optimization is demonstrated, and the performance of the proposed algorithm is discussed. Chapter 6 concludes the dissertation by highlighting the main findings and suggesting some future research directions.

## CHAPTER 2:

### REVIEW OF GAUSSIAN PROCESS MODEL

#### 2.1 Introduction

Gaussian process model has been widely used for regression and classification problems in machine learning communities (Rasmussen 2004). Compared to other surrogate models, the Gaussian process model has been gaining popularity due to its flexibility in modeling complex functions and ability to provide closed-form predictive distributions. In engineering fields, we are often more interested in responses which are continuous quantities, and thus this dissertation mainly focuses on Gaussian process model for regression tasks. A traditional regression problem typically aims at modeling the relationship between responses (i.e., outputs) and explanatory variables (i.e., inputs), and describing the relationship by a function which fits the given training data the best (Ghanizadeh et al. 2021). In Gaussian process model, the fundamental hypothesis is that there are infinite possible functions that can describe the input-output relationship, and any finite number of these functions follow a multivariate Gaussian distribution (Rasmussen 2004). Note that the probability distribution corresponds to the prior distribution over the functions. After conditioning on training data, the prior distribution over the functions is updated to the posterior distribution which is then used to establish a predictive distribution over any new input (Ghanizadeh et al. 2021).

#### 2.2 Gaussian Process Model

For an unknown expensive function (i.e., an alternative way of describing high-fidelity model) with input vector  $\mathbf{x} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{n_x}$ , a Gaussian process model can be adopted to approximate the deterministic function output  $y(\mathbf{x}) \in \mathbb{R}$  given a set of *observations*. Here, the observations

are the function responses or the model outputs over a set of selected inputs (i.e., experiments). The core principle is to assume the target function  $y(\mathbf{x})$  as a realization of a global model  $m(\mathbf{x})$  and a Gaussian process  $\varepsilon(\mathbf{x})$ . The Gaussian process model is written as

$$y(\mathbf{x}) = m(\mathbf{x}) + \varepsilon(\mathbf{x}) \quad \text{where } \varepsilon(\mathbf{x}) \sim \mathcal{GP}(0, k(\mathbf{x}, \mathbf{x}')) \quad (2.1)$$

This indicates that the residual between the function and the global model is modeled by a Gaussian process, and the mean of the Gaussian process prior over the target function is also  $m(\mathbf{x})$ . The global model  $m(\mathbf{x})$  is typically assumed as a regression model, constant or even zero. For generalization and interpretability purpose, this dissertation uses a regression model, expressed by  $m(\mathbf{x}) = \mathbf{f}(\mathbf{x})^T \boldsymbol{\beta}$ , where  $\mathbf{f}(\mathbf{x})$  is the  $q$ -dimensional vector of basis function and  $\boldsymbol{\beta} = [\beta_1, \beta_2, \dots, \beta_q]^T$  is the regression coefficients vector. Typically, polynomials of  $\mathbf{x}$  are used as the basis function, e.g., linear or quadratic functions of  $\mathbf{x}$ . For the linear and full quadratic cases,  $q$  equals to  $(n_x + 1)$  and  $(n_x + 1)(n_x + 2)/2$ , respectively.

After the mean function is set, the property of the Gaussian process model is fully determined by the covariance function or the so-called kernel,  $k(\mathbf{x}, \mathbf{x}') = \sigma^2 \Psi(\mathbf{x}, \mathbf{x}')$ , where  $\sigma^2$  is the variance and  $\Psi(\mathbf{x}, \mathbf{x}')$  is the correlation function. The following equations show several commonly used kernel functions:

- Exponential kernel

$$k(\mathbf{x}, \mathbf{x}') = \sigma^2 \prod_{i=1}^{n_x} \exp\left(-\frac{|x_i - x'_i|}{\theta_i}\right) \quad (2.2)$$

- Squared exponential kernel

$$k(\mathbf{x}, \mathbf{x}') = \sigma^2 \prod_{i=1}^{n_x} \exp\left(-\frac{|x_i - x'_i|^2}{2\theta_i^2}\right) \quad (2.3)$$

- Matérn 5/2 kernel

$$k(\mathbf{x}, \mathbf{x}') = \sigma^2 \prod_{i=1}^{n_x} \left(1 + \frac{\sqrt{5}|x_i - x'_i|}{\theta_i} + \frac{\sqrt{5}|x_i - x'_i|^2}{3\theta_i^2}\right) \exp\left(-\frac{\sqrt{5}|x_i - x'_i|}{\theta_i}\right) \quad (2.4)$$

- Matérn 3/2 kernel

$$k(\mathbf{x}, \mathbf{x}') = \sigma^2 \prod_{i=1}^{n_x} \left(1 + \frac{\sqrt{3}|x_i - x'_i|}{\theta_i}\right) \exp\left(-\frac{\sqrt{3}|x_i - x'_i|}{\theta_i}\right) \quad (2.5)$$

- Rational quadratic kernel

$$k(\mathbf{x}, \mathbf{x}') = \sigma^2 \prod_{i=1}^{n_x} \left(1 + \frac{|x_i - x'_i|^2}{2\alpha\theta_i^2}\right)^{-\alpha} \quad (2.6)$$

where  $x_i$  and  $x'_i$  are the  $i$ th component of the input  $\mathbf{x}$  and  $\mathbf{x}'$ , respectively.  $\theta_i$  is the scale correlation parameter (i.e., length-scale) in the correlation function for the  $i$ th component of the inputs. Note that we can also use the same  $\theta$  for all components, and the corresponding kernel is the isotropic kernel. For the rational quadratic kernel,  $\alpha$  is the scale mixture parameter. If Gaussian process is used for interpolation/regression purpose, typically positive correlation functions are used considering that closer points will behave more similarly than points that are further away and the correlation between two points is expected to decay with the distance between the points.

### 2.3 Model Parameter Calibration

The prior distribution over the functions in Eq. (2.1) is then updated by a set of given data points, i.e., the calibration of the uncertain model parameters in the mean function and the covariance function. The calibration of the Gaussian process model first requires the creation of a database of  $n$  observations based on the high-fidelity model, corresponding to evaluations of the response vector  $\mathbf{Y} = \{y(\mathbf{x}^h); h = 1, \dots, n\}$  for different inputs  $\mathbf{X} = \{\mathbf{x}^h; h = 1, \dots, n\}$ . The selection of these inputs is frequently referenced as *design of experiments* (DOE) and will impact the accuracy of the surrogate model. Typically, to ensure the surrogate model have good accuracy over the design space, selections that can evenly fill the design space is used, and one popular choice is the Latin Hypercube Sampling (LHS). Also, some adaptive sampling schemes can be applied to improve the accuracy in target regions. More details on the DOE will be discussed later. The set  $\{\mathbf{X}, \mathbf{Y}\}$  is the so-called training set.

The unknown model parameters introduced by  $m(\mathbf{x})$  and  $k(\mathbf{x}, \mathbf{x}')$  include regression coefficients  $\beta$ , variance  $\sigma^2$ , and length-scale  $\theta$ . Here,  $\sigma^2$  and  $\theta$  introduced by  $k(\mathbf{x}, \mathbf{x}')$  are frequently referred to as *hyperparameters*. In order to calibrate these parameters using the training set, maximum likelihood estimation is usually used to obtain point estimates of the parameters. Due to the Gaussian properties, the likelihood function  $p(\mathbf{Y}|\beta, \sigma^2, \theta)$  is represented by

$$L(\beta, \sigma^2, \theta|\mathbf{Y}) = \frac{1}{(2\pi\sigma^2)^{n/2}|\Psi|^{1/2}} \exp \left[ -\frac{1}{2\sigma^2}(\mathbf{Y} - \mathbf{F}\beta)^T \Psi^{-1}(\mathbf{Y} - \mathbf{F}\beta) \right] \quad (2.7)$$

where  $\mathbf{F}(\mathbf{x}) = [\mathbf{f}(\mathbf{x}^1)^T, \mathbf{f}(\mathbf{x}^2)^T, \dots, \mathbf{f}(\mathbf{x}^n)^T]^T$ , and  $\Psi$  is the correlation matrix with the  $ij$ th element calculated by  $\Psi(\mathbf{x}_i, \mathbf{x}_j)$ . The derivatives of the likelihood function with respect to  $\beta$  and  $\sigma^2$  can be directly calculated, and by setting them to zero the maximum likelihood estimates (MLE) of both

parameters can be represented by  $\boldsymbol{\theta}$ . Then optimization algorithms such as genetic algorithm can then be employed to find the MLE of the parameter  $\boldsymbol{\theta}$ .

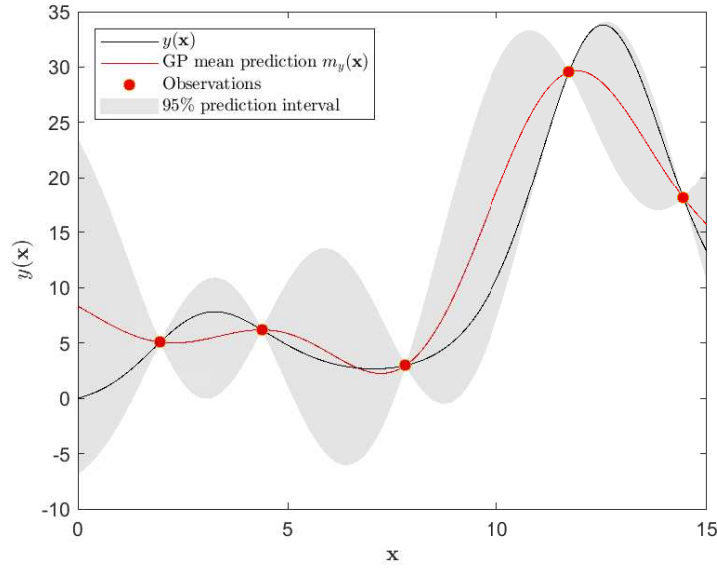
## 2.4 Prediction

Conditional on the given observations  $\mathbf{Y}$  and the optimal selection of the model parameters, the Gaussian process model gives a prediction at any new input  $\mathbf{x}^0$  that follows a Gaussian distribution with mean  $m_y(\mathbf{x}^0)$  and variance  $s_y^2(\mathbf{x}^0)$  given by

$$m_y(\mathbf{x}^0) = \mathbf{f}(\mathbf{x}^0)^T \boldsymbol{\beta}^* + \mathbf{k}(\mathbf{x}^0)^T \mathbf{K}^{-1}(\mathbf{Y} - \mathbf{F}\boldsymbol{\beta}^*) \quad (2.8)$$

$$s_y^2(\mathbf{x}^0) = k(\mathbf{x}^0, \mathbf{x}^0) - \mathbf{k}(\mathbf{x}^0)^T \mathbf{K}^{-1} \mathbf{k}(\mathbf{x}^0) + \mathbf{u}(\mathbf{x}^0)^T (\mathbf{F}^T \mathbf{K}^{-1} \mathbf{F})^{-1} \mathbf{u}(\mathbf{x}^0) \quad (2.9)$$

where  $\mathbf{k}(\mathbf{x}^0)$  is the  $n$ -dimensional covariance vector between the new input  $\mathbf{x}^0$  and each of the inputs in  $\mathbf{X}$  and  $\mathbf{k}(\mathbf{x}^0) = [k(\mathbf{x}^0, \mathbf{x}^1), k(\mathbf{x}^0, \mathbf{x}^2), \dots, k(\mathbf{x}^0, \mathbf{x}^n)]^T$ .  $\mathbf{K}$  is the covariance matrix evaluated over  $\mathbf{X}$  and  $\mathbf{K} = \mathbf{K}(\mathbf{X}, \mathbf{X})$ . For simplicity of the expression,  $\mathbf{u}(\mathbf{x}^0) = \mathbf{F}^T \mathbf{K}^{-1} \mathbf{k}(\mathbf{x}^0) - \mathbf{f}(\mathbf{x}^0)$ . Here  $\boldsymbol{\beta}^* = (\mathbf{F}^T \mathbf{K}^{-1} \mathbf{F})^{-1} (\mathbf{F}^T \mathbf{K}^{-1} \mathbf{Y})$  corresponds to the MLE of the parameter  $\boldsymbol{\beta}$ . Note that the mean prediction in Eq. (2.8) and the associated variance in Eq. (2.9) are typically used when applying Gaussian process model. If needed, we can also calculate higher-order moments based on the posterior predictive distribution. For illustration purpose, Figure 2.1 shows an example of the Gaussian process model prediction, where the black line represents the unknown function to be predicted, the red line is the mean prediction by the constructed Gaussian process model (i.e.,  $m_y(\mathbf{x})$ ), the red dots are the observations, and the shaded area corresponds to the 95% confidence interval.



**Figure 2.1:** Illustrative example of Gaussian process model (GP in the figure stands for Gaussian process model).

The predictive mean in Eq. (2.8) is also known as the Best Linear Unbiased Predictor (BLUP). For prediction at any new input point, it is based on only matrix manipulations, and can provide an exact interpolation, which means that when the input is the same as one of the training data, the prediction will have the exact value of the actual (i.e., expensive high-fidelity model based) output. More importantly, the Gaussian process model not only gives the prediction but also the associated uncertainty, namely the local variance of the prediction error. This local variance can be explicitly incorporated in guiding adaptive design of experiments to adaptively improve the accuracy of the surrogate model and further facilitating effective analysis and design (e.g., sensitivity analysis or optimization).

## 2.5 Design of Experiment

The preparation of the training set  $\{\mathbf{X}, \mathbf{Y}\}$  requires the selection of a set of inputs, i.e., DOE (Sacks et al. 1989). These experiments may significantly impact the prediction accuracy of the

Gaussian process model. Typically, to ensure the surrogate model have good accuracy over the entire design space, selection that can evenly fill the design space is used, i.e., space-filling techniques. Such techniques usually obtain the optimal spread of the experiments based on some geometric criteria which are defined appropriately (Johnson et al. 1990; McKay et al. 2000), and are model-independent (i.e., do not rely on the information provided by the model). One popular choice of space-filling DOE is the Latin Hypercube Sampling (LHS). For prediction tasks, LHS is commonly used due to its space-filling property and also convenience in implementation.

Alternatively, adaptive DOE schemes (Shewry and Wynn 1987; Sacks et al. 1989) can be applied to sequentially add samples to the training set. Under such schemes, information extracted from the already trained surrogate model can be used to develop some criteria to determine where the new sample(s) should be located. For example, the Gaussian process model provides predictive variance (see Eq. (2.9)) which measures the level of uncertainty of the predictions, and taking advantage of this information the surrogate model can be guided to adaptively add more samples in locations where predictive uncertainty is higher. More importantly, in many situations where we are more concerned about the prediction accuracy at certain region rather than the global prediction accuracy (e.g., optimization and reliability analysis), adaptive DOE strategies are often applied to achieve a balance between reducing the global prediction uncertainty and exploring the regions of interest (Picheny et al. 2010). Such adaptive sampling schemes can enable more wise allocation of the computational budget when evaluating high-fidelity model for the selected inputs for the training data and adaptively improve the prediction accuracy of the Gaussian process model.

In the following chapters, LHS is used for prediction tasks, while for the optimization task, an adaptive sampling scheme is used.

## 2.6 Multidimensional outputs

The Gaussian process model developed above is specifically for scalar model output (i.e.,  $y(\mathbf{x}) \in \mathbb{R}$ ). In many engineering problems, the interested outputs are of multiple (or even high) dimensions (i.e.,  $\mathbf{y}(\mathbf{x}) \in \mathbb{R}^{n_y}$ ) where  $n_y > 1$ . If there is no correlation or there are weak correlations between different components of the model outputs and also  $n_y$  is not large, we can separately build one Gaussian process model for each dimension of the outputs (Jia and Taflanidis 2013). For all the Gaussian process models, same type of regression model and same type of kernel function can be selected while the uncertain model parameters for each Gaussian process model is calibrated separately (Zhang et al. 2020a). In the end, the prediction for the multidimensional outputs are assembled sequentially by the mean predictions from all the calibrated Gaussian process model. However, if different components of the model outputs show strong correlations with each other and also  $n_y$  is large, building a separate Gaussian process model for each output component becomes computationally expensive and may lead to poor prediction accuracy due to the negligence of the correlations between the outputs. To address the issue, dimension reduction techniques can be adopted to explore the correlations within the high-dimensional outputs and establish a low-dimensional latent outputs representation (Blatman and Sudret 2010; Jia et al. 2016), and the Gaussian process model is then conveniently built for the low-dimensional latent outputs. Due to the low dimensionality of the latent outputs, a single Gaussian process model can be built with respect to each of the latent outputs, and finally, the prediction of the original multidimensional outputs can be obtained by transforming back the predicted latent outputs.

## 2.7 Derivatives of model outputs

Derivative information about the unknown function  $y(\mathbf{x})$  is sometimes important in engineering analysis and design, e.g., in reliability analysis, the first-order reliability method and various optimization algorithms generally require gradients of the limit-state function (Gong et al. 2014). To ensure the performance of such analysis and design, the derivatives should be preserved when building a surrogate model to approximate the unknown function.

In training of the Gaussian process model, the derivative information of the output is typically not available and not used, rather the training aims to match the model prediction to the output. Therefore, theoretically the established Gaussian process model cannot guarantee its derivatives will match exactly the derivatives of the underlying expensive function. However, when derivative information is available in the training data, we can incorporate it into the training of the Gaussian process model to enhance its accuracy (including preserving the derivative information). This corresponds to the idea of gradient-enhanced surrogate model, and is one form of the idea of co-kriging (Forrester et al. 2008; Han et al. 2013).

On the other hand, if the Gaussian process model is accurate or as the prediction accuracy improves, the derivative information is expected to be captured with increasing accuracy as well. If we assume that the established Gaussian process model has enough accuracy, then because the predictive distribution has analytical expressions, we can efficiently compute its first or even higher-order derivatives. But note that the existence of the derivatives is determined by the differentiability of its mean and kernel functions. For example, if squared exponential kernel is used, the Gaussian process model will have infinitely many derivatives since the kernel is infinitely differentiable. Note that the exponential kernel Eq. (2.2) should not be used in situations where differentiability

is required because a Gaussian process model with such a kernel is not differentiable, although it is integrable.

## CHAPTER 3:

### PHYSICS-CONSTRAINED GAUSSIAN PROCESS MODEL THROUGH KERNEL DESIGN

#### 3.1 Introduction

Standard way of building surrogate model solely relies on input/output data, i.e., purely data-driven model. However, this may lead to predictions which are not physically consistent or plausible (Karniadakis et al. 2021), especially when the training data is not sufficient and extrapolation occurs, resulting in poor prediction accuracy and generalization ability. Although enriching the training set is an intuitive way to improve the performance of the surrogate model, it is still not efficient since extracting interpretable information from a large number of data is difficult for most surrogate models (Karniadakis et al. 2021). Moreover, obtaining so many training data is often impractical due to the high cost of running the high-fidelity model.

In recent years, physics-informed or physics-constrained methods have merged as an effective way to address the aforementioned issue in surrogate modeling. In many engineering problems, we have some prior knowledge about the behavior or fundamental physical laws of the interested system. Physics-constrained surrogate model integrates our available prior knowledge about the physical constraints in the model development, and in such way, a more “informative prior” is provided on top of the training data, guiding the surrogate model to learn the underlying physical laws. This way of building surrogate model can potentially reduce the required training cost (Rasmussen 2004), and improve the prediction accuracy and generalization ability of the surrogate model under the same number of training data (Haasdonk and Burkhardt 2007).

A physics-constrained Gaussian process model through kernel design is proposed in this chapter to explicitly incorporate in the surrogate model two types of physical constraints: (i) symmetry

and invariance features, and (ii) additivity feature. More specifically, we first mathematically formulate the physical constraints including symmetry, invariance, and additivity. An invariant kernel is then developed to incorporate the invariance and symmetry information and integrated into the Gaussian process model. This developed invariant kernel enables efficient incorporation of the prior knowledge in the Gaussian process model, and eliminates the need of handcrafting the input features. Then, taking advantage of the many-body expansion principle, a new kernel is developed to embed the additive feature in the Gaussian process model. The kernel exhibits a similar additive feature to the model output, and is thus named as additive kernel. Finally, the proposed invariant kernel is combined with the additive kernel to establish an integrated kernel, which is used to construct more informative Gaussian process models that explicitly include all the physical constraints. To demonstrate the performance of the proposed algorithm, it is applied to predict hydrodynamic characteristics of wave energy converters (WECs) in an array which exhibit both symmetry and invariance features and the additivity feature. The results show that compared to the standard Gaussian process models, the proposed physics-constrained Gaussian process models require less training data to achieve desired accuracy in predicting the hydrodynamic characteristics, and is less vulnerable to the curse of dimensionality.

## **3.2 Physical Constraints: Invariance, Symmetry, and Additivity**

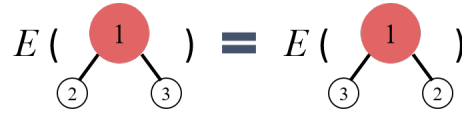
### *3.2.1 Invariance and symmetry*

In many problems, the interested system often exhibits invariance or symmetry features, i.e., some transformations on the model input do not change the model output. For example, in a chemical environment, the interatomic potential ( $E$ ) of a molecule or crystal is permutation invariant with respect to the ordering of the atoms in the same species (Bartók et al. 2013) (see Figure 3.1).

For image recognition, the label of an image such as a handwritten digit can be invariant to the translation and small rotation of the image pixels (LeCun et al. 1995). Following the general definition (Kondor 2008; Ginsbourger et al. 2012), such invariance and symmetry knowledge about a function/model  $y(\mathbf{x})$  can be formulated as

$$y(\mathbf{x}) = y(t(\mathbf{x})) \quad \forall \mathbf{x} \in \mathcal{X} \quad \forall t \in \mathcal{T} \quad (3.1)$$

where  $t(\mathbf{x})$  is an operation/transformation on the input  $\mathbf{x}$  that determines the invariance/symmetry, and  $\mathcal{T}$  represents a finite group of all possible such operations. Note that the function  $y(\mathbf{x})$  is also invariant to the compositions of the operations (van der Wilk et al. 2018), and thus the group  $\mathcal{T}$  includes both the operations and their possible compositions. For engineering problems, possible operations typically include permutation, reflection, translation, and rotation which can be discrete or continuous, finite or infinite transformations. Note that this chapter is mainly concerned with discrete and finite transformations, and also problems with a small to medium number of transformations in the operation group.



**Figure 3.1:** Permutation invariance property of interatomic potential  $E$ .

### 3.2.2 Additivity

For problems with multidimensional model input, the impact of the model input on the model output may be independent or cooperative (Li et al. 2000) (i.e., the model output may depend on single dimensions of the input, or interactions between different dimensions of the input). Math-

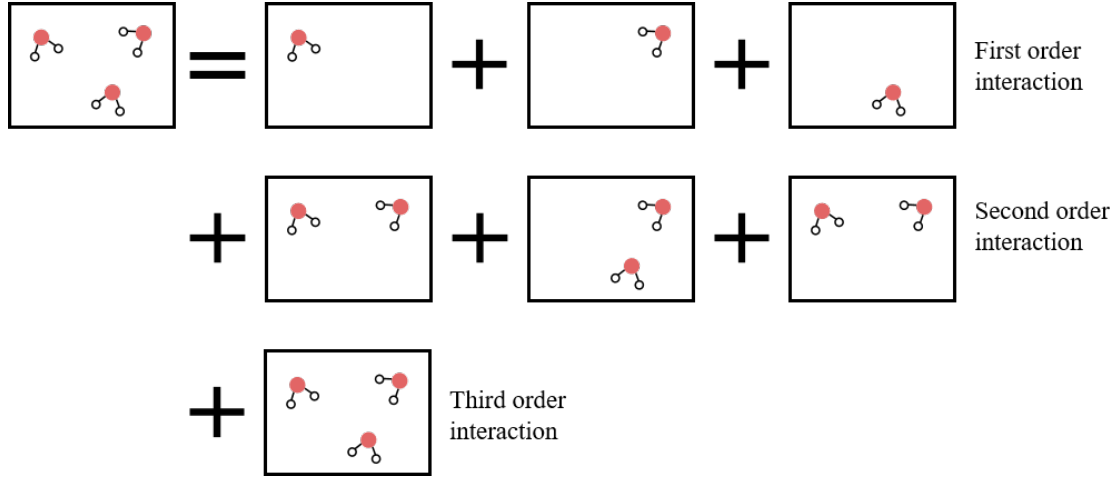
ematically, for a function  $y(\mathbf{x})$  with input vector  $\mathbf{x} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{n_x}$ , two extreme cases can be expressed by: (i)  $y(\mathbf{x}) = y_1(x_1) + y_2(x_2) + \dots + y_{n_x}(x_{n_x})$ , i.e., each dimension of the input has independent influence on the model output and (ii)  $y(\mathbf{x}) = y_{1,2,\dots,n_x}(x_1, x_2, \dots, x_{n_x})$ , i.e., only the interactions between all the input dimensions impact the the model output. In general, the model function  $y(\mathbf{x})$  can be decomposed/expanded as the sum of the contributions from all orders of interactions between different input dimensions, and this concept is similar to the well-known high-dimensional model representation (HDMR) (Sobol' 1993; Li et al. 2001; Sobol' 2003). The decomposition can be expressed by

$$y(\mathbf{x}) = \sum_{i=1}^{n_x} y_i(x_i) + \sum_{i=1}^{n_x} \sum_{j=i+1}^{n_x} y_{i,j}(x_i, x_j) + \dots + y_{1,2,\dots,n_x}(x_1, x_2, \dots, x_{n_x}) \quad (3.2)$$

where  $y_i(x_i)$  represents the independent contribution from the  $i$ th input dimension,  $y_{i,j}(x_i, x_j)$  represents the cooperative contribution from the interaction between the  $i$ th and  $j$ th input dimension, and  $y_{1,2,\dots,n_x}(x_1, x_2, \dots, x_{n_x})$  represents cooperative contribution from the interaction between all input dimensions.

The function decomposition described by Eq. (3.2) is now based on mathematical principles, but can be physically interpretable for many engineering problems. A good example is the classical many-body interaction problems in the field of quantum mechanics and molecular dynamics. For these problems, the system model output (i.e., the total potential) can be decomposed to the sum of the output from sub-systems, represented by many-body terms (i.e., individual terms and interaction terms) (Yao et al. 2017; Zhang et al. 2020a). Figure 3.2 shows an example of this additivity property for a three-body example. Moreover, from the physical model perspective, the expansion usually converges very fast (i.e., the high-order of interactions do not contribute too much) and thus can be

approximated by the series truncated at some finite orders. In this chapter, the additive feature of the model output shown in Eq. (3.2) is named as additivity, and is regarded as a physical constraint when it has physical interpretability.



**Figure 3.2:** Illustration of additivity property of a three-body example.

### 3.3 Physics-constrained Gaussian Process Model through Kernel Design

Consider a high-fidelity model  $y(\mathbf{x})$ , which demonstrates the aforementioned invariance, symmetry, and additivity properties, is expensive to evaluate. A Gaussian process model is developed to replace the high-fidelity model and improve the computational efficiency. In order to incorporate the physical constraints, we propose a physics-constrained Gaussian process model through kernel design. An invariant kernel is first developed based on the invariance and symmetry information on the high-fidelity model and a base kernel (e.g., commonly used kernel functions). After that, an additive kernel is designed to incorporate the additive characteristics of the high-fidelity model. Finally, the additive kernel is combined with the invariant kernel, and the integrated kernel can explicitly consider the invariance, symmetry, and additivity features, and be used to build Gaussian process models with high interpretability and generalization ability to predict the model response.

### 3.3.1 Basics of kernel function

As discussed in Chapter 2, kernel is an crucial ingredient of a Gaussian process model. Consider two input locations  $\mathbf{x}$ ,  $\mathbf{x}'$  and their corresponding model outputs  $y(\mathbf{x})$ ,  $y(\mathbf{x}')$ , a kernel  $k(\cdot)$  is used to measure the covariance (representing similarity or distance) between the model outputs at two input locations, written as

$$\text{Cov}(y(\mathbf{x}), y(\mathbf{x}')) = k(\mathbf{x}, \mathbf{x}') \quad (3.3)$$

The kernel function in the Gaussian process model typically assumes that the covariance between two model outputs decays smoothly as the distance between their inputs increases (Lophaven et al. 2002; Rasmussen 2004), i.e., if  $\mathbf{x}$  are close to  $\mathbf{x}'$ ,  $y(\mathbf{x})$  will be similar to  $y(\mathbf{x}')$ . Such kernel functions lie in the category of stationary kernel function, which is a function of the distance between inputs  $\mathbf{x} - \mathbf{x}'$ . Commonly used stationary kernel functions include the squared exponential kernel (shown in Eq. (2.3)), Matérn kernels (shown in Eq. (2.4)~(2.5)), rational quadratic kernel (shown in Eq. (2.6)), and more details about these kernel function can be found in Rasmussen (2004). Note that the distances between inputs are calculated in the Euclidean space in this chapter. A kernel function have several hyperparameters, which have significant influence on the established Gaussian process model. Take the squared exponential kernel in Eq. (2.3) as an example,  $\sigma^2$  is the variance which tunes the amplitude of the model and  $\boldsymbol{\theta} = [\theta_1, \dots, \theta_{n_x}]$  are the length-scales which control the wiggleness of the model.

In order to obtain desired prediction performance, selecting an expressive valid kernel function is especially important, since the kernel determines the prior function assumed for the Gaussian process model. Therefore, here we propose to impose the already known physical constraints (in-

cluding invariance, symmetry, and additivity described in Section 3.2) into the kernel, and in this way a more “informative prior” can be provided for the Gaussian process model construction. In the end, the established Gaussian process model is able to exhibit invariance, symmetry, and additivity features.

### 3.3.2 Invariant kernel

First, in order to build a Gaussian process model that is invariant under the operations in  $\mathcal{T}$  (see Eq. (3.1)), the proposed invariant kernel function should be invariant under the same operations. Ginsbourger et al. (2013) have shown that a Gaussian process model is invariant to the operations in the group  $\mathcal{T}$  if and only if  $k$  is argument-wise invariant to these operations, i.e., satisfying the following property:

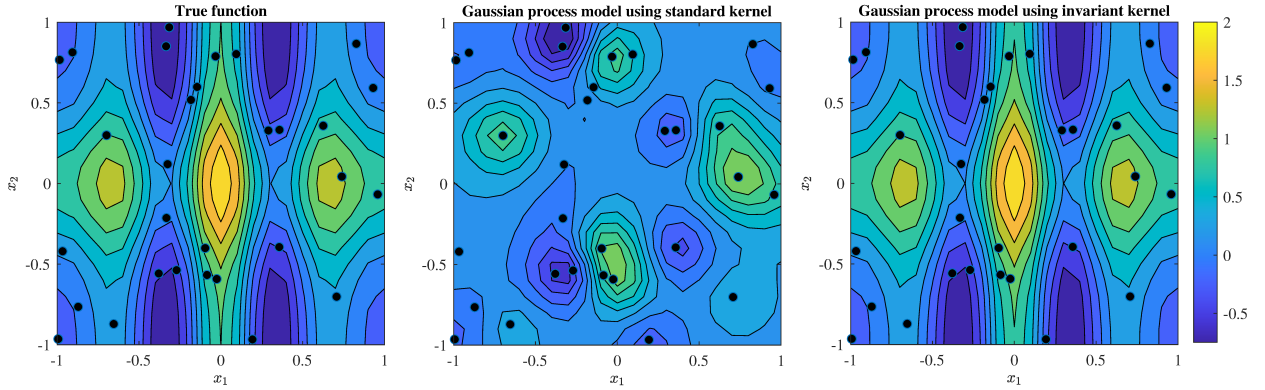
$$k(\mathbf{x}, \mathbf{x}') = k(t(\mathbf{x}), t'(\mathbf{x}')) \quad \forall \mathbf{x}, \mathbf{x}' \in \mathcal{X} \quad \forall t, t' \in \mathcal{T} \quad (3.4)$$

This kernel will be established by double summing a base kernel (i.e., commonly used kernels such as squared exponential kernel) over the orbits of the inputs, where the orbit of  $\mathbf{x}$  is the set of all transformed inputs obtained by applying each possible operation in  $\mathcal{T}$  to  $\mathbf{x}$  and can be represented by the set  $\mathcal{A}(\mathbf{x}) = \{t(\mathbf{x}); t \in \mathcal{T}\}$ . Formally, the invariant kernel is given by

$$k(\mathbf{x}, \mathbf{x}') = \sum_{\mathbf{x} \in \mathcal{A}(\mathbf{x})} \sum_{\mathbf{x}' \in \mathcal{A}(\mathbf{x}')} k_{base}(\mathbf{x}, \mathbf{x}') \quad (3.5)$$

To this end, the Gaussian process model constructed using the invariant kernel function in Eq. (3.5) will be capable of integrating our prior knowledge (i.e., invariance and symmetry information) about the physical characteristics of the problem. In order to show the performance of

the proposed invariant kernel, it is applied to build a Gaussian process model for a toy problem  $y = e^{-2x_1^2} \times \cos(9x_1) + e^{-5x_2^2}$  where  $x_1, x_2 \in [-1, 1]$ . Figure 3.3(a) shows the contour map of the true response, and as can be observed, the response is symmetric with respect to  $x_1$ -axis and  $x_2$ -axis. For comparison purpose, a Gaussian process model using a standard kernel is also constructed, and the predicted response is illustrated in Figure 3.3(b). Figure 3.3(c) shows the predicted response by the Gaussian process model using the invariant kernel. The comparison between the true response and the predictions from two Gaussian process models indicates that using the invariant kernel can significantly improve the prediction performance and generalization ability of the Gaussian process model when the high-fidelity model exhibits symmetry/invariance features.



**Figure 3.3:** Comparison between true response and predictions by Gaussian process models using standard kernel and invariant kernel (black dots correspond to the training samples).

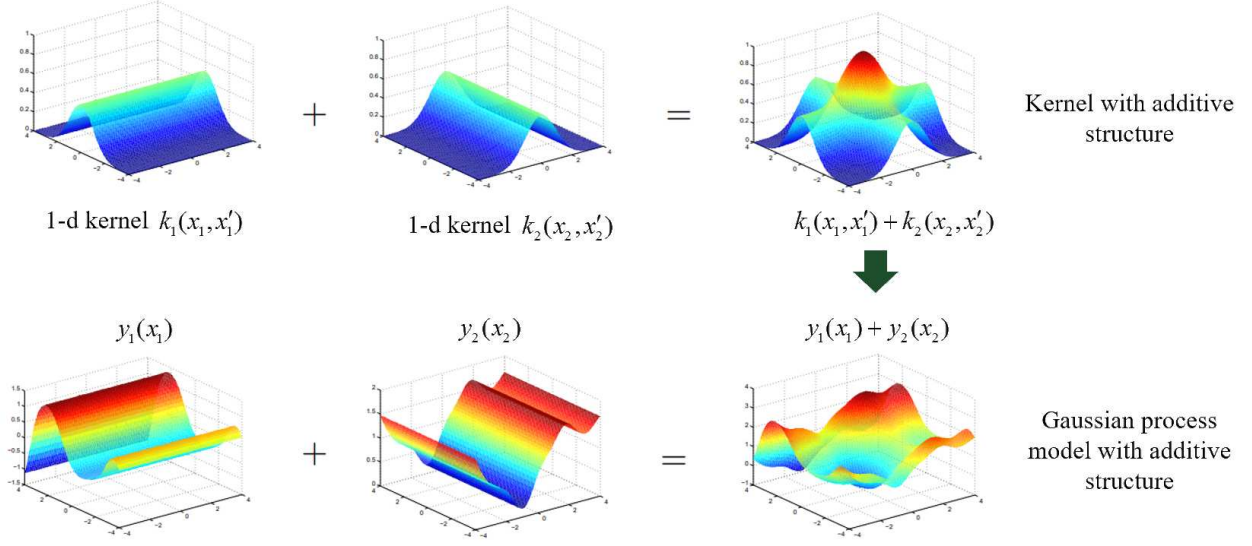
### 3.3.3 Additive kernel

Second, we embed our knowledge about the additive features of the high-fidelity model output. To build a Gaussian process model with the additive features, a new class of kernel is defined. It has been proved that the additivity of the Gaussian process model can be expressed by the additivity of the corresponding kernels (Durrande et al. 2012; Duvenaud 2014). Take the two-dimensional

model input case considering the first order of interaction as an example (see Figure 3.4), assume we apply one-dimensional kernel for each input dimension, and by summing them a new kernel can be defined. Then a Gaussian process model is constructed using the new kernel, and as a result the model will exhibit the same additive structure across two dimensions. Based on this property, the Gaussian process model in our problem is specified by a kernel which can be decomposed into a series of sub-kernels with each sub-kernel corresponding to the contribution from the interaction between subsets of the inputs. Such kernel is defined as the additive kernel, and it has the following expression

$$\begin{aligned}
k_{add}(\mathbf{x}, \mathbf{x}') &= \sigma_1^2 \sum_{i_1=1}^{n_x} k_{i_1}(x_{i_1}, x'_{i_1}) + \sigma_2^2 \sum_{i_1=1}^{n_x} \sum_{i_2=i_1+1}^{n_x} k_{i_1}(x_{i_1}, x'_{i_1}) k_{i_2}(x_{i_2}, x'_{i_2}) + \cdots \\
&\quad + \sigma_j^2 \sum_{1 \leq i_1 < \cdots < i_j \leq n_x} \prod_{d=1}^j k_{i_d}(x_{i_d}, x'_{i_d}) + \cdots
\end{aligned} \tag{3.6}$$

where  $\{k_{i_d}(x_{i_d}, x'_{i_d}), d = 1, 2, \dots, n_x\}$  is a base kernel operating on the  $i_d$ th dimension of the input  $\mathbf{x}$  and  $\mathbf{x}'$ , and  $\sigma_j^2$  is the variance assigned to the  $j$ th interaction term. In order to specify the additive kernel, one needs to select a base kernel function first, and then optimize the hyperparameters including the length-scale of the base kernel and the variance  $\sigma_j^2$  in each order based on the given training data. Note that different variance values can be specified for each sub-kernel of the additive kernel which also helps the Gaussian process model control the variance assigned to each order of interaction of the input.



**Figure 3.4:** Additive structure in kernel and Gaussian process model (adapted from Duvenaud (2014)).

### 3.3.4 Combined invariant and additive kernel

Finally, the additive kernel in Eq. (3.6) is combined with the invariant kernel in Eq. (3.5), and more specifically, the base kernel of the invariant kernel takes the additive kernel, i.e.,

$$k(\mathbf{x}, \mathbf{x}') = \sum_{\mathbf{x} \in \mathcal{A}(\mathbf{x})} \sum_{\mathbf{x}' \in \mathcal{A}(\mathbf{x}')} k_{add}(\mathbf{x}, \mathbf{x}') \quad (3.7)$$

The integrated kernel is then used to develop Gaussian process models for high-fidelity model which demonstrates invariance, symmetry, and additivity features. It is noteworthy to point out that the computational cost of evaluating the proposed physics-constrained kernel in Eq. (3.7) is expensive when the dimension of the input is large. However, if the decomposition converges fast, we can truncate the interaction series to a low order and this will help reduce the computational effort in building Gaussian process model. This assumption naturally suggests that, when the computational resources are limited, one can limit the maximum considered order of interaction for the additive

kernel without significantly impacting the prediction accuracy. In the end, the computational cost of evaluating Eq. (3.7) can be significantly reduced.

### **3.4 Illustrative Example: Hydrodynamic Characteristics of WECs in an Array**

To demonstrate the performance of the proposed algorithm, it is applied to predict hydrodynamic characteristics of wave energy converters (WECs) in an array.

#### *3.4.1 Motivation*

According to the U.S. Department of Energy, the estimated total renewable wave energy resource in the U.S. has the potential to power over 100 million homes each year (U.S. Department of Energy 2016). The deployment of WECs in large-scale arrays, also known as wave farms, offers great prospects for harnessing such renewable wave energy and facilitating electrical power transmission. Unlike the single-WEC configuration, a wave farm typically involves complex physical laws (i.e., WECs interact with scattered and radiated waves) and the hydrodynamic interactions between neighbouring WECs could have a significant impact on the total power generation of the wave farm. By arranging the positions of WECs in array properly, the total power generation can be greater than the power generated by the same number of isolated devices (Budal et al. 1977; Borgarino et al. 2012). In order to improve the efficiency of wave farms and achieve the maximum power generation, the layout of WECs needs to be carefully designed so that the hydrodynamic interactions can be positively exploited. Therefore, hydrodynamic modeling has drawn extensive attention in this field.

Most previous research on modeling hydrodynamic interaction has been conducted based on linear potential flow theory (Götteman 2020), and various numerical methods have been proposed within this context for calculating hydrodynamic interactions within arrays of WECs (McIver 2002).

The main methods can be classified into two classes : (i) analytical or semi-analytical methods, such as point-absorber approximation (Falnes 1980), plane-wave approximation (McIver 1994), and multi-scattering method (Twersky 1952); and (ii) numerical methods, such as boundary element method and finite element method. However, if high accuracy is desired, the computational efforts may still be burdensome, especially when the number of WECs is large and higher-fidelity models are used. Moreover, additional challenges also arise when dealing with problems such as uncertainty quantification and design optimization which typically require a large number of model evaluations.

To address the above computational challenges, this chapter proposes a Gaussian process model based approach for predicting hydrodynamic interactions between WECs. To the author's best knowledge, few research has investigated the prediction of hydrodynamic interactions using surrogate models. The main difficulties are two-folds. First, standard way of building surrogate model fails to incorporate available prior physical knowledge about the problem or input-output relationship. The direct surrogate modeling option is to take the layout of a WEC array (i.e., represented by a vector consisting of the coordinates of the WECs) as the model input, and the hydrodynamic characteristics as the model output. However, the hydrodynamic characteristics could be permutation invariant with respect to the ordering of the WECs given an array layout, or be symmetric with respect to the axes. In addition, the WEC interaction problem is similar to the classical many-body interaction problem (Zhang et al. 2020a), and the hydrodynamic characteristics can demonstrate the additivity feature. Directly taking the coordinates as inputs and applying commonly used product kernels for training Gaussian process model cannot incorporate the above prior knowledge into the surrogate modeling. Instead, these features (i.e., invariance, symmetry, and additivity) can only be learned by a large number of training data if directly using the common way to train the surrogate

models. However, a lot of times obtaining many training data for the hydrodynamic characteristics is impractical due to the high computational cost of running expensive models (e.g., MS solver, or boundary element models). The limited training data typically will lead to lower prediction accuracy and generalization ability of the established surrogate model. Second, as the size of the WEC array increases, there are additional challenges in constructing a surrogate model stemming from the increased input dimension, which typically requires more training data to obtain desired prediction accuracy, i.e., suffering from curse of dimensionality. However, the computational effort to calculate the hydrodynamic characteristics typically increases significantly with the number of the WECs, which means obtaining training data is more costly. In the end, building surrogate models for prediction of hydrodynamic characteristics becomes a challenging task. To tackle these issues, the most related research is conducted by Zhang et al. (2020a). They proposed a surrogate model based approach to predict the hydrodynamic characteristics of multiple bodies through a hierarchical interaction decomposition method. The key idea is to decompose hydrodynamic characteristics into contributions from clusters with fewer bodies, and then separately build lower-order surrogate models for these components/clusters instead of the total response. This mitigates the curse of dimensionality of surrogate modeling. In addition, by carefully designing the model input according to the input-output relationship, these lower-order surrogate models can include the invariance and symmetry principles. Overall, this approach addresses some of the surrogate modeling challenges by reducing the size of the problem and selecting the proper inputs and outputs (which involves some relatively complex transformations on the model inputs/outputs), while no effort was placed on modifying the Gaussian process model itself. Also, more importantly, to apply the approach, information on the contributions from subsets of the multiple-body system is needed; however, such

information is not always available, and its availability (i.e., whether it can be calculated) depends on the (numerical) model used.

Alternatively, this section applies the proposed physics-constrained Gaussian process model to efficiently predict the hydrodynamic characteristics of WEC array with different layouts. The proposed physics-constrained Gaussian process model is able to address the challenges discussed above within the constructed Gaussian process model itself, and does not involve complex design or transformation of the model input or output. Moreover, it is general and can be employed for cases where only total response of the array is available from the adopted numerical model. In the current problem, the prior knowledge about the hydrodynamic characteristics (i.e., invariance, symmetry, and additivity) can be encoded as the physical constraints of the hydrodynamic model. These constraints can be enforced in the surrogate model by adding operations inside the surrogate model that gives rise to the desired features. In particular, an invariant kernel is first developed to incorporate the invariance and symmetry information and integrated into the Gaussian process model. This developed invariant kernel enables efficient incorporation of the prior knowledge in the Gaussian process model, and eliminates the need of handcrafting the input features. Second, taking advantage of the similarity to the classical many-body interaction problem (Zhang et al. 2020a), the hydrodynamic characteristics is decomposed to the sum of the outputs from sub-systems. An additive kernel is developed to embed the additive feature in the Gaussian process model, and the kernel exhibits a similar additive feature to the hydrodynamic characteristics. Finally, the proposed invariant kernel is combined with the additive kernel to establish an integrated kernel, which is used to construct more informative Gaussian process models that explicitly include the known physical constraints on the hydrodynamic characteristics.

### 3.4.2 Problem formulation: hydrodynamic interaction between WECs

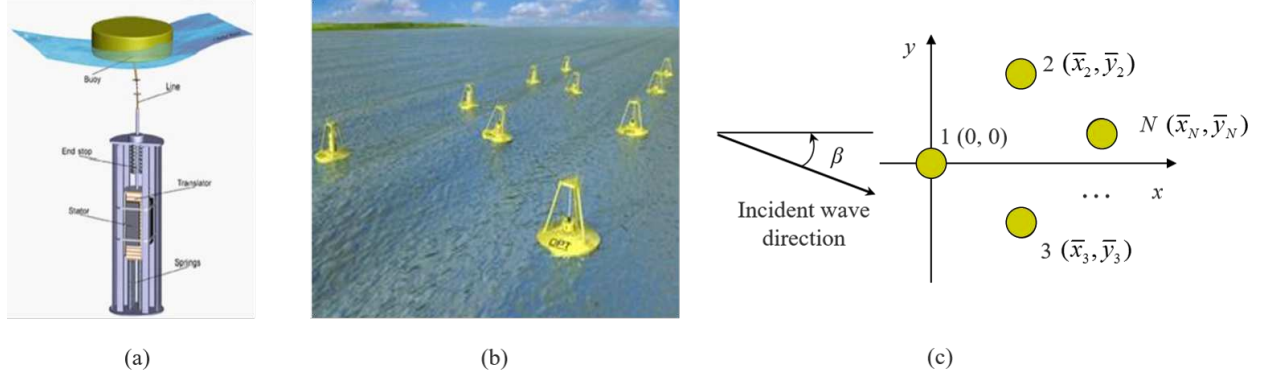
#### Representation of WEC array layout

Assume the considered WEC array has  $N$  identical WECs floating in the water, which are able to oscillate in  $M$  modes of motion. Among different configurations for WECs, this chapter mainly focuses cylindrical heave converters, which are heaving floating resonant buoys connected to a power take-off (PTO) moored to the seafloor. A typical cylindrical heave converter is shown in Figure 3.5(a). Therefore, the number of oscillating modes  $M$  is reduced to 1 here.

A two-dimensional Cartesian coordinate system is defined to express the locations of the WECs. The WECs in the array are under the incident wave propagating along the positive  $x$ -axis, and the corresponding incident angle is an arbitrary number  $\beta$ . Figure 3.5(c) shows an example layout of the array with  $N$  WECs, and note that throughout the paper, an array of WEC is numbered as in the figure. Without loss of generality, the leftmost buoy of the array is assumed to be located at the origin of the coordinate system. The layout of the array can then be characterized by the locations of the remaining  $N - 1$  buoys, i.e.,  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \dots, \bar{x}_N, \bar{y}_2, \bar{y}_3, \dots, \bar{y}_N] \in \mathcal{X} \subset \mathbb{R}^{2(N-1)}$ , where the pair  $(\bar{x}_i, \bar{y}_i)$  represents the coordinates of the  $i$ th WEC and  $\mathcal{X}$  denotes the admissible layout space. For convenience, the layouts of the WEC arrays are transformed into equivalent layouts with all buoys on the right half plane with  $\beta = 0$  beforehand, which can be established by rotating the coordinate system and adjusting the origin.

#### Diffraction and radiation problem

Following the assumption of classical hydrodynamics, i.e., small displacements, inviscid and incompressible fluid and irrotational flow, the fluid motion can be described by a velocity potential function based on linear potential theory of waves (Mavrakos and Koumoutsakos 1987). By further



**Figure 3.5:** (a) Wave energy converter (WEC) (adapted from (Prakash et al. 2016)), (b) wave farm (adapted from (Drew et al. 2009)), and (c) Cartesian representation of an array of  $N$  WECs.

assuming that all motions are time-harmonic with angular frequency  $\omega$ , we can extract the time-dependence of the velocity potential (Mavrakos and McIver 1997), and the time-dependent velocity potential function  $\Phi$  is then written as

$$\Phi = \text{Re} [\phi \cdot e^{-i\omega t}] \quad (3.8)$$

where  $i = \sqrt{-1}$ , and  $\text{Re}[\cdot]$  indicates that the real part is to be taken.  $\omega$  is the wave frequency, and should satisfy the dispersion equation  $\omega^2 = gk \tanh(kh)$ , where  $g$  is the gravity acceleration,  $k$  is the wave number and  $h$  is the water depth. Due to the linear potential theory of waves, the complex-valued velocity  $\phi$  can be expressed by the superposition of the incident wave potential, the scattered wave potential and the radiated wave potential (Mavrakos and McIver 1997), i.e.,

$$\phi = \phi_I + \phi_D + \sum_{p=1}^N U^{(p)} \phi_R^{(p)} \quad (3.9)$$

where  $\phi_I$  is the incident wave velocity potential,  $\phi_D$  describes the velocity potential of the diffracted wave field, and  $\phi_R^{(p)}$  is the velocity potential of the wave field induced by the oscillation of the  $p$ th body in the heave motion, and  $U^{(p)}$  is the corresponding velocity amplitude. The complex-

valued potential  $\phi$  must satisfy the Laplace equation within the entire fluid domain, expressed by  $\nabla^2\phi = 0$ . Also, several boundary conditions should be satisfied, and interested readers are referred to (Mavrakos and Koumoutsakos 1987; Mavrakos 1991) for detailed descriptions of the boundary conditions. By solving this boundary value problem, the diffracted wave potential  $\phi_D$  and the radiated wave potential  $\phi_R$  can be determined.

In this chapter, we are interested in the hydrodynamic characteristics of the array of WECs, and they can be computed based on the hydrodynamic pressure acting on the devices. More specifically, after the velocity potential in the entire fluid domain is solved, we can calculate the hydrodynamic pressure based on the linearized Bernoulli's equation (Hsu and Wu 1997). Ultimately, the hydrodynamic forces on the WECs may be calculated by integrating the hydrodynamic pressure over the submerged surface of the WECs (Mavrakos and McIver 1997). More specifically, the wave excitation force on the  $p$ th body in the direction of heave motion due to the incident and diffracted wave, and the hydrodynamic reaction force on the  $p$ th body in the direction of heave motion induced by the oscillation of the  $q$ th body in the direction of heave motion are written as

$$F^{(p)} = i\omega\rho \int \int_{S_p} (\phi_I + \phi_D)\nu dS \quad (3.10)$$

$$f^{(pq)} = i\omega\rho \int \int_{S_p} \left[ U^{(q)}\phi_R^{(q)} \right] \nu dS \quad (3.11)$$

where  $\rho$  is the fluid density,  $S_p$  is the submerged surface of the  $p$ th body, and  $\nu$  is the generalized normal component with respect to the body  $p$ . It should be noted that the wave reaction force,  $f^{(pq)}$ , can be reformulated to obtain the radiation-related hydrodynamic characteristics, i.e., the

added mass coefficient  $a^{(pq)}$ , and the added damping coefficient  $b^{(pq)}$ . The reformulation is written as the following equation:

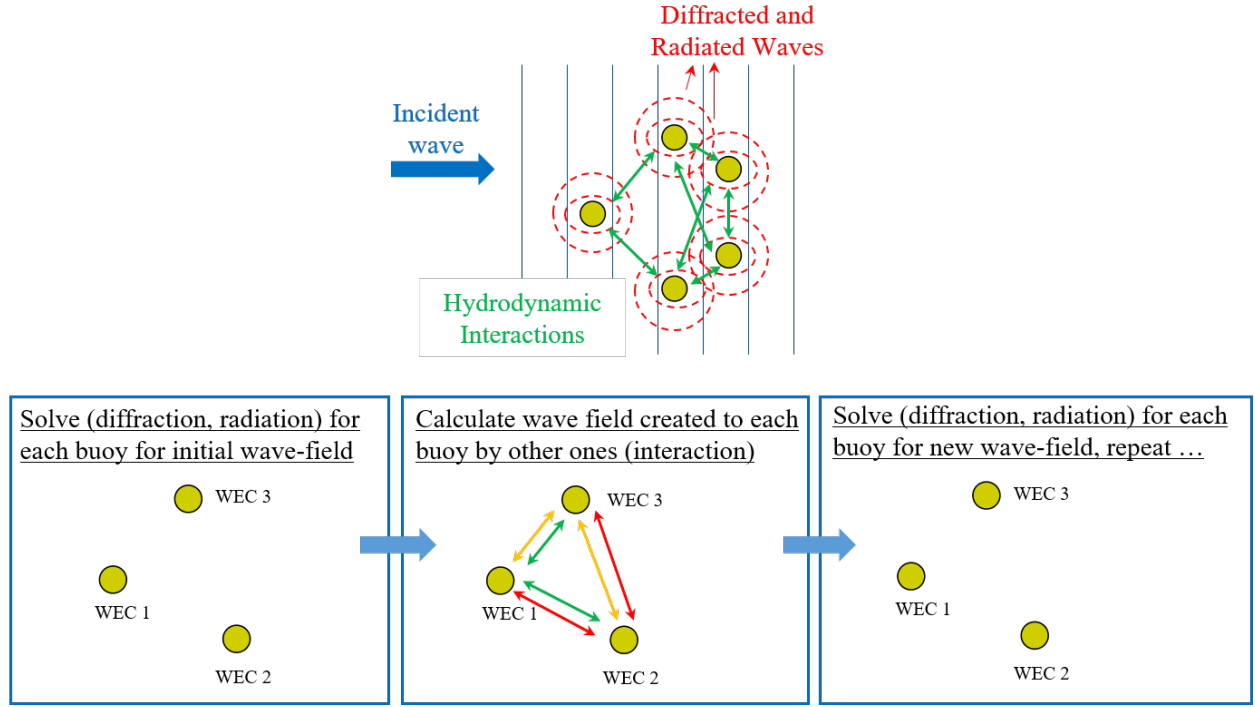
$$f^{(pq)} = i\omega U^{(q)} \left[ a^{(pq)} + i \frac{b^{(pq)}}{\omega} \right] \quad (3.12)$$

In the end, for a WEC array with  $N$  identical buoys (shown in Figure 3.5(c)), our main purpose is to calculate the hydrodynamic characteristics of the array, including wave excitation force  $F^{(p)}$  (corresponding to the diffraction problem), added mass coefficient  $a^{(pq)}$ , and added damping coefficient  $b^{(pq)}$  (corresponding to the radiation problem), where  $p, q = 1, \dots, N$ .

### Multiple scattering

In this chapter, the multiple-scattering (MS) method is adopted to obtain the solution of velocity potential of the wave field and calculate the hydrodynamic interaction within the array of WECs due to its versatility in achieving enhanced accuracy. This method relies on estimating the single-body hydrodynamic characteristics, and describes the interaction between the different bodies by superimposing the incident wave potential and various orders of propagating and evanescent modes that are scattered and radiated by all the devices in the array, so that an accurate representation of the total wave field around each device can be obtained. The concept of the MS method is demonstrated in Figure 3.6, and the steps are briefly reviewed here.

The method is applied separately to the diffraction and radiation problem. In the first step, this approach requires the solution of the single (isolated) body problem for each body  $p$  within the array. This solution can be established in frequency domain by solving the boundary value problem for each body through an eigenfunction expansion approach (Mavrakos and Koumoutsakos 1987; Mavrakos 1991; Kokkinowrachos et al. 1986). This entails an infinite series representation, using



**Figure 3.6:** Illustration of hydrodynamic interactions between WECs and concept of MS approach.

Bessel functions as eigenfunctions, truncated to a proper range. It is important to mention that for arrays of identical buoys, this solution is common for all the bodies. We refer the reader to (Scruggs et al. 2013) for more details on the derivation for the case of upright cylindrical bodies under heave oscillation.

In the second step, the interaction between the bodies is then addressed by considering the implications to the other bodies (e.g, body  $q$ ) of the diffracted field generated by the initial body  $p$ . The interaction between all bodies within the array, interchanging subsequently the roles of  $p$  and  $q$  as shown in Figure 3.6, needs to be considered this way. This new diffracted wave field represents a new excitation for each body  $p$  in response to which it generates a new diffracted field which influences the other bodies. This way the order of interaction is increased till the new wave field generated to all bodies  $q$  is small. The accuracy of the MS method is influenced by both the truncation order (i.e., for deriving the solution for specific body) as well as the interaction order

(i.e., for addressing the coupling between the bodies). Further details about the MS approach can be found in (Mavrakos 1991).

Ultimately, for each considered frequency  $\omega$ , the MS method provides the wave excitation forces exerted on each body by the incident wave, the added mass and damping coefficients exerted on body  $p$  in the direction of heave motion due to the oscillation of body  $q$  in the direction of heave motion. In this chapter, we denote  $\mathbf{F}(\omega) \in \mathbb{C}^N$  as the response vector relating the incident and diffracted wave to the wave excitation force on each body in the heave direction, and each element of the vector corresponds to  $F^{(p)}$ . Also, we denote  $\mathbf{a}(\omega) \in \mathbb{R}^{N \times N}$  and  $\mathbf{b}(\omega) \in \mathbb{R}^{N \times N}$  as the added mass and damping matrices, with the  $pq$ th element (i.e.,  $a^{(pq)}$  and  $b^{(pq)}$ ) of these matrices relating the heave oscillations of bodies  $p$  and  $q$ . It is noteworthy to point out that the computational cost of the MS solver can be expensive (i.e., modeling hydrodynamic interactions between WECs takes a lot of computational time), especially when the number of WECs in the array is large. Considering the significant computational effort in calculation of the hydrodynamic characteristics, typically, one can limit the maximum order of interaction or the maximum number of eigenfunction series to a relatively small number. Although such implementation gives rise to faster computation of the hydrodynamic characteristics, it trades off the accuracy for computational efficiency.

### 3.4.3 *Physics-constrained Gaussian process model for predicting hydrodynamic characteristics*

To alleviate the computational burden in calculating the hydrodynamic characteristics of the WEC array, Gaussian process surrogate modeling is used. In this problem, Gaussian process models are constructed to predict the relationship between the WEC layout and the hydrodynamic characteristics of each buoy. As discussed in Section 3.4.2, the layout is characterized by the locations of the buoys, i.e.,  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \dots, \bar{x}_N, \bar{y}_2, \bar{y}_3, \dots, \bar{y}_N]$ , where the pair  $(\bar{x}_i, \bar{y}_i)$  represents the

coordinates of the  $i$ th WEC. Therefore,  $\mathbf{x}$  is taken as the input of the Gaussian process models. For the model output, we are interested in wave excitation force vector  $\mathbf{F}$ , added mass coefficient matrix  $\mathbf{a}$ , and added damping coefficient matrix  $\mathbf{b}$  described in Section 3.4.2. In this chapter, we build separate Gaussian process models to predict the elements of  $\mathbf{F}$ ,  $\mathbf{a}$ , and  $\mathbf{b}$  (i.e.,  $F^{(p)}$ ,  $a^{(pq)}$ , and  $b^{(pq)}$ , where  $p, q = 1, \dots, N$ ), and the reason will be discussed later. Note that the wave excitation forces are complex values, and we separately predict the real and imaginary parts. With the model input and output selected, corresponding Gaussian process models can be constructed.

However, standard way of building Gaussian process model faces significant challenges. Directly taking coordinates as inputs cannot incorporate our prior knowledge about the input-output relationship, i.e., physical constraints including invariance, symmetry, and additivity (will be discussed in details later), into the surrogate modeling. These features can only be learned by a large number of training data, but a lot of times we cannot obtain so many training data due to the high cost of running the numerical model (e.g., the MS solver). As a result, the prediction accuracy and generalization ability of the Gaussian process model under limited number of training data may be significantly reduced. This is especially the case for arrays with large number of WECs, since the large input space typically requires more training data to obtain desired prediction accuracy. However, the computational time of calculating the hydrodynamic characteristics increases dramatically with the number of WECs, which means obtaining training data is more costly as well.

To address the challenges in constructing Gaussian process model to predict the hydrodynamic characteristics, this section applies the proposed physics-constrained Gaussian process model which explicitly embeds the available physical constraints into the surrogate modeling process and eliminates the need of preparing a large training data set. This section first summarizes the physical characteristics/constraints of the relationship between the WEC layout and the hydrodynamic char-

acteristics of WECs. Then, for different hydrodynamic characteristics, different model inputs are selected so that the physical constraints can be encoded appropriately and conveniently. Finally, in order to provide “informative prior” to the Gaussian process models, a physics-constrained Gaussian process model through designing specific kernels is established.

### Physical constraints

In the hydrodynamic interaction problem, we also have some prior knowledge about the invariance and symmetry features that the hydrodynamic characteristics exhibit, and they are essentially derived based on the physical laws of hydrodynamic interaction problem and based on the underlying principles of the numerical model (i.e., MS solver) for calculating the hydrodynamic interaction. For the WEC array with  $N$  buoys shown in Figure 3.5(c) represented by  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \dots, \bar{x}_N, \bar{y}_2, \bar{y}_3, \dots, \bar{y}_N]$ , the physical constraints are summarized as below:

(1) The wave excitation force of the buoy  $p$  (i.e.,  $\text{Re}[F^{(p)}]$  and  $\text{Im}[F^{(p)}]$  where  $p = 1 \dots N$ ) is permutation invariant to the ordering of the rest of the buoys in the array;

(2) The added mass and damping coefficient matrices (i.e.,  $\mathbf{a}$  and  $\mathbf{b}$ ) are symmetric, which means the added mass and damping coefficients of buoy  $p$  due to the heave oscillation of buoy  $q$  are equal to the added mass and damping coefficients of buoy  $q$  due to the heave oscillation of buoy  $p$ ;

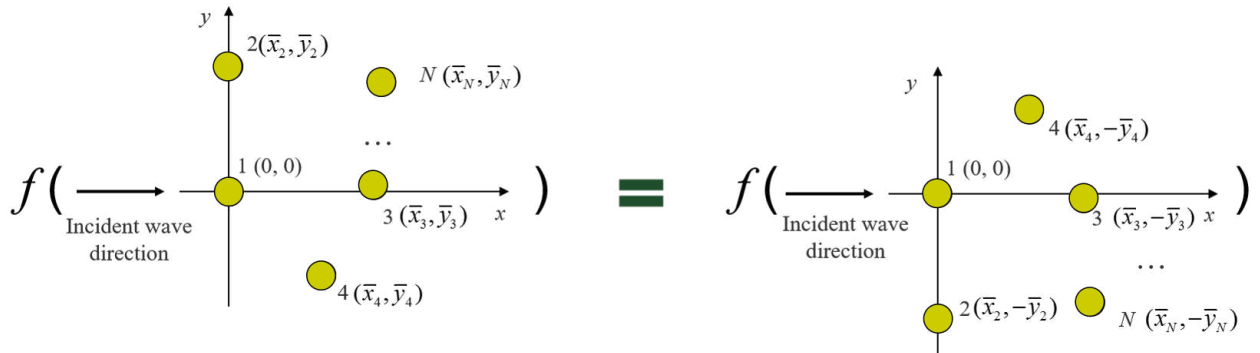
(3) The diagonal terms of added mass and damping coefficient matrices, i.e.,  $a^{(pp)}$  and  $b^{(pp)}$  where  $p = 1 \dots N$ , are permutation invariant with respect to the ordering of the rest of the buoys in the array. The off-diagonal terms of the matrices, i.e.,  $a^{(pq)}$  and  $b^{(pq)}$  where  $p, q = 1 \dots N$  and  $p \neq q$ , are permutation invariant with respect to the ordering of the rest of the buoys in the array (i.e., all the buoys in the array excluding buoy  $p$  and buoy  $q$  themselves);

(4) Since any layout of the WEC array can be transformed to a new one with incident wave angle  $\beta$  equal to zero, under zero incident angle, the hydrodynamic characteristics of all the buoys stay the same when the layout is reflected with respect to the  $x$ -axis. This symmetry property is illustrated in Figure 3.7.

The invariance and symmetry information summarized above can be encoded mathematically through Eq. (3.1). For example, suppose the array has three buoys and the coordinate vector used to characterize the layout is  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]$ . Then the operations for the wave excitation force of buoy 1 (i.e.,  $F^{(1)}$ ) can be expressed by

$$\begin{aligned} g_1([\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]) &= [\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3], & g_2([\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]) &= [\bar{x}_3, \bar{x}_2, \bar{y}_3, \bar{y}_2] \\ g_3([\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]) &= [\bar{x}_2, \bar{x}_3, -\bar{y}_2, -\bar{y}_3], & g_4([\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]) &= [\bar{x}_3, \bar{x}_2, -\bar{y}_3, -\bar{y}_2] \end{aligned} \quad (3.13)$$

where  $g_1(\cdot)$  represents “no operation”,  $g_2(\cdot)$  represents swapping WEC 2 and 3,  $g_3(\cdot)$  represents reflecting the WEC layout with respect to  $x$ -axis, and  $g_4(\cdot)$  represents swapping WEC 2 and 3 first and then reflecting the entire layout with respect to  $x$ -axis.

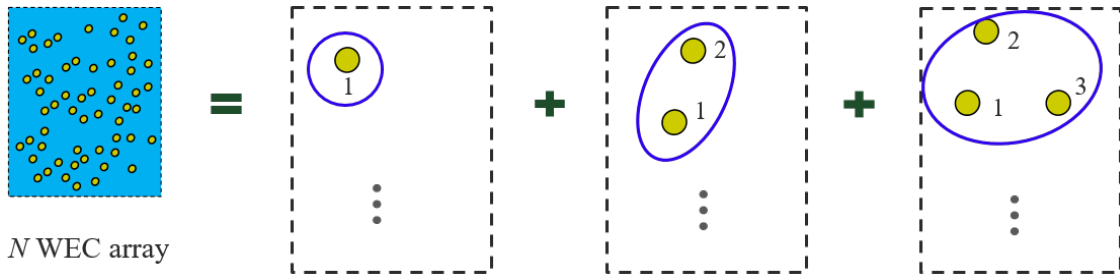


**Figure 3.7:** Illustration of symmetry property of hydrodynamic characteristics.

Second, the hydrodynamic interaction problem involves wave interactions between multiple floating bodies, and resembles the classical many-body interaction problem in many fields such as quantum mechanics and molecular dynamics (Gordon et al. 2011; Yao et al. 2017). In some systems such as molecule or crystal, the model output can be decomposed to the sum of the output from sub-systems, represented by many-body terms (i.e., individual terms and interaction terms), and approximated by the series truncated at some finite orders (Zhang et al. 2020a; Yao et al. 2017). Inspired by this idea, we assume that the quantities of interest in our problem also have such additive characteristics, i.e., the hydrodynamic characteristics can be decomposed into contributions from subsets of the buoys. Mathematically, we formulate such decomposition as

$$y(\mathbf{x}) = \sum_{i=2}^N y_i(\mathbf{x}_i) + \sum_{i=2}^N \sum_{j=i+1}^N y_{ij}(\mathbf{x}_i, \mathbf{x}_j) + \cdots + \sum_{2 \leq i < j < \cdots < l \leq N} y_{ij \dots l}(\mathbf{x}_i, \mathbf{x}_j, \dots, \mathbf{x}_l) + \cdots \quad (3.14)$$

where  $\mathbf{x}_i = (\bar{x}_i, \bar{y}_i)$  is the coordinate vector of the  $i$ th WEC.  $y$  represents the quantity of interest, i.e., any element of force vector  $\mathbf{F}$ , added mass coefficient matrix  $\mathbf{a}$ , or added damping coefficient matrix  $\mathbf{b}$ .  $y_i$  denotes the response attributed to the  $i$ th buoy,  $y_{ij}$  denotes the response attributed to the interaction between the  $i$ th and the  $j$ th buoys, and  $y_{ij \dots l}$  denotes the response attributed to the interaction between the  $i$ th, the  $j$ th, ..., and the  $l$ th buoys. The additivity feature is illustrated in Figure 3.8.



**Figure 3.8:** Illustration of decomposition of many-body systems (i.e., array of  $N$  WECs).

The invariance, symmetry, and additivity features discussed above are deemed the physical constraints of the hydrodynamic interaction model, and if these available prior knowledge can be incorporated when the Gaussian process models are constructed (e.g., encoded in the prior function assumed for the Gaussian process model), the required number of training data to reach good prediction accuracy is expected to be reduced. This reduction is especially important for building accurate Gaussian process models for arrays with relatively large number of WECs, where the computational cost to obtain the training data is typically quite expensive.

### **Selection of model input**

As summarized in the previous section, the permutation invariance features of wave excitation force, added mass and added damping coefficient may involve different buoys (i.e., corresponding to different model input), therefore we separately predict these hydrodynamic characteristics. For different quantities of interest, we need to do some simple transformation (e.g., shift the coordinate system by some distances) on the original model input  $\mathbf{x}$  so that the invariance and symmetry information can be appropriately and conveniently included in surrogate modeling.

First, we consider the model output  $y = \text{Re} [F^{(p)}]$ , or  $\text{Im} [F^{(p)}]$ , or  $a^{(pp)}$ , or  $b^{(pp)}$ , where  $p = 1, \dots, N$ . In order to include the permutation invariance of the model outputs with respect to the ordering of all buoys except buoy  $p$  in the Gaussian process model, we move the origin of the coordinate system to buoy  $p$  and obtain a new coordinate vector which will be used as the model input. Take the array with three WECs as an example, if we are interested in the added mass coefficient of WEC 2 due to the heave motion of itself (i.e.,  $a^{(22)}$ ), we first augment the coordinate vector  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \bar{y}_2, \bar{y}_3]$  to  $\mathbf{x}_{aug} = [0, \bar{x}_2, \bar{x}_3, 0, \bar{y}_2, \bar{y}_3]$  and then shift the coordinate system so that the second buoy is the new origin. Finally, the new coordinate vector used to characterize

the array layout becomes  $\mathbf{x}_T = [-\bar{x}_2, \bar{x}_3, -\bar{y}_2, \bar{y}_3]$  (i.e., represented by the positions of WEC 1 and 3). It is noteworthy to point out that the added mass and damping coefficients are invariant with respect to any translational shift of the coordinate system, and thus the above transformation on the model input does not change the model outputs. However, any translational shift of the coordinate system along the wave propagation direction (i.e.,  $x$  direction when  $\beta = 0$ ) can cause changes to the wave excitation forces. Therefore, in practice, when predicting the wave excitation force of buoy  $p$  where  $p = 2, \dots, N$ , we select the original coordinate vector as the model input (i.e.,  $\mathbf{x}_T = \mathbf{x}$ ), and release the constraint on the permutation invariance, i.e., switching the ordering of buoys except buoy 1 and buoy  $p$  does not change  $F^{(p)}$ .

Second, for the off-diagonal terms of added mass and damping coefficient matrices, i.e.,  $y = a^{(pq)}$  or  $b^{(pq)}$ , where  $p, q = 1, \dots, N$  and  $p \neq q$ , we also move the origin of the coordinate system to the body  $p$  so that the permutation invariance with respect to the ordering of the buoys except buoy  $p$  and buoy  $q$  can be considered conveniently, and the transformed coordinate vector is also represented by  $\mathbf{x}_T$ . It should be noted that in this case we also release the constraint on the permutation invariance since switching the ordering of buoy  $p$  and buoy  $q$  also does not change the value of  $a^{(pq)}$  or  $b^{(pq)}$ .

### Physics-constrained kernel

Once the model inputs are selected, we can establish the physics-constrained kernels based on Eq. (3.6)~(3.7). It is noteworthy to point out that the definition of the model input in this example is  $\mathbf{x} = [\bar{x}_2, \bar{x}_3, \dots, \bar{x}_N, \bar{y}_2, \bar{y}_3, \dots, \bar{y}_N]$ , where the pair  $(\bar{x}_i, \bar{y}_i)$  represents the coordinates of the  $i$ th WEC. When calculating the additive kernel in Eq. (3.6), the  $i_d$ th dimension of the input (i.e.,

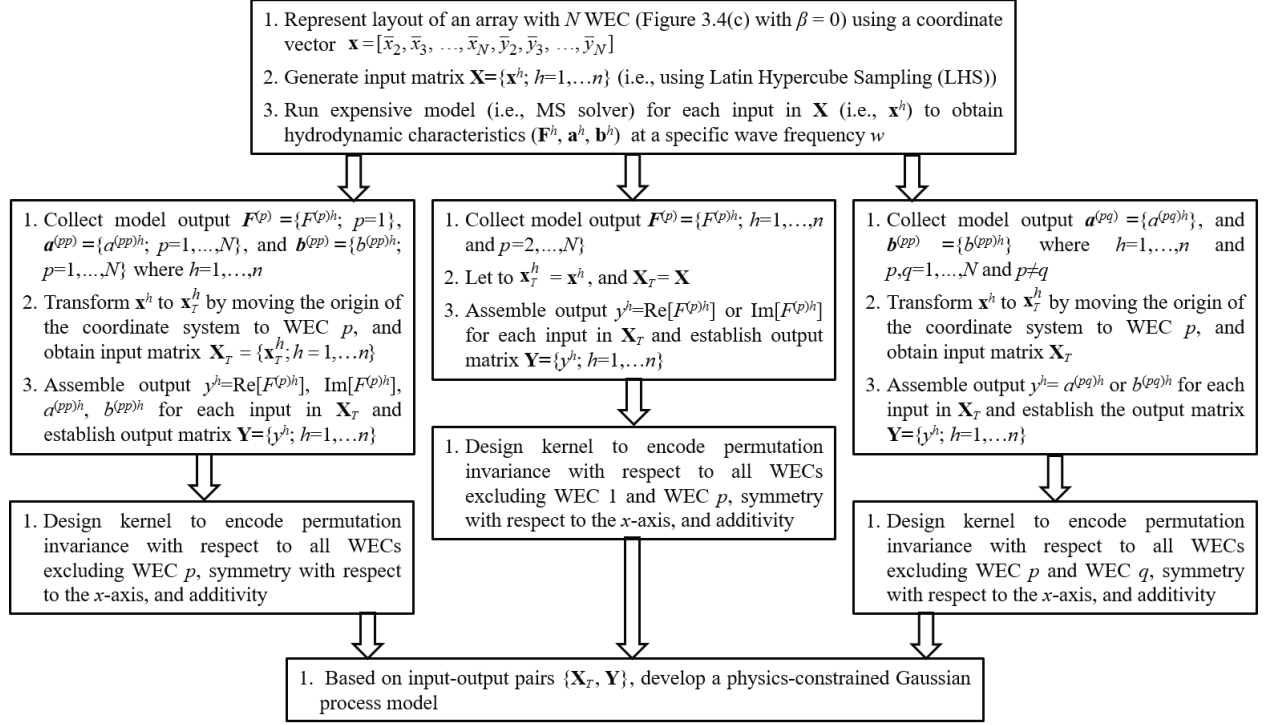
$x_{id}$  in Eq. (3.6)) should correspond to  $(\bar{x}_{id}, \bar{y}_{id})$  in this example and  $d$  takes the value from 2 to  $N$ .

The physics-constrained kernels are the used to establish the physics-constrain GP models.

### Overall algorithm

The flowchart of the proposed physics-constrained Gaussian process model for predicting hydrodynamic characteristics of WEC arrays is illustrated in the Figure 5.3, which shows the key steps. In the flowchart, the output of the Gaussian process model is scalar, i.e., build a separate Gaussian process model for each element of  $\mathbf{F}$ ,  $\mathbf{a}$  and  $\mathbf{b}$ . Therefore, for an array of  $N$  WECs, we need to train  $2N + 2(N^2 + N)/2 = N^2 + 3N$  Gaussian process models in total. However, in practice, we can train a model to predict some elements together if they depend on the same model input, e.g.,  $\text{Re}[F^{(11)}]$ ,  $\text{Im}[F^{(11)}]$ ,  $a^{(11)}$ , and  $b^{(11)}$ . In this way, the number of Gaussian process model needed to be trained is reduced to  $N + N - 1 + (N^2 - N)/2 = (N^2 + 3N)/2 - 1$  at most, which is less than half of the original required number of Gaussian process models.

For any new WEC array shown in Fig 3.5(c), we first transform the layout to a new one with incident angle  $\beta = 0$ , and the corresponding coordinate vector is denoted  $\mathbf{x}^0$ . For any model output of interest, we then transform  $\mathbf{x}^0$  to appropriate  $\mathbf{x}_T^0$  by simply moving the coordinate system according to Section 3.4.3. Then the trained Gaussian process models can be used to directly predict the elements of the hydrodynamic characteristics for the considered wave frequency  $\omega$ . Finally, the elements are assembled to establish the corresponding vector  $\mathbf{F}$  or matrices  $\mathbf{a}$  and  $\mathbf{b}$ . Note that the proposed algorithm focuses on the hydrodynamic characteristics given a specific wave frequency. In many cases, we are interested in the hydrodynamic characteristics under a range of wave frequencies, and the responses could correspond to high-dimensional model outputs. Instead of building Gaussian models for the hydrodynamic characteristics under each frequency which is



**Figure 3.9:** Flowchart of the proposed physics-constrained Gaussian process model for prediction of hydrodynamic characteristics.

computational demanding, we can apply dimension reduction technique to explore the correlations of the characteristics between different frequencies and build a low-dimensional representation for the original high-dimensional outputs. Then the Gaussian process model can be efficiently constructed for low-dimensional latent outputs. More details can be found in Section 2.6.

### 3.4.4 Implementation details

The array is in a rectangular domain with 127.5 m along the  $x$ -axis and 255 m along the  $y$ -axis, and the water depth is 60 m. The buoys in Jia et al. (2015), identical cylinders oscillating in heave direction only, are considered in this example. The radius for each buoy is  $r_b=3$  meters and its mass is  $m_b=1.8e5$  kg, corresponding to a draft of  $D_r=6.37$  m and period of oscillation of 5.06 s. The considered wave frequency is between 0.3 rad/s and 1.4 rad/s. The incident angle  $\beta$  is selected as 0 in this example, but generally it can take any values. In order to investigate the scalability of the

proposed algorithm, it is applied to arrays with different number of WECs, and here  $N$  in the range of  $3 \sim 10$  is considered.

To establish the training data set,  $n = 1000$  inputs are generated by Latin Hypercube Sampling (LHS) and the corresponding outputs are calculated by the hydrodynamic interaction model using MS approach. For the MS solver, the order of interaction is set as 5, whereas the eigenfunction series are truncated at 5 and 40 for the main fluid and the fluid below the cylinder, respectively, to ensure adequate accuracy. These values are selected through a convergence study of the model. In the proposed physics-constrained Gaussian process model, the kernel is the invariant kernel in Eq. (3.5) taking the additive kernel in Eq. (3.6) as base kernel, and for the base kernel of additive kernel, we select the commonly used Matérn 5/2 kernel. Since decomposition in Eq. (3.14) is expected to converge fast, we infer that only first several orders of interaction in Eq. (3.6) contribute much to the model response. To validate this assumption with the currently available computational resource, we used the additive kernel with full orders of interaction for the WEC arrays with 3 buoys and 4 buoys, and the contribution from each order is characterized by the assigned variance  $\sigma_j^2$ , where  $j = 1, 2$  for 3-WEC array and  $j = 1, 2, 3$  for 4-WEC array. After optimization of the hyperparameters, the results for all the diffraction and radiation related hydrodynamic characteristics show that for the 3-WEC array,  $\sigma_1^2 \approx 1$  while  $\sigma_2^2 \approx 0$ . Similarly, for the 4-WEC array,  $\sigma_1^2 \approx 1$  while  $\sigma_2^2, \sigma_3^2 \approx 0$ . Both cases indicate that only the first order of interaction between the buoys contribute most to the hydrodynamic coefficients predicted by the physics-constrained Gaussian process model. Therefore, for all the other cases, we limit the maximum considered order of interaction for the additive kernel to the first order and assume that this implementation will not significantly impacting the prediction accuracy. This significantly reduces the computational effort in evaluating the kernel functions. For comparison purpose, a standard Gaussian process model,

which is the same with the proposed model but uses the standard base kernel, is constructed to predict the same output. For the kernels in the proposed Gaussian process model and the standard Gaussian process model, we use the same length-scale value for all input dimensions. The reason of this selection here is that the computational effort in optimizing single length-scale is typically much less than optimizing multiple length-scales especially when the input dimensionality is high, and more importantly in this example using the same length-scale for all the input dimensions can already inform an accurate Gaussian process model.

To assess the accuracy of the Gaussian process models in predicting the hydrodynamic characteristics, validation metric are calculated over a testing set for each case. The test size is set as  $n_t = 1000$ , and the testing data is also generated by LHS. Here the coefficient of determination  $R^2$  is used as the validation metric,

$$R^2 = 1 - \frac{\sum_{i=1}^{n_t} (y^i - \hat{y}^i)^2}{\sum_{i=1}^{n_t} (y^i - \sum_{i=1}^{n_t} y^i / n_t)^2} \quad (3.15)$$

where  $\hat{y}^i$  is the prediction from the established Gaussian process model for the  $i$ th data in testing set. Large  $R^2$  values (e.g., closer to 1) indicate that the trained model has good accuracy.

#### 3.4.5 Prediction accuracy

To investigate the prediction accuracy of the proposed physics-constrained Gaussian process model, we calculate the coefficient of determination  $R^2$  for all cases (i.e., arrays with different numbers of WECs). Here we only list the results for arrays of 3, 5, 8 and 10 WECs, while the results for other cases show similar pattern. The prediction accuracy metrics are calculated separately for different elements of the hydrodynamic characteristic matrices  $\mathbf{F}$ ,  $\mathbf{a}$  and  $\mathbf{b}$ . For illustration purpose, the elements are classified into four groups: (1) real and imaginary parts of  $F^{(1)}$ , (2) real and

imaginary parts of  $F^{(p)}$  where  $p = 2, \dots, N$ , (3)  $a^{(pp)}$  and  $b^{(pp)}$  where  $p = 1, \dots, N$ , and (4)  $a^{(pq)}$  and  $b^{(pq)}$  where  $p, q = 1, \dots, N$  and  $p \neq q$ . The reason for such classification is that the prediction accuracies for the real and imaginary parts of  $\mathbf{F}$  are close, and also the prediction accuracies for  $\mathbf{a}$  and  $\mathbf{b}$  are close. Additionally, when modeling  $F^{(1)}$ ,  $a^{(pp)}$ , and  $b^{(pp)}$ , the strict invariance information is incorporated while only partial invariance information is considered when modeling  $F^{(p)}$ ,  $a^{(pq)}$ , and  $b^{(pq)}$ . For each group, the mean of the accuracy metrics is calculated, and to show the spread of prediction accuracy between different elements within a group, the minimum and maximum of the accuracy metrics are also calculated. Since the variation of wave frequency  $\omega$  may impact the prediction accuracy of the Gaussian process model, we show the results for three different cases:  $\omega = 0.3$  rad/s,  $\omega = 0.85$  rad/s, and  $\omega = 1.4$  rad/s. When  $\omega = 0.3$  rad/s, all the calculated accuracy metrics mentioned above are over 0.989 for each case (i.e., under different number of WECs). Similar observation can be made for  $\omega = 0.85$  rad/s: all the calculated  $R^2$  are over 0.980. This indicates excellent prediction accuracy of the physics-constrained Gaussian process model for the hydrodynamic characteristics when  $\omega = 0.3$  or  $0.85$  rad/s, and also there is little difference between all cases in terms of prediction accuracy. For this reason, the detailed prediction accuracy metrics are not shown for  $\omega = 0.3$  or  $0.85$  rad/s. Here we only report the detailed prediction accuracy metrics for  $\omega = 1.4$  rad/s, which are shown in Table 3.1.

Compared to  $\omega = 0.3$  or  $0.85$  rad/s, the prediction accuracy metrics are relatively lower when  $\omega = 1.4$  rad/s, especially for arrays of 8 or 10 WECs. For diffraction-related characteristics (i.e., elements of  $\mathbf{F}$ ), as can be observed from the table, the  $R^2$  for all cases are close to one, which reveals very good prediction accuracy of the proposed physics-constrained Gaussian process model. In each case, by comparing the mean, minimum and maximum of the  $R^2$  within a group of the hydrodynamic characteristics, we can find that the values are very close, which means there is

**Table 3.1:** Prediction accuracy metrics of physics-constrained Gaussian process model for  $\omega = 1.4$  rad/s

| # of WECs | $\text{Re}[F^{(1)}](\text{Im}[F^{(1)}])$ |              |              | $\text{Re}[F^{(p)}](\text{Im}[F^{(p)}])$ |              |              |
|-----------|------------------------------------------|--------------|--------------|------------------------------------------|--------------|--------------|
|           | mean                                     | min          | max          | mean                                     | min          | max          |
| 3         | 0.986(0.987)                             | 0.986(0.987) | 0.986(0.987) | 0.975(0.978)                             | 0.975(0.975) | 0.975(0.980) |
| 5         | 0.972(0.966)                             | 0.972(0.966) | 0.972(0.966) | 0.940(0.939)                             | 0.935(0.936) | 0.951(0.942) |
| 8         | 0.935(0.939)                             | 0.935(0.939) | 0.935(0.939) | 0.886(0.877)                             | 0.871(0.864) | 0.895(0.888) |
| 10        | 0.910(0.925)                             | 0.910(0.925) | 0.910(0.925) | 0.848(0.845)                             | 0.817(0.827) | 0.862(0.860) |
| # of WEC  | $a^{(pp)}(b^{(pp)})$                     |              |              | $a^{(pq)}(b^{(pq)})$                     |              |              |
|           | mean                                     | min          | max          | mean                                     | min          | max          |
| 3         | 0.931(0.920)                             | 0.899(0.885) | 0.978(0.976) | 0.956(0.958)                             | 0.934(0.939) | 0.967(0.970) |
| 5         | 0.857(0.847)                             | 0.805(0.803) | 0.950(0.943) | 0.866(0.880)                             | 0.819(0.845) | 0.927(0.924) |
| 8         | 0.784(0.767)                             | 0.727(0.729) | 0.906(0.875) | 0.750(0.758)                             | 0.669(0.672) | 0.864(0.858) |
| 10        | 0.755(0.731)                             | 0.719(0.681) | 0.868(0.879) | 0.674(0.687)                             | 0.536(0.594) | 0.817(0.815) |

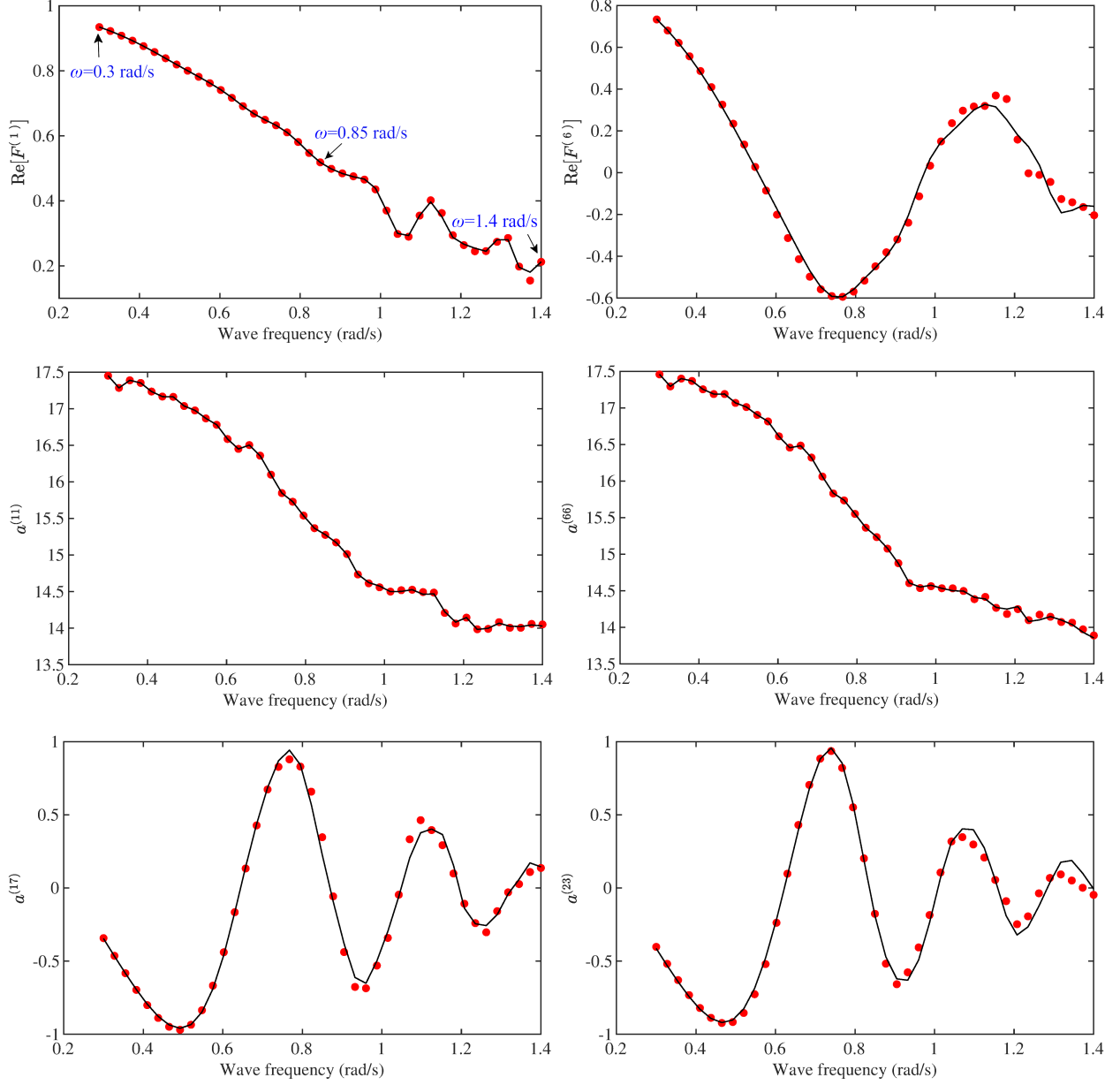
little variation in the prediction accuracy for different elements of  $\mathbf{F}$ . In addition, we can observe slightly lower prediction accuracy for  $F^{(p)}$  compared to  $F^{(1)}$  for all cases. This is because the physics-constrained Gaussian process model for  $F^{(1)}$  encodes all available permutation invariance information into the kernel, while the kernel for  $F^{(p)}$  incorporates only part of the the permutation invariance, i.e., switching the ordering of buoys except buoy 1 and buoy  $p$  does not change  $F^{(p)}$ . For radiation-related characteristics (i.e., elements of  $\mathbf{a}$  and  $\mathbf{b}$ ), excellent prediction accuracy can also be obtained for the arrays of 3 and 5 WECs, while for arrays of 8 and 10 WECs the prediction accuracy sees an obvious drop. This is expected since the dimensionality of the model input is proportional to the number of WECs in the array, and higher-dimensional input space requires more training data to inform an accurate Gaussian process model. Although we can also observe a drop in the prediction accuracy for the diffraction-related characteristics as the number of WECs increases, the magnitude of decrease is not as large as that for radiation-related characteristics. The reason behind might be that radiation-related characteristics show stronger non-linearity with respect to the model input than the diffraction-related characteristics. It is noteworthy to point out that the mean, minimum and maximum values of the  $R^2$  for radiation-related hydrodynamic characteristics

are relatively different, especially for arrays with 8 and 10 buoys. This means that the prediction accuracy for different elements of  $a^{(pp)}/b^{(pp)}$  and  $a^{(pq)}/b^{(pq)}$  show relatively large variability when the number of WECs is large in the array.

Based on the above results, we can find that the wave frequency has large impact the prediction accuracy of the proposed physics-constrained Gaussian process model. To investigate the impact of the wave frequency, Figure 3.10 shows the variation of hydrodynamic characteristics as the wave frequency changes for an array of 10 WECs from the testing set. Note that we only pick some representatives from all the elements and show them in the figures. The other elements of the hydrodynamic characteristics show similar trend in terms of the prediction accuracy based on the results in Table 3.1, and thus are not illustrated here. From the figure, we can observe that when the wave frequency is low (e.g., less than 1 rad/s), the prediction for all hydrodynamic characteristics demonstrate a good agreement with the ground truth values. However, as the wave frequency increases, it is sometimes more challenging to predict the hydrodynamic characteristics accurately, especially when the response reaches its maximum or minimum. This might be because the input-output relationship is more nonlinear when the wave frequency becomes higher. Therefore, when predicting the hydrodynamic characteristics under relatively higher wave frequencies, more training data can be used to increase the prediction accuracy. For this particular array, the predictions for  $\text{Re}[F^{(6)}]$ ,  $a^{(17)}$  and  $a^{(23)}$  have relatively large errors when the wave frequency is high, which might still be due to the kernel design (i.e., partial invariance is used).

#### 3.4.6 Performance of physics-constrained kernel

In this section, the performance of the kernel in the proposed physics-constrained Gaussian process model is investigated. For comparison purpose, a standard Gaussian process model with



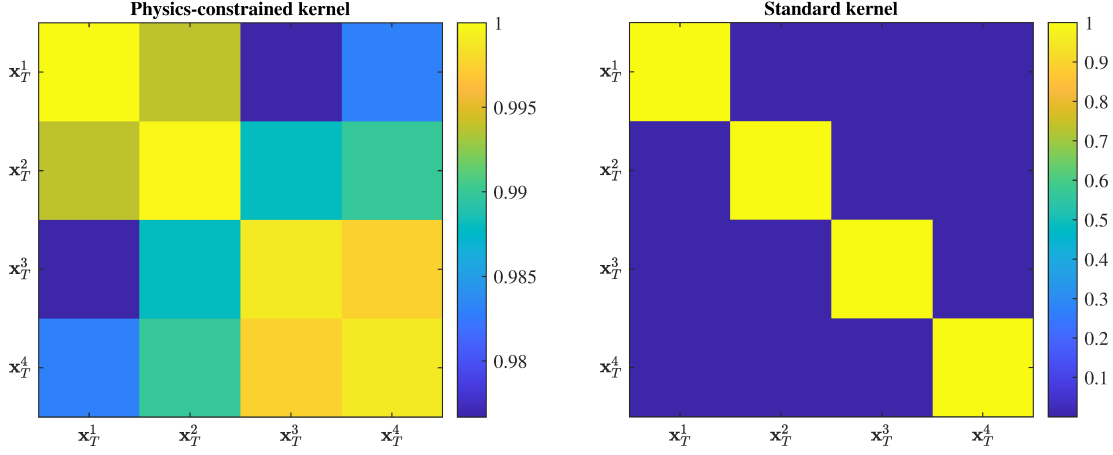
**Figure 3.10:** Comparison of the hydrodynamic characteristics calculated by numerical model (black lines) and predicted by physics-constrained Gaussian process model (red dots).

the Matérn 5/2 kernel is also constructed with the same training data set and validated using the same testing data as the physics-constrained Gaussian process model. For the convenience, the kernel in the physics-constrained Gaussian process model and the standard Gaussian process models are referred as physics-constrained kernel and standard kernel respectively in this section. The main

purpose of this section is to compare the performance of the physics-constrained and the standard Gaussian process model, and thus the prediction accuracy for a single case is less of a concern here. As discussed in the previous section, the prediction task is less challenging when wave frequency is small, therefore we select  $\omega = 0.3$  rad/s for demonstrating the results in this section.

The kernel function in a Gaussian process model specifies the covariance between the model response at two input locations, which implicitly describes the distance between two input locations. In order to visualize the capability of the physics-constrained kernel and the standard kernel in describing the similarity of the model response between the input points, Figure 3.11 illustrates the covariance function values evaluated at a data set  $\mathbf{X}_T$ , i.e., covariance matrix  $k(\mathbf{X}_T, \mathbf{X}_T)$ . The data set contains four arrays of 3 WECs, characterized by  $\mathbf{X}_T = \{\mathbf{x}_T^1, \mathbf{x}_T^2, \mathbf{x}_T^3, \mathbf{x}_T^4\}$ , where  $\mathbf{x}_T^1 = [123.5 \ 93.2 \ -59.2 \ -116.4]$ ,  $\mathbf{x}_T^2 = [93.2 \ 123.5 \ -116.4 \ -59.2]$ ,  $\mathbf{x}_T^3 = [123.5 \ 93.2 \ 59.2 \ 116.4]$ , and  $\mathbf{x}_T^4 = [93.2 \ 123.5 \ 116.4 \ 59.2]$ . Among these four inputs,  $\mathbf{x}_T^1$  is selected from the training set, and the other three are transformed from  $\mathbf{x}_T^1$  based on the invariance and symmetry properties, as shown in Eq. (5.5). For illustration purpose, we added some Gaussian noises to the locations of WECs. Note that here we focus on the physics-constrained kernel used for  $F^{(1)}$ ,  $a^{(pp)}$  and  $b^{(pp)}$ , which includes the complete permutation invariance and symmetry information available. In Figure 3.11, the value of each grid describes the covariance of the model response between two inputs, and for comparison purpose, the values are normalized so that the diagonal terms of the covariance matrices are equal to one. As a result, the covariance values close to one means that two input points are close to each other and their responses are also highly similar. As can be observed from the figure, the covariance calculated by the standard kernel between  $\mathbf{x}_T^i$  and  $\mathbf{x}_T^j$  ( $i \neq j$ ) is nearly zero, while the corresponding covariance calculated by the physics-constrained kernel is close to one. However, as we mentioned earlier,  $\mathbf{x}_T^i$  and  $\mathbf{x}_T^j$  ( $i \neq j$ ) should give us similar model responses

due to the invariance and symmetry features of the input-output relationship. This indicates that the physics-constrained kernel has the capability of encoding the invariance and symmetry features into the Gaussian process model, while the standard kernel is not able to capture such properties.

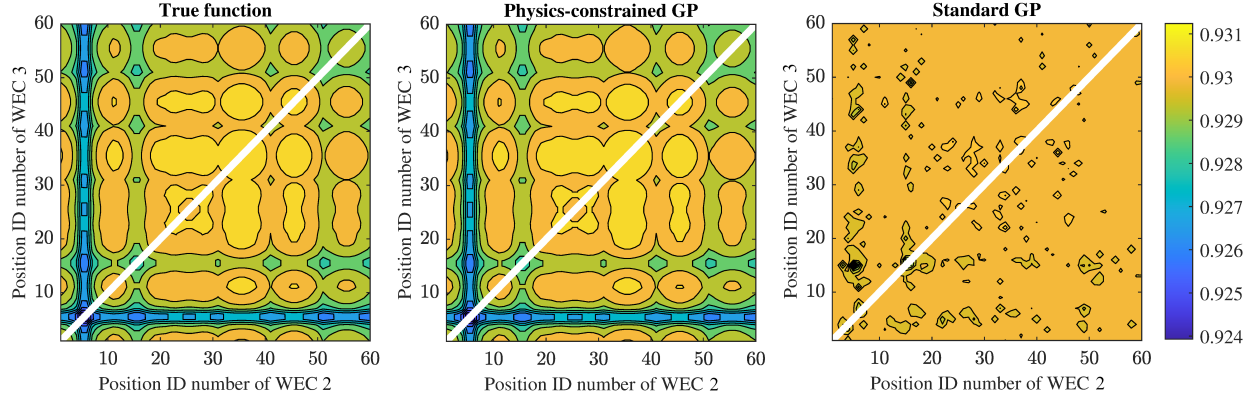


**Figure 3.11:** Covariance matrix calculated by physics-constrained kernel and standard kernel.

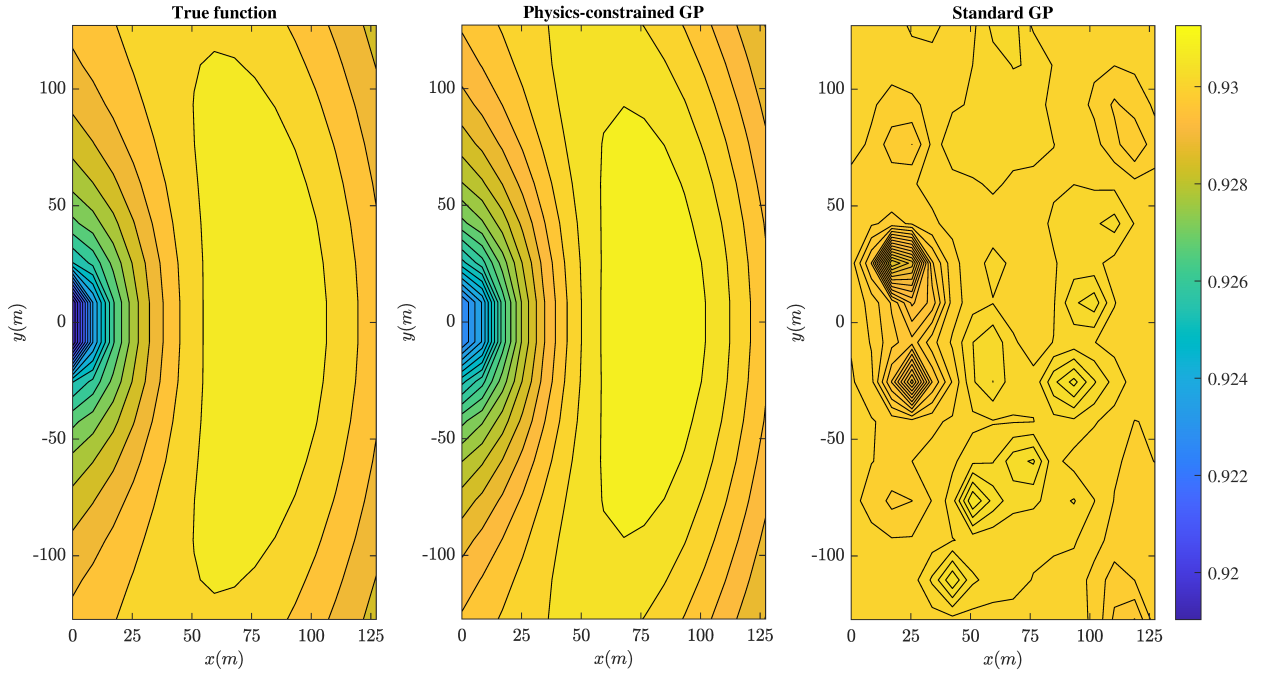
To further demonstrate the performance of the physics-constrained Gaussian process model and the standard Gaussian process model in recovering the ground truth input-output relationship, Figure 3.12 and Figure 3.13 show the model response  $F^{(1)}$  calculated by the true function, predicted by the physics-constrained Gaussian process model and the standard Gaussian process model. In order to visualize the permutation invariance and symmetry features of the model response, we pick the array of 3 WECs and plot the contour map of  $F^{(1)}$  with respect to  $\mathbf{x}_{T,2}$  and  $\mathbf{x}_{T,3}$ , where  $\mathbf{x}_{T,2}$  and  $\mathbf{x}_{T,3}$  are the positions of the two WECs (i.e., WEC 2 and WEC 3) except the one at the origin of the coordinate system. First, Figure 3.12 shows the contour map of  $F^{(1)}$  in terms of the positions of WEC 2 and WEC 3 characterized by some ID numbers. More specifically, the input domain is divided into a 5-by-9 grid, and the nodes of the grid are assigned a sequence of ID numbers. As a result, each ID number represents a position that WEC 2 and WEC 3 may take in the input domain,

and thus in Figure 3.12 the horizontal and vertical axes correspond to the positions of WEC 2 and WEC 3, respectively. As can be observed from the figure, the true function of  $F^{(1)}$  is symmetrical with respect to the diagonal line, which means switching the position of WEC 2 and WEC 3 does not change the model response. Note that the diagonal line in the figure is blank because the two buoys cannot be overlapped with each other physically. The physics-constrained Gaussian process model can correctly recover such permutation invariance feature since such feature has been entirely coded in to the kernel, while the standard Gaussian process model fails to capture the permutation invariance. However, we can find some slight symmetry property with respect to the diagonal line predicted by the standard Gaussian process model, which might have been learned by the training data, and if more training data is used, it is expected that the constructed standard Gaussian process might be able to capture more of the permutation invariance. Second, Figure 3.13 shows the contour map of  $F^{(1)}$  in terms of the position of WEC 2 characterized by its  $[x, y]$  coordinate. Here WEC 3 moves symmetrically with respect to  $x$ -axis, and also in order to make sure the whole array is not symmetrical with respect to the  $x$ -axis, we set the distance of WEC 3 to the  $x$ -axis as  $2/3$  of the corresponding distance of WEC 2 to the  $x$ -axis. In this case, if we reflect WEC 2 with respect to the  $x$ -axis, the whole layout of the array will also be reflected with respect to the  $x$ -axis. Based on the true function in Figure 3.13, we know the model response is symmetrical with respect to  $x$ -axis. Again, the physics-constrained Gaussian process models perfectly captures the symmetry feature because of the use of kernels that explicitly incorporate symmetry, while the standard Gaussian process model can not capture the symmetry feature.

Table 3.2 reports the accuracy metrics of the predictions from the physics-constrained and the standard Gaussian process models. Note that the metrics are averaged over the elements of each hydrodynamic characteristics (i.e.,  $\mathbf{F}$ ,  $\mathbf{a}$ , and  $\mathbf{b}$ ), and these hydrodynamic characteristics are calcu-



**Figure 3.12:** Permutation invariance of  $F^{(1)}$  from true function, physics-constrained Gaussian process model and standard Gaussian process model (GP in the figure stands for Gaussian process model).



**Figure 3.13:**  $x$ -axis symmetry of  $F^{(1)}$  from true function, physics-constrained Gaussian process model and standard Gaussian process model (GP in the figure stands for Gaussian process model).

lated for  $\omega = 0.3$  rad/s. From Table 3.2, we can conclude that the physics-constrained Gaussian process model significantly outperforms the standard Gaussian process model in terms of the prediction accuracy, especially when the array has a large number of WECs. When the array has 3 WECs, the standard Gaussian process model trained with 1000 data can also obtain a relatively good prediction. However, when the number of WECs in the array reaches 5, its prediction accu-

racy drops drastically, and 1000 training data is not enough for a standard Gaussian process model to correctly learn the input-output relationship. Therefore, the standard Gaussian process model cannot explain the highly nonlinear relationship with limited training data and also suffers severely from the curse of dimensionality. In comparison, for physics-constrained Gaussian process model, the  $R^2$  values calculated in all cases are close to or equal to 1, which indicates excellent prediction accuracy. This validates that encoding the invariance, symmetry, and additivity features into the kernel according to the input-output relationship can significantly improve the prediction accuracy. More importantly, with the same number of training data, the prediction accuracy for 10 WECs almost does not change/reduce compared to that for 3 WECs, which means the prediction is less prone to the curse of dimensionality when wave frequency  $\omega = 0.3$  rad/s. One possible reason is that using the physics-constrained kernel can build interpretable Gaussian process model directly rather than learning the features from a large number of training data and thus can help reduce the required number of training data. Also, it has been proved that kernels which include lower-order additive structures sometimes allow us to make predictions over data far away from the training data (i.e., extrapolation) (Duvenaud 2014). For example, additive kernels of first order give high covariance between model response at input locations that are similar in any one dimension. This is of significant importance to arrays with a large number of WECs, where the computational effort in running the MS solver is quite expensive and obtaining a large number of training data is sometimes impractical. Additionally, the results prove that the physics-constrained kernel has included the most contributed part of the model response, and this ensures that our computation in evaluating the proposed physics-constrained kernel can be efficient. It is noteworthy to point out that if the wave frequency is high, the model input-output relationship may become more nonlinear and thus the prediction accuracy for both the physics-constrained and the standard

**Table 3.2:** Prediction accuracy metrics from physics-constrained Gaussian process model and standard Gaussian process model

| # of WECs | Physics-constrained Gaussian process model |          |          | Standard Gaussian process model |          |          |
|-----------|--------------------------------------------|----------|----------|---------------------------------|----------|----------|
|           | <b>F</b>                                   | <b>a</b> | <b>b</b> | <b>F</b>                        | <b>a</b> | <b>b</b> |
| 3         | 1.000                                      | 0.999    | 1.000    | 0.709                           | 0.695    | 0.750    |
| 5         | 1.000                                      | 0.999    | 0.999    | 0.089                           | 0.082    | 0.091    |
| 8         | 1.000                                      | 0.999    | 0.999    | 0.000                           | 0.000    | 0.000    |
| 10        | 1.000                                      | 0.999    | 0.999    | 0.000                           | 0.000    | 0.000    |

Gaussian process models might decrease with the same set of training data. For example, from Table 3.1, we can observe a drop in the prediction accuracy for radiation-related characteristics as the number of WEC in the array increases when  $\omega = 1.4$  rad/s. The physics-constrained kernel alleviates the curse of dimensionality to some extent, however due to the high non-linearity, 1000 training data is still not enough to obtain accurate predictions for arrays with 8 and 10 WECs. In this case, the straightforward way is to increase the number of training data to increase prediction accuracy. However, overall the physics-constrained Gaussian process model still outperforms the standard Gaussian process model in terms of obtaining ideal prediction accuracy with a relatively small number of training data.

#### 3.4.7 Computational efficiency

In this section, we discuss the efficiency gain provided by the physics-constrained Gaussian process model. On average, one evaluation of the numerical model for calculating the hydrodynamic characteristics of arrays of 3, 5, 8, and 10 WECs takes 2.5s, 9.6s, 37.7s, 73.1s, respectively. Overall, the computational time increases dramatically with the number of WECs in the numerical model, further highlighting the computational challenges in modeling large array. On the other hand, the physics-constrained Gaussian process model trained using 1000 training data takes only around 0.001s, 0.020s, 0.081s, 0.163s to obtain all the hydrodynamic characteristics. Therefore,

several orders-of-magnitude speedup can be obtained. Note that when more expensive numerical models (such as boundary element models) are used, the computational gain by using the physics-constrained model will be even greater.

### **3.5 Conclusions**

This chapter proposed a physics-constrained Gaussian process model to efficiently predict responses that show invariance, symmetry, and additivity features. Instead of building a standard Gaussian process model and learning the physical constraints (i.e., invariance, symmetry, and additivity features) of the problem through a large training data set which is computationally inefficient or even prohibitive, the proposed physics-constrained Gaussian process model directly encodes these physical constraints/features into the model development. More specifically, the known physical constraints are encoded into the kernel by designing and integrating the invariant kernel and the additive kernel, and in this way a more “informative prior” can be provided for the Gaussian process model construction. Once trained, the physics-constrained Gaussian process model can be employed to directly and efficiently predict the responses.

Application to prediction of the hydrodynamic characteristics of arrays with different number of wave energy converters (WECs) demonstrates the high accuracy and efficiency of the proposed approach. The results show that the designed integrated kernel is able to correctly capture the invariance, symmetry, and additivity features of the problem. The proposed physics-constrained Gaussian process model can accurately predict the hydrodynamic characteristics with a relatively small number of training data, especially when the wave frequency is low. More importantly, the proposed physics-constrained Gaussian process model is much less vulnerable to curse of dimen-

sionality compared to standard Gaussian process model, and such good scalability is crucial for analyzing arrays with relatively large number of WECs.

One limitation of the proposed physics-constrained Gaussian process model is that if the decomposition of the many-body system does not converge fast (e.g., requires much more than first three orders of interaction), the computational cost of evaluating the proposed kernel might be high and thus it might take longer time to train the Gaussian process model. Also, the application only investigated arrays with up to 10 WECs. Future research work will investigate the scalability of the proposed model to arrays with even larger sizes (e.g., with 20 to 50 WECs). Another future research topic of interest is to use the hydrodynamic characteristics predicted by the proposed physics-constrained Gaussian process model to calculate the total power generation of the WEC array and also optimize the layout of the WEC array.

## CHAPTER 4:

### MULTI-FIDELITY GAUSSIAN PROCESS MODEL INTEGRATING LOW- AND HIGH-FIDELITY DATA CONSIDERING CENSORING<sup>1</sup>

#### 4.1 Introduction

The prediction accuracy of surrogate models highly depends on the given training data. However, for many systems, the accurate evaluation of high-fidelity system models is time consuming and/or costly, which limits the size of the training data. Smaller training data will lead to incomplete coverage of the input space and reduced accuracy of the established Gaussian process model in making predictions at new inputs. On the other hand, while low-fidelity system models can be evaluated more efficiently, they lack accuracy in predicting the system response/output. However, the low-fidelity system models do provide some information on the system behavior. Therefore, to leverage the high accuracy of high-fidelity data and the high efficiency of low-fidelity data, multi-fidelity Gaussian process have been proposed in the literature (Kennedy and O'Hagan 2000; Forrester et al. 2007; Qian and Wu 2008; Le Gratiet 2013a). The basic idea is to integrate information from a small number of high-fidelity data with a large number of low-fidelity data and leverage the correlations between the outputs from different fidelity to efficiently establish an accurate Gaussian process model that can be used for prediction at new design sites in place of the high-fidelity model. However, existing multi-fidelity Gaussian process models mainly use constant scaling factors to accommodate the difference between the high-fidelity part and the low-fidelity part, and also measurement error is usually not taken into account when constructing the model. On the other hand, for many engineering applications, the outputs of high-fidelity models are not

---

<sup>1</sup>This chapter is adapted from a published paper by the author (Li and Jia 2020).

always the exact values of the interested output, rather, censored data, which only provide bounds for the interested output, are given. Gaussian process models that can address censored data and make best use of limited number of high-fidelity data and leverage the information from data with different levels of accuracy are needed.

To bridge the research gap, this chapter proposes a general multi-fidelity Gaussian process model integrating low-fidelity data and high-fidelity data considering censoring in high-fidelity data. To alleviate the cost associated with establishing high-fidelity data, the proposed model integrates information from a small number of expensive high-fidelity data and information provided by a large number of cheap low-fidelity data to efficiently inform a more accurate surrogate model. Censored data are explicitly considered in calibration of the multi-fidelity Gaussian process model. Posterior distributions of the model parameters are established in the context of Bayesian updating to explicitly take into account the uncertainties in the model parameters and the measurement errors in the data. To address the computational challenges in estimating the likelihood function of model parameters when considering censored data, data augmentation algorithm is adopted where the censored data are treated as additional uncertain model parameters and closed form conditional posterior distributions are derived for the model parameters. Gibbs sampling is then used to efficiently generate samples from the posterior distributions for the model parameters, which are used to establish the posterior statistics for the output predictions at new inputs. The effectiveness of the proposed model is illustrated in an example to establish predictive model for the deformation capacity of reinforced concrete (RC) columns using limited number of high-fidelity experimental data (the majority of which are censored data) and a large number of low-fidelity data established from analytical and numerical modeling.

## 4.2 Multi-Fidelity Gaussian Process Prediction Model

The Gaussian process model has been reviewed in Chapter 2, and the formulation discussed only works well when the observations have single fidelity. This section extends the model formulation to a multi-fidelity version, i.e., multi-fidelity Gaussian process model. This section presents the proposed multi-fidelity Gaussian process model and establishes the posterior prediction equation when the multi-fidelity Gaussian process model is used for response prediction at new inputs.

### 4.2.1 Model formulation

For a system model with input  $\mathbf{x} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{n_x}$  and corresponding output/response  $y(\mathbf{x}) \in \mathbb{R}$ , a multi-fidelity Gaussian process model can be trained to approximate the system input-output relationship based on some training set. Let  $y_h(\mathbf{x})$  and  $y_l(\mathbf{x})$  denote the outputs from high-fidelity model and low-fidelity model, respectively. The multi-fidelity Gaussian process model can be written as

$$y_h(\mathbf{x}) = \rho(\mathbf{x})y_l(\mathbf{x}) + \delta(\mathbf{x}) + \epsilon \quad (4.1)$$

where  $\rho(\mathbf{x})$  is the scaling factor, and  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$  are modeled as Gaussian processes, and  $\epsilon$  denotes the measurement error term and is assumed to be Gaussian with zero mean and variance of  $\sigma_\epsilon^2$ . Existing multi-fidelity Gaussian process models usually use a constant scaling factor (Kennedy and O'Hagan 2000; Forrester et al. 2008; Kuya et al. 2011), i.e.,  $\rho(\mathbf{x}) = \rho$ . Also, measurement error is usually not explicitly modeled. Here we use a non-constant scaling factor  $\rho(\mathbf{x})$  that varies with the input  $\mathbf{x}$  to accommodate more complex relationships between  $y_h(\mathbf{x})$  and  $y_l(\mathbf{x})$ , and also use  $\epsilon$  to explicitly consider the measurement error. For the model in Eq. (4.1),  $\rho(\mathbf{x})$  corresponds to the scaling correction term, while  $\delta(\mathbf{x})$  corresponds to the additive bias correction term; both terms are used to correct the low-fidelity model to match the high-fidelity model. Compared to

Gaussian process models in the literature where typically only scaling correction (e.g., constant scaling) or additive correction is used, the model proposed here includes both correction terms to make the overall model more general and have more flexibility in capturing the potentially complex correction relationship between the low-fidelity model and the high-fidelity model over different input values  $\mathbf{x}$ . More specifically, here  $\rho(\mathbf{x})$  will be modeled as  $\rho(\mathbf{x}) = \mathbf{f}_\rho(\mathbf{x})^T \boldsymbol{\beta}_\rho$  where  $\mathbf{f}_\rho(\mathbf{x})$  is the  $q_\rho$ -dimensional vector of basis function and  $\boldsymbol{\beta}_\rho = [\beta_{\rho 1}, \beta_{\rho 2}, \dots, \beta_{\rho q_\rho}]^T$  is the regression coefficients vector. The use of  $\rho(\mathbf{x})$  does increase the complexity in the calibration of the multi-fidelity Gaussian process, and steps to calibrate  $\boldsymbol{\beta}_\rho$  will be developed and discussed in detail in later sections.

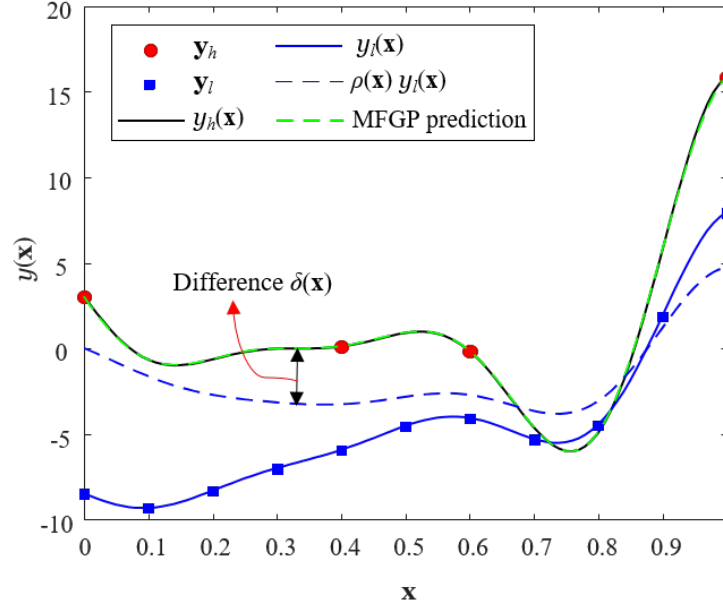
In general, when analytical expressions for  $y_l(\mathbf{x})$  are available, they can be directly used instead of using Gaussian process model; however, in many cases, analytical models may not exist and even low-fidelity models may still take a lot of computational effort. Therefore, considering the more general case for  $y_l(\mathbf{x})$ , in the context of multi-fidelity Gaussian process model, Gaussian process is applied to model  $y_l(\mathbf{x})$  as well. For  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$ , since they are modeled as Gaussian processes, we can write  $y_l(\mathbf{x}) \sim N(\mathbf{f}_l(\mathbf{x})^T \boldsymbol{\beta}_l, \sigma_l^2)$  and  $\delta(\mathbf{x}) \sim N(\mathbf{f}_\delta(\mathbf{x})^T \boldsymbol{\beta}_\delta, \sigma_\delta^2)$ . Similar to the notations in the previous section,  $\mathbf{f}_l(\mathbf{x})$  and  $\mathbf{f}_\delta(\mathbf{x})$  are  $q_l$ -dimensional and  $q_\delta$ -dimensional vectors of basis functions, and  $\boldsymbol{\beta}_l$  and  $\boldsymbol{\beta}_\delta$  are the regression coefficient vectors. The correlation parameters of the Gaussian process models for  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$  are denoted by  $\boldsymbol{\theta}_l$  and  $\boldsymbol{\theta}_\delta$ , respectively. Note that in the model in Eq. (4.1),  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$  are two separate terms in the multi-fidelity Gaussian process model. Instead of explicitly modeling them as dependent processes, here their dependence is implicitly modeled through (i) their common dependence on  $\mathbf{x}$ , and (ii)  $\delta(\mathbf{x})$  depends on the difference between  $y_h(\mathbf{x})$  and  $y_l(\mathbf{x})$ . Later for model calibration, the model parameters in  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$  will also be calibrated jointly to capture the dependence between them.

The multi-fidelity Gaussian process model can be calibrated using a combination of high-fidelity data and low-fidelity data. Let  $\mathbf{Y} = [\mathbf{y}_l; \mathbf{y}_h] = [y_l(\mathbf{x}_l^1), \dots, y_l(\mathbf{x}_l^{n_l}), y_h(\mathbf{x}_h^1), \dots, y_h(\mathbf{x}_h^{n_h})]^T$  denote the training data/observations including  $n_l$  low-fidelity data  $\mathbf{y}_l = [y_l(\mathbf{x}_l^1), \dots, y_l(\mathbf{x}_l^{n_l})]^T$  and  $n_h$  high-fidelity data  $\mathbf{y}_h = [y_h(\mathbf{x}_h^1), \dots, y_h(\mathbf{x}_h^{n_h})]^T$  with the corresponding design sites denoted by  $\mathbf{X} = [\mathbf{x}_l; \mathbf{x}_h] = [\mathbf{x}_l^1, \dots, \mathbf{x}_l^{n_l}, \mathbf{x}_h^1, \dots, \mathbf{x}_h^{n_h}]^T$ . Here we use the common assumption that  $\mathbf{x}_h$  is a subset of  $\mathbf{x}_l$ . Intuitively, this assumption allows direct comparison between high-fidelity data and low-fidelity data for the design sites in  $\mathbf{x}_h$ , guiding how the low-fidelity data should be “corrected” to match the high-fidelity data. This assumption will also facilitate deriving analytical expressions for the posterior of the model parameters, and this derivation will be discussed in later section. Given the training data, all the model parameters in the multi-fidelity Gaussian process model, denoted here by  $\boldsymbol{\eta} = [\beta_\rho, \beta_l, \sigma_l^2, \boldsymbol{\theta}_l, \beta_\delta, \sigma_\delta^2, \boldsymbol{\theta}_\delta, \sigma_\epsilon^2]$ , can be calibrated. Using the calibrated model, the prediction or more specifically the predictive distribution of  $y_h(\mathbf{x}^0)$  at a new design point  $\mathbf{x}^0$  can be obtained, which is discussed next.

As an illustration, Figure 4.1 shows an example of the multi-fidelity Gaussian process modeling, where the red dots and blue squares are the high-fidelity and low-fidelity training data, respectively. The black solid line represents the unknown function to be predicted (i.e., high-fidelity model), the blue solid line represents the low-fidelity model, the blue dashed line is the scaled low-fidelity model, and the green dashed line corresponds the prediction by the multi-fidelity Gaussian process model.

#### 4.2.2 Posterior predictions at new inputs

In this section we will discuss the prediction of  $y_h$  at new input  $\mathbf{x}^0$ . Based on the observations  $\mathbf{y}_l$  and  $\mathbf{y}_h$ ,  $\boldsymbol{\eta}$  can be calibrated. Let  $p(\boldsymbol{\eta}|\mathbf{y}_h, \mathbf{y}_l)$  denote the posterior distribution for  $\boldsymbol{\eta}$ , which will be



**Figure 4.1:** Illustrative example of multi-fidelity Gaussian process model (MFGP in the figure stands for multi-fidelity Gaussian process model).

discussed in detail in next section. Conditional on given observations  $\mathbf{y}_l$  and  $\mathbf{y}_h$  and specific values of the model parameters  $\boldsymbol{\eta}$ , the multi-fidelity Gaussian process model gives a prediction of  $y_h$  that follows a Gaussian distribution and can be denoted as  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}]$ . Therefore, propagating the uncertainty in  $\boldsymbol{\eta}$ , we can establish the predictive distribution for  $y_h$ , denoted  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l]$ ,

$$p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l] = \int p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}]p(\boldsymbol{\eta}|\mathbf{y}_h, \mathbf{y}_l)d\boldsymbol{\eta} \quad (4.2)$$

With  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l]$  or samples from it, we can then easily calculate the statistics (e.g., mean and variance) for  $y_h(\mathbf{x}^0)$  as well. To establish  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l]$ , we need to figure out (i) what is the expression for  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}]$ , and (ii) what is the posterior distribution  $p(\boldsymbol{\eta}|\mathbf{y}_h, \mathbf{y}_l)$  and how to sample from it. These two aspects are discussed in detail in the following sections with a focus on (ii).

First, for the conditional distribution  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}]$  in Eq. (4.2), it corresponds to a Gaussian with mean  $m_{y_h}(\mathbf{x}^0)$  and variance  $s_{y_h}^2(\mathbf{x}^0)$ , which are given by

$$m_{y_h}(\mathbf{x}^0) = \rho(\mathbf{x}^0)\mathbf{f}_l(\mathbf{x}^0)^T\boldsymbol{\beta}_l + \mathbf{f}_\delta(\mathbf{x}^0)^T\boldsymbol{\beta}_\delta + \mathbf{k}^T\mathbf{V}^{-1}(\mathbf{Y} - \mathbf{M}) \quad (4.3)$$

$$s_{y_h}^2(\mathbf{x}^0) = k_{hh}(\mathbf{x}^0, \mathbf{x}^0) - \mathbf{k}^T\mathbf{V}^{-1}\mathbf{k} \quad (4.4)$$

with  $\rho(\mathbf{x}^0) = \mathbf{f}_\rho(\mathbf{x}^0)^T\boldsymbol{\beta}_\rho$ ,  $\mathbf{k} = [\mathbf{k}_{hl}(\mathbf{x}^0, \mathbf{x}_l)^T, \mathbf{k}_{hh}(\mathbf{x}^0, \mathbf{x}_h)^T]^T$ , and

$$\mathbf{V} = \begin{pmatrix} \mathbf{K}_{ll}(\mathbf{x}_l, \mathbf{x}_l) & \mathbf{K}_{lh}(\mathbf{x}_l, \mathbf{x}_h) \\ \mathbf{K}_{hl}(\mathbf{x}_h, \mathbf{x}_l) & \mathbf{K}_{hh}(\mathbf{x}_h, \mathbf{x}_h) \end{pmatrix}, \mathbf{M} = \begin{pmatrix} \mathbf{F}_l(\mathbf{x}_l)\boldsymbol{\beta}_l \\ \mathbf{P}(\mathbf{x}_h) \cdot (\mathbf{F}_l(\mathbf{x}_h)\boldsymbol{\beta}_l) + \mathbf{F}_\delta(\mathbf{x}_h)^T\boldsymbol{\beta}_\delta \end{pmatrix}$$

The analytical form of the conditional distribution  $p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}]$  is derived from the following joint distribution,

$$\begin{pmatrix} y_h(\mathbf{x}^0) \\ \mathbf{y}_l(\mathbf{x}_l) \\ \mathbf{y}_h(\mathbf{x}_h) \end{pmatrix} \sim N \left( \begin{pmatrix} \rho(\mathbf{x}^0)(\mathbf{f}_l(\mathbf{x}^0)^T\boldsymbol{\beta}_l) + \mathbf{f}_\delta(\mathbf{x}^0)^T\boldsymbol{\beta}_\delta \\ \mathbf{F}_l(\mathbf{x}_l)\boldsymbol{\beta}_l \\ \mathbf{P}(\mathbf{x}_h)(\mathbf{F}_l(\mathbf{x}_h)\boldsymbol{\beta}_l) + \mathbf{F}_\delta(\mathbf{x}_h)\boldsymbol{\beta}_\delta \end{pmatrix}, \begin{pmatrix} k_{hh}(\mathbf{x}^0, \mathbf{x}^0) & \mathbf{k}_{hl}(\mathbf{x}^0, \mathbf{x}_l)^T & \mathbf{k}_{hh}(\mathbf{x}^0, \mathbf{x}_h)^T \\ \mathbf{k}_{lh}(\mathbf{x}_l, \mathbf{x}^0) & \mathbf{K}_{ll}(\mathbf{x}_l, \mathbf{x}_l) & \mathbf{K}_{lh}(\mathbf{x}_l, \mathbf{x}_h) \\ \mathbf{k}_{hh}(\mathbf{x}_h, \mathbf{x}^0) & \mathbf{K}_{hl}(\mathbf{x}_h, \mathbf{x}_l) & \mathbf{K}_{hh}(\mathbf{x}_h, \mathbf{x}_h) \end{pmatrix} \right) \quad (4.5)$$

with  $\mathbf{F}_l(\mathbf{x}_l) = [\mathbf{f}_l(\mathbf{x}_l^1)^T, \dots, \mathbf{f}_l(\mathbf{x}_l^{n_l})^T]^T$ ,  $\mathbf{F}_l(\mathbf{x}_h) = [\mathbf{f}_l(\mathbf{x}_h^1)^T, \dots, \mathbf{f}_l(\mathbf{x}_h^{n_h})^T]^T$ ,  $\mathbf{F}_\delta(\mathbf{x}_h) = [\mathbf{f}_\delta(\mathbf{x}_h^1)^T, \dots, \mathbf{f}_\delta(\mathbf{x}_h^{n_h})^T]^T$ ,  $\mathbf{P}(\mathbf{x}_h) = \text{diag}[\rho(\mathbf{x}_h^1), \dots, \rho(\mathbf{x}_h^{n_h})]$ ,  $\mathbf{k}_{st}(\mathbf{x}_s, \mathbf{x}^0) = \mathbf{k}_{ts}(\mathbf{x}^0, \mathbf{x}_s) = [k_{st}(\mathbf{x}_s^1, \mathbf{x}^0), \dots, k_{st}(\mathbf{x}_s^{n_s}, \mathbf{x}^0)]^T$ , and  $k_{st}(\mathbf{x}_s^i, \mathbf{x}_t^j)$  denotes the covariance between  $\mathbf{y}_s(\mathbf{x}_s^i)$  and  $\mathbf{y}_t(\mathbf{x}_t^j)$ , where  $s$  and  $t$

represent low-fidelity or high-fidelity. Additionally, we have

$$\mathbf{K}_{st}(\mathbf{x}_s, \mathbf{x}_t) = \begin{pmatrix} \mathbf{k}_{st}(\mathbf{x}_s^1, \mathbf{x}_t^1) & \cdots & \mathbf{k}_{st}(\mathbf{x}_s^1, \mathbf{x}_t^{n_t}) \\ \vdots & \ddots & \vdots \\ \mathbf{k}_{st}(\mathbf{x}_s^{n_s}, \mathbf{x}_t^1) & \cdots & \mathbf{k}_{st}(\mathbf{x}_s^{n_s}, \mathbf{x}_t^{n_t}) \end{pmatrix} \quad (4.6)$$

### 4.3 Posterior Distributions of Model Parameters

To obtain the predictive distribution of  $y_h(\mathbf{x}^0)$ , we also need to establish the posterior distributions of the unknown parameters  $\boldsymbol{\eta}$ , which is proportional to the prior  $p(\boldsymbol{\eta})$  multiplied by the likelihood function  $L(D|\boldsymbol{\eta}) = p(\mathbf{y}_h, \mathbf{y}_l|\boldsymbol{\eta})$ , i.e.,  $p(\boldsymbol{\eta}|\mathbf{y}_h, \mathbf{y}_l) \propto p(\boldsymbol{\eta})p(\mathbf{y}_h, \mathbf{y}_l|\boldsymbol{\eta})$ , where  $D$  represents the training data including both high-fidelity data and low-fidelity data. This section first discusses the selection of the prior distributions for  $\boldsymbol{\eta}$ , and then derives the expressions for the likelihood functions when there is censored data in the observations. The latter is part of the novelty of this chapter, since existing models only consider accurate (or equality) data (i.e., without censored data). In this section, the unknown parameters are collected into two groups:  $\boldsymbol{\eta}_1 = [\boldsymbol{\beta}_l, \sigma_l^2, \boldsymbol{\theta}_l]$  associated with low-fidelity response  $y_l(\mathbf{x})$ , and  $\boldsymbol{\eta}_2 = [\boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_\delta^2, \boldsymbol{\theta}_\delta, \sigma_\epsilon^2]$  associated with high-fidelity response  $y_h(\mathbf{x})$ .

#### 4.3.1 Prior distributions of model parameters

Assuming independence between the model parameters, the prior distribution of the unknown model parameter  $\boldsymbol{\eta}$  can be written as

$$p(\boldsymbol{\eta}) = p(\boldsymbol{\eta}_1, \boldsymbol{\eta}_2) = p(\boldsymbol{\eta}_1)p(\boldsymbol{\eta}_2) = p(\boldsymbol{\beta}_l, \sigma_l^2, \boldsymbol{\theta}_l)p(\boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_\delta^2, \boldsymbol{\theta}_\delta, \sigma_\epsilon^2) \quad (4.7)$$

Either informative and non-informative priors can be used. To establish analytical posterior distributions of the parameters, the following priors are adopted:  $\beta_l, \beta_\rho, \beta_\delta$  are assumed to follow the multivariate normal distribution,  $\sigma_l^2, \sigma_\delta^2, \sigma_\epsilon^2$  are assumed to follow the inverse gamma distribution and  $\theta_l, \theta_\delta$  are assumed to follow the gamma distribution. More specifically, we have

$$\begin{aligned} \beta_l &\sim N(\mathbf{m}_l, \mathbf{V}_l^{-1}), \sigma_l^2 \sim IG(\alpha_l, \gamma_l), \theta_{l,i} \sim G(a_l, b_l), \beta_\rho \sim N(\mathbf{m}_\rho, \mathbf{V}_\rho^{-1}), \\ \beta_\delta &\sim N(\mathbf{m}_\delta, \mathbf{V}_\delta^{-1}), \sigma_\delta^2 \sim IG(\alpha_\delta, \gamma_\delta), \theta_{\delta,i} \sim G(a_\delta, b_\delta), \sigma_\epsilon^2 \sim IG(\alpha_\epsilon, \gamma_\epsilon) \end{aligned} \quad (4.8)$$

where  $i = 1, 2, \dots, n_x$ ,  $\mathbf{m}_l, \mathbf{m}_\rho, \mathbf{m}_\delta$  are the mean vectors of  $\beta_l, \beta_\rho, \beta_\delta$ , and  $\mathbf{V}_l, \mathbf{V}_\rho, \mathbf{V}_\delta$  are the precision matrices of  $\beta_l, \beta_\rho, \beta_\delta$ . One option to select values for the mean vectors is to set  $\mathbf{m}_l = m_l \mathbf{e}_{q_l}$ ,  $\mathbf{m}_\rho = m_\rho \mathbf{e}_{q_\rho}$ ,  $\mathbf{m}_\delta = m_\delta \mathbf{e}_{q_\delta}$ , and  $\mathbf{e}_k$  is a  $k \times 1$  vector of ones and  $k = q_l$  or  $q_\rho$  or  $q_\delta$ . For the precision matrices,  $\mathbf{V}_l = (v_l/\sigma_l^2) \mathbf{I}_{q_l \times q_l}$ ,  $\mathbf{V}_\rho = v_\rho \mathbf{I}_{q_\rho \times q_\rho}$ ,  $\mathbf{V}_\delta = (v_\delta/\sigma_\delta^2) \mathbf{I}_{q_\delta \times q_\delta}$  can be used where  $\mathbf{I}_{k \times k}$  is a  $k \times k$  identity matrix. To establish more informed priors, another option is to use least square regression to get a better estimate/guess of the values for  $\mathbf{m}_l, \mathbf{m}_\rho, \mathbf{m}_\delta$ . For example,  $\mathbf{m}_l$  can be taken as the least square regression values for  $\beta_l$  based on low-fidelity data. Similarly,  $\mathbf{m}_\delta$  can be taken as the least square regression values for  $\beta_\delta$  by setting the scaling factor  $\rho(\mathbf{x})$  equal to 1 and using the difference between high-fidelity and low-fidelity data. For  $\mathbf{m}_\rho$ , it can be taken as a  $q \times 1$  vector with value of 1 for the first element and 0 for the rest of the elements or more generally uniform random values between 0 and 1 for all the elements. Note that the procedure affects the selection of priors, which in turn will to some extent affect the posterior distribution of the uncertain model parameters since the posterior is proportional to product of the prior and the likelihood. Depending on the amount of data, the impact of such selection of priors varies, e.g., when the data is extremely limited the selection of the prior is expected to have larger impact on the posterior, but when there

is enough data and the likelihood dominates the posterior, the selection of the prior will have less impact on the posterior.

#### 4.3.2 Likelihood functions of model parameters

In this section, we derive the likelihood functions of the unknown parameters based on the training data set  $D$ . Depending on the type of data (e.g., accurate data or censored data) in the training data, the likelihood function will be different. Here we will derive the likelihood function for two cases. The first case is when the training data  $D$  only includes accurate responses, and the second case is when the training set includes both censored responses and accurate responses. Also, here we only consider censored responses in the high-fidelity data/responses, which is the more common case, while the derivation process can be easily extended to consider censored responses in low-fidelity data/responses.

##### **Likelihood function considering only accurate data**

The training data  $D$  can be represented by  $\mathbf{Y} = [\mathbf{y}_l; \mathbf{y}_h]$ , where  $\mathbf{y}_l = [y_l^1, y_l^2, \dots, y_l^{n_l}]^T$ ,  $\mathbf{y}_h = [y_h^1, y_h^2, \dots, y_h^{n_h}]^T$ . To simplify the notation, we also define  $\xi(\mathbf{x}) = \delta(\mathbf{x}) + \epsilon = y_h(\mathbf{x}) - \rho(\mathbf{x})y_l(\mathbf{x})$ . Then the likelihood function of the unknown model parameters  $\boldsymbol{\eta} = [\boldsymbol{\eta}_1, \boldsymbol{\eta}_2]$  can be written as

$$\begin{aligned} L(D|\boldsymbol{\eta}) &= p(\mathbf{y}_h, \mathbf{y}_l | \boldsymbol{\eta}_1, \boldsymbol{\eta}_2) = p(\mathbf{y}_l | \boldsymbol{\eta}_1) p(\mathbf{y}_h | \mathbf{y}_l, \boldsymbol{\eta}_2) \\ &= L(y_l^1, y_l^2, \dots, y_l^{n_l} | \boldsymbol{\beta}_l, \sigma_l^2, \boldsymbol{\theta}_l) L(\xi^1, \xi^2, \dots, \xi^{n_h} | \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_\delta^2, \boldsymbol{\theta}_\delta, \sigma_\epsilon^2) \end{aligned} \quad (4.9)$$

Since  $y_l(\mathbf{x})$  and  $\delta(\mathbf{x})$  are Gaussian process models and  $\epsilon$  corresponds to Gaussian random error, using the joint Gaussian PDF (e.g., over the low-fidelity data and high-fidelity data), we have

$$\begin{aligned} L(y_l^1, y_l^2, \dots, y_l^{n_l} | \beta_l, \sigma_l^2, \theta_l) \\ = \frac{1}{(2\pi)^{n_l/2} |\Psi_l|^{1/2}} \exp \left[ -\frac{1}{2} (\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \beta_l)^T \Psi_l^{-1} (\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \beta_l) \right] \end{aligned} \quad (4.10)$$

$$\begin{aligned} L(\xi^1, \xi^2, \dots, \xi^{n_h} | \beta_\rho, \beta_\delta, \sigma_\delta^2, \theta_\delta, \sigma_\epsilon^2) \\ = \frac{1}{(2\pi)^{n_h/2} |\Psi_\xi|^{1/2}} \exp \left[ -\frac{1}{2} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \beta_\delta)^T \Psi_\xi^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \beta_\delta) \right] \end{aligned} \quad (4.11)$$

with  $\mathbf{A}_{y_l} = \text{diag}[y_l(\mathbf{x}_h^1), \dots, y_l(\mathbf{x}_h^{n_h})]$ ,  $\mathbf{F}_\rho(\mathbf{x}_h) = [\mathbf{f}_\rho(\mathbf{x}_h^1)^T, \dots, \mathbf{f}_\rho(\mathbf{x}_h^{n_h})^T]^T$ .  $\Psi_l$  and  $\Psi_\xi$  are the covariance matrices of  $\mathbf{y}_l$  and  $\boldsymbol{\xi}(\mathbf{x}_h) = [\xi(\mathbf{x}_h^1), \xi(\mathbf{x}_h^2), \dots, \xi(\mathbf{x}_h^{n_h})]^T$ , and have the expressions  $\Psi_l = \sigma_l^2 \mathbf{R}_l$  and  $\Psi_\xi = \sigma_\delta^2 \mathbf{R}_\delta + \sigma_\epsilon^2 \mathbf{I}_{n_h \times n_h}$ , where  $\mathbf{R}_l$  and  $\mathbf{R}_\delta$  are the correlation matrices of  $\mathbf{y}_l$  and  $\boldsymbol{\delta}(\mathbf{x}_h) = [\delta(\mathbf{x}_h^1), \delta(\mathbf{x}_h^2), \dots, \delta(\mathbf{x}_h^{n_h})]^T$ , respectively.

### Likelihood function considering both censored data and accurate data

Now suppose the high-fidelity data set includes both censored and accurate responses, denoted by  $\mathbf{y}_{h,I}$  and  $\mathbf{y}_{h,J}$ , respectively. In this case, the likelihood function of the unknown parameters  $\boldsymbol{\eta} = [\boldsymbol{\eta}_1, \boldsymbol{\eta}_2]$  can be written as  $L(D|\boldsymbol{\eta}) = p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\eta}_1, \boldsymbol{\eta}_2)$ , which can be expanded as

$$L(D|\boldsymbol{\eta}) = p(\mathbf{y}_l | \boldsymbol{\eta}_1) p(\mathbf{y}_{h,J} | \mathbf{y}_l, \boldsymbol{\eta}_2) \int_{\prod_{i \in I} A_i} p(\mathbf{y}_{h,I} | \mathbf{y}_{h,J}, \mathbf{y}_l, \boldsymbol{\eta}_2) d\mathbf{y}_{h,I} \quad (4.12)$$

where  $i$  represents the  $i^{\text{th}}$  censored response,  $I$  is the index set for the censored responses, and  $A_i$  represents the censoring interval. Evaluation of the likelihood in Eq. (4.12) for given  $\boldsymbol{\eta}$  involves evaluation of a potentially high-dimensional integral over the “failure region” defined by  $\prod_{i \in I} A_i$ .

While Monte Carlo simulation (MCS) can be used to evaluate the integral, it typically requires large number of simulations, entailing huge computational efforts. Within the context of Bayesian updating for establishing or sampling from the posterior distributions for  $\eta$ , where many repeated evaluations of the likelihood function for different values of  $\eta$  is required, direct adoption of MCS for each evaluation of the likelihood function is computationally expensive.

To address the above challenge, here we propose a data augmentation approach where the censored data are treated as unknown model parameters (e.g., augmented on top of the existing unknown model parameters). This will eliminate the need to evaluate the high-dimensional integrals when sampling from the posterior distribution. Details of using data augmentation to facilitate sampling from the posterior distribution are presented in the next section, and the posterior sampling using data augmentation within the context of multi-fidelity Gaussian process model considering both censored data and accurate data is a novel contribution of this chapter, which has not been considered in existing research.

## **4.4 Sampling from Posterior Distributions of Model Parameters**

### *4.4.1 Posterior distributions of model parameters*

Using the prior distribution and likelihood function of model parameters established in the previous section, we can establish the posterior distributions of the unknown model parameters. While Qian and Wu (2008) have derived the posterior distributions of model parameters for the multi-fidelity Gaussian process model based on only accurate data, here we derive the posterior distributions of model parameters for the multi-fidelity Gaussian process model considering data with both censored and accurate data (in the high-fidelity data). To facilitate the sampling from the posterior distributions, here we use data augmentation to augment the posterior distribution

by treating the censored data as unknown uncertain model parameters. This will be the focus of this section. More specifically, the censored responses,  $\mathbf{y}_{h,I} = [y_h^1, y_h^2, \dots, y_h^{n_I}]^T$ , are regarded as uncertain parameters, denoted by  $\tilde{\mathbf{y}}_{h,I} = [\tilde{y}_h^1, \tilde{y}_h^2, \dots, \tilde{y}_h^{n_I}]^T$ . Note that for notation convenience, we arrange the data so that the elements of  $y_{h,I}$  are the first  $n_I$  elements of  $y_h$ . The posterior joint distribution for  $[\tilde{\mathbf{y}}_{h,I}, \boldsymbol{\eta}]$  can be written as

$$p(\tilde{\mathbf{y}}_{h,I}, \boldsymbol{\eta} | \mathbf{y}_h, \mathbf{y}_l) \propto p(\boldsymbol{\eta}) p(\tilde{\mathbf{y}}_{h,I} | \boldsymbol{\eta}) p(\mathbf{y}_h, \mathbf{y}_l | \tilde{\mathbf{y}}_{h,I}, \boldsymbol{\eta}) \quad (4.13)$$

where  $p(\tilde{\mathbf{y}}_{h,I} | \boldsymbol{\eta}) = \prod_{i \in I} \mathbf{1}_{\{\tilde{y}_h^i \in A_i\}}$  and  $A_i$  is defined by the actual censored data/responses  $y_{h,I}$ .

Although Eq. (4.13) provides the full joint posterior distribution for  $[\tilde{\mathbf{y}}_{h,I}, \boldsymbol{\eta}]$ , it is not practical to directly generate samples from it. On the other hand, with the selection of the priors discussed earlier and using the Gaussian nature of the likelihood functions, we can derive analytical expressions for the conditional posterior distributions of the model parameters. With analytical expressions, we can directly and efficiently sample from these conditional posterior distributions. Within this context, we will use Gibbs sampling and the analytical expressions for the conditional posterior distributions to efficiently generate samples from the joint posterior distribution. Next, we present the derived conditional posterior distributions for all the parameters.

(1)  $\tilde{y}_h^i$

$$p(\tilde{y}_h^i | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \boldsymbol{\eta}, \tilde{\mathbf{y}}_{h,I}^{(i)}) \propto N(m_{y_h}(\mathbf{x}_h^i), s_{y_h}^2(\mathbf{x}_h^i)) \mathbf{1}_{\{\tilde{y}_h^i \in A_i\}} \quad (4.14)$$

where  $\tilde{y}_h^i (i = 1, 2, \dots, n_I)$  is the  $i^{th}$  element of  $\tilde{\mathbf{y}}_{h,I}$ , and

$$m_{y_h}(\mathbf{x}_h^i) = \rho(\mathbf{x}_h^i) \mathbf{f}_l(\mathbf{x}_h^i)^T \boldsymbol{\beta}_l + \mathbf{f}_\delta(\mathbf{x}_h^i)^T \boldsymbol{\beta}_\delta + [\mathbf{k}_{hl}(\mathbf{x}_h^i, \mathbf{x}_l)^T, \mathbf{k}_{hh}(\mathbf{x}_h^i, \mathbf{x}_h^{(i)})^T] [\mathbf{V}^{(i)}]^{-1} \left( [\mathbf{y}_l, \tilde{\mathbf{y}}_{h,I}^{(i)}, \mathbf{y}_{h,J}]^T - \mathbf{M}^{(i)} \right)$$

$$s_{y_h}^2(\mathbf{x}_h^i) = k_{hh}(\mathbf{x}_h^i, \mathbf{x}_h^i) - [\mathbf{k}_{hl}(\mathbf{x}_h^i, \mathbf{x}_l)^T, \mathbf{k}_{hh}(\mathbf{x}_h^i, \mathbf{x}_h^{(i)})^T] [\mathbf{V}^{(i)}]^{-1} [\mathbf{k}_{lh}(\mathbf{x}_l, \mathbf{x}_h^i), \mathbf{k}_{hh}(\mathbf{x}_h^{(i)}, \mathbf{x}_h^i)]^T$$

and  $\rho(\mathbf{x}_h^i) = \mathbf{f}_\rho(\mathbf{x}_h^i)^T \boldsymbol{\beta}_\rho$ ,  $\mathbf{x}_h^i$  corresponds to  $\mathbf{x}_h$  with the  $i^{th}$  input  $\mathbf{x}^i$  removed,  $\tilde{\mathbf{y}}_{h,I}^{(i)}$  corresponds to  $\tilde{\mathbf{y}}_{h,I}$  with the  $i^{th}$  element  $\tilde{y}_h^i$  removed,  $\mathbf{V}^{(i)}$  is  $\mathbf{V}$  with the  $(n_l + i)^{th}$  column and  $(n_l + i)^{th}$  row removed,  $\mathbf{M}^{(i)}$  is  $\mathbf{M}$  with the  $(n_l + i)^{th}$  row removed.

(2)  $\boldsymbol{\beta}_l$

$$p(\boldsymbol{\beta}_l | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_l) = p(\boldsymbol{\beta}_l | \mathbf{y}_h, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_l) = N(\tilde{\mathbf{m}}_l, \tilde{\mathbf{V}}_l^{-1}) \quad (4.15)$$

where  $\bar{\boldsymbol{\beta}}_l$  represents the rest of the parameters excluding  $\boldsymbol{\beta}_l$  (similar notation is used for the following derivations), and  $\tilde{\mathbf{m}}_l = \tilde{\mathbf{V}}_l^{-1}[\mathbf{V}_l \mathbf{m}_l + \mathbf{F}_l(\mathbf{x}_l)^T \boldsymbol{\Psi}_l^{-1} \mathbf{y}_l]$ ,  $\tilde{\mathbf{V}}_l = \mathbf{V}_l + \mathbf{F}_l(\mathbf{x}_l)^T \boldsymbol{\Psi}_l^{-1} \mathbf{F}_l(\mathbf{x}_l)$ .

(3)  $\boldsymbol{\beta}_\rho$

$$p(\boldsymbol{\beta}_\rho | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_\rho) = p(\boldsymbol{\beta}_\rho | \mathbf{y}_h, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_\rho) = N(\tilde{\mathbf{m}}_\rho, \tilde{\mathbf{V}}_\rho^{-1}) \quad (4.16)$$

where  $\tilde{\mathbf{m}}_\rho = \tilde{\mathbf{V}}_\rho^{-1}[\mathbf{V}_\rho \mathbf{m}_\rho + (\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h))^T \boldsymbol{\Psi}_\xi^{-1}(\mathbf{y}_h - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)]$  and  $\tilde{\mathbf{V}}_\rho = \mathbf{V}_\rho + (\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h))^T \boldsymbol{\Psi}_\xi^{-1}(\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h))$ .

(4)  $\boldsymbol{\beta}_\delta$

$$p(\boldsymbol{\beta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_\delta) = p(\boldsymbol{\beta}_\delta | \mathbf{y}_h, \mathbf{y}_l, \bar{\boldsymbol{\beta}}_\delta) = N(\tilde{\mathbf{m}}_\delta, \tilde{\mathbf{V}}_\delta^{-1}) \quad (4.17)$$

where  $\tilde{\mathbf{m}}_\delta = \tilde{\mathbf{V}}_\delta^{-1}[\mathbf{V}_\delta \mathbf{m}_\delta + \mathbf{F}_\delta(\mathbf{x}_h)^T \boldsymbol{\Psi}_\xi^{-1}(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho)]$  and  $\tilde{\mathbf{V}}_\delta = \mathbf{V}_\delta + \mathbf{F}_\delta(\mathbf{x}_h)^T \boldsymbol{\Psi}_\xi^{-1} \mathbf{F}_\delta(\mathbf{x}_h)$ .

(5)  $\sigma_l^2$

$$p(\sigma_l^2 | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\sigma}_l^2) = p(\sigma_l^2 | \mathbf{y}_h, \mathbf{y}_l, \bar{\sigma}_l^2) = IG(\tilde{\alpha}_l, \tilde{\gamma}_l) \quad (4.18)$$

where  $\tilde{\alpha}_l = \frac{n_l}{2} + \frac{q_l}{2} + \alpha_l$  and  $\tilde{\gamma}_l = \frac{v_l}{2}(\boldsymbol{\beta}_l - \mathbf{m}_l)^T(\boldsymbol{\beta}_l - \mathbf{m}_l) + \frac{1}{2}[\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l]^T \mathbf{R}_l^{-1}[\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l] + \gamma_l$ .

Note that the derivation of the posterior distributions for  $\boldsymbol{\beta}_l$ ,  $\boldsymbol{\beta}_\rho$ ,  $\boldsymbol{\beta}_\delta$  and  $\sigma_l^2$  is based on the property of conjugate prior (Gelman et al. 2013). For  $\boldsymbol{\beta}_l$ ,  $\boldsymbol{\beta}_\rho$  and  $\boldsymbol{\beta}_\delta$ , the likelihood is a normal distribution with known variance and the prior distributions of these parameters are normal distributions, which are the conjugate priors, therefore the posterior distributions will be in the same probability density

family as the priors (i.e., normal distributions). Similarly, the posterior distribution for  $\sigma_l^2$  is an inverse gamma because the likelihood is a normal distribution with known mean and the conjugate prior distribution is an inverse gamma.

(6)  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$

$$p(\sigma_\delta^2, \sigma_\epsilon^2 | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \overline{\sigma_\delta^2}, \overline{\sigma_\epsilon^2}) \propto (\sigma_\delta^2)^{-n_h/2 - \alpha_\delta - 1} (\sigma_\epsilon^2)^{-\alpha_\epsilon - 1} \exp \left[ -\frac{\gamma_\epsilon}{\sigma_\epsilon^2} - \frac{\gamma_\delta}{\sigma_\delta^2} \right] \frac{1}{|\mathbf{R}_{\xi\delta}|^{1/2}} \quad (4.19)$$

$$\times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)}{2\sigma_\delta^2} \right]$$

where  $\mathbf{R}_{\xi\delta} = \mathbf{R}_\delta + (\sigma_\epsilon^2 / \sigma_\delta^2) \mathbf{I}_{n_h \times n_h}$ .

(7)  $\theta_l$  and  $\theta_\delta$

$$p(\theta_l | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\theta}_l) \propto p(\theta_l) \frac{1}{(\sigma_l^2)^{n_l/2} |\mathbf{R}_l|^{1/2}} \exp \left[ -\frac{(\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l)^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l)}{2\sigma_l^2} \right] \quad (4.20)$$

$$p(\theta_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\theta}_\delta) \propto p(\theta_\delta) \frac{1}{(\sigma_\delta^2)^{n_h/2} |\mathbf{R}_{\xi\delta}|^{1/2}} \quad (4.21)$$

$$\times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)}{2\sigma_\delta^2} \right]$$

#### 4.4.2 Sampling from posterior distributions of model parameters

As can be seen from the above expressions, the posterior distributions for  $\tilde{y}_h^i, \beta_l, \beta_\rho, \beta_\delta, \sigma_l^2$  are of standard forms, which means samples for these parameters can be directly generated from the corresponding distributions. By contrast, the samples for  $\sigma_\delta^2, \sigma_\epsilon^2, \theta_l$  and  $\theta_\delta$  cannot be directly generated since the conditional posterior distributions for these parameters do not have closed-form analytical expressions (or do not correspond to standard distributions, e.g., Gaussian or inverse Gamma or other known distributions). To generate samples from these distributions, the general

Metropolis-Hastings algorithm will be applied. For ease of the posterior sampling, we will use point estimates for the correlation parameters  $\theta_l$  and  $\theta_\delta$  instead of sampling from the corresponding posterior distributions. In this case, for selected values of  $\theta_l$  and  $\theta_\delta$ , the posterior predictive distribution of  $y_h(\mathbf{x}^0)$  is given by

$$p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l] = \int p[y_h(\mathbf{x}^0)|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\eta}] p(\mathbf{y}_{h,I}, \bar{\boldsymbol{\theta}}|\mathbf{y}_h, \mathbf{y}_l, \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta) d\mathbf{y}_{h,I} d\bar{\boldsymbol{\theta}} \quad (4.22)$$

where we use the notation  $\boldsymbol{\eta} = [\boldsymbol{\theta}, \bar{\boldsymbol{\theta}}]$ ,  $\boldsymbol{\theta} = [\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta]$ , and  $\bar{\boldsymbol{\theta}} = [\boldsymbol{\beta}_l, \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_l^2, \sigma_\delta^2, \sigma_\epsilon^2]$ .

### Point estimates for correlation parameters

To use point estimates for  $\theta_l$  and  $\theta_\delta$  in Eq. (4.22), we need to first establish these point estimates. For this, we use the marginal posterior distribution for  $\theta_l$  and  $\theta_\delta$ , which has the following expression,

$$\begin{aligned} p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l) &\propto \int_{\sigma_\delta^2, \sigma_\epsilon^2} p(\boldsymbol{\theta}_l) p(\boldsymbol{\theta}_\delta) (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\epsilon\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \\ &\times \left( \gamma_l + \frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8} \right)^{-\alpha_l - n_l/2} \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} \right] d\sigma_\delta^2 d\sigma_\epsilon^2 \end{aligned} \quad (4.23)$$

The derivation of this marginal distribution involves complicated mathematical operations to integrate out each of the parameters in  $\bar{\boldsymbol{\theta}} = [\boldsymbol{\beta}_l, \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_l^2, \sigma_\delta^2, \sigma_\epsilon^2]$ . The detailed derivations are presented in Appendix A.

Then the Maximum A Posteriori (MAP) estimate for  $\boldsymbol{\theta}$  can be established by

$$\boldsymbol{\theta}_{MAP} = \max_{\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta} p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l) \quad (4.24)$$

This optimization problem is complicated. To evaluate the objective function, which corresponds to the multi-dimensional integral with respect to  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  in Eq. (4.23), Eq. (4.23) is rewritten by multiplying and dividing the probability distribution  $p(\sigma_\delta^2, \sigma_\epsilon^2) = p(\sigma_\delta^2)p(\sigma_\epsilon^2)$  with  $p(\sigma_\delta^2) \propto (\sigma_\delta^2)^{-\alpha_\delta-1} \exp(-\gamma_\delta/\sigma_\delta^2)$  and  $p(\sigma_\epsilon^2) \propto (\sigma_\epsilon^2)^{-\alpha_\epsilon-1} \exp(-\gamma_\epsilon/\sigma_\epsilon^2)$ . That is, the integral is explicitly written as expectation with respect to  $p(\sigma_\delta^2, \sigma_\epsilon^2)$ , which allows the use of stochastic simulation to evaluate the integral. In the end, we consider the following equivalent optimization,

$$\begin{aligned}\boldsymbol{\theta}_{MAP} &= \max_{\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta} \int_{\sigma_\delta^2, \sigma_\epsilon^2} p(\boldsymbol{\theta}_l) p(\boldsymbol{\theta}_\delta) (\sigma_\delta^2)^{-(n_h - q_\rho)/2 + \alpha_\delta + 1} (\sigma_\epsilon^2)^{\alpha_\epsilon + 1} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \\ &\quad \times \left( \gamma_l + \frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8} \right)^{-\alpha_l - n_l/2} \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} + \frac{\gamma_\delta}{\sigma_\delta^2} + \frac{\gamma_\epsilon}{\sigma_\epsilon^2} \right] p(\sigma_\delta^2, \sigma_\epsilon^2) d\sigma_\delta^2 d\sigma_\epsilon^2 \\ &= \max_{\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta} \int_{\sigma_\delta^2, \sigma_\epsilon^2} F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2) p(\sigma_\delta^2, \sigma_\epsilon^2) d\sigma_\delta^2 d\sigma_\epsilon^2 = \max_{\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta} E[F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2)]\end{aligned}\tag{4.25}$$

The optimization problem in Eq. (4.25) can be solved by a sample based approach, i.e., *sample average approximation* (Ruszczynski and Shapiro 2003). We first generate  $K$  samples for  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  from the prior distribution  $p(\sigma_\delta^2, \sigma_\epsilon^2)$  and then evaluate the expectation of  $F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2)$  based on these sample, i.e.,  $E[F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2)] \approx \frac{1}{K} \sum_{k=1}^K F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_{\delta,k}^2, \sigma_{\epsilon,k}^2)$ . Therefore, the optimization problem in Eq. (4.25) can be replaced by

$$\boldsymbol{\theta}_{MAP} = \max_{\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta} \frac{1}{K} \sum_{k=1}^K F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_{\delta,k}^2, \sigma_{\epsilon,k}^2)\tag{4.26}$$

which can be solved using general nonlinear optimization algorithms. To address the estimation error of  $E[F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2)]$  for different values of  $\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta$  due to randomness in the samples for  $[\sigma_\delta^2, \sigma_\epsilon^2]$ , Common Random Numbers (CRNs) will be used (i.e., the same set of samples for  $[\sigma_\delta^2, \sigma_\epsilon^2]$  will be used for estimation of  $E[F(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \sigma_\delta^2, \sigma_\epsilon^2)]$  in the optimization in Eq. (4.26)).

### Sampling for other parameters

With point estimates for  $\theta_l$  and  $\theta_\delta$ , we can use Gibbs sampling to generate samples for other parameters where samples for  $\tilde{y}_h^i, \beta_l, \beta_\rho, \beta_\delta, \sigma_l^2$  can be directly generated from their conditional posterior distributions which are of standard forms. However, for  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$ , since their conditional posteriors are not of standard form, Metropolis-Hastings algorithm will be used to generate samples for  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  within each step of Gibbs sampling. In the end, we can generate samples from the joint posterior distributions for the model parameters, which can be used to establish the posterior predictive distribution for  $y_h(\mathbf{x}^0)$  (as shown in Eq. (4.22)).

### 4.5 Illustrative Example: Probabilistic Deformation Capacity of RC Columns

Accurate prediction of deformation capacity of reinforced concrete (RC) columns under repeated cyclic loading is critical for assessing performance of RC bridges under seismic loading (e.g., for predicting damage, establishing fragility curves, assessing failure probability). Traditional deformation capacity model for RC columns is derived based on the mechanics rules, and the uncertainties in model parameters are not taken into account. The deterministic model often gives conservative predictions for the deformation capacity (Gardoni et al. 2002). Some previous studies (Gardoni et al. 2002, 2003) have proposed probabilistic capacity models to account for the uncertainty in predicting the capacity of RC columns. These models are calibrated using existing experimental data. In this section, we will apply the multi-fidelity Gaussian process model to establish a probabilistic predictive model for the deformation capacity of RC columns and demonstrate the performance of the proposed multi-fidelity Gaussian process model. This application will integrate limited high-fidelity experimental data with large number of low-fidelity data (that is easier, cheaper to obtain compared to high-fidelity data), and will explicitly incorporate the cen-

sored data in the high-fidelity data and the uncertainties in the Gaussian process model parameters as well as the measurement errors in the data. In the end, the established predictive model can be used to predict the deformation capacity (or posterior distribution of the deformation capacity, i.e., probabilistic deformation capacity) for any new RC column that is not in the existing database.

Similar to Gardoni et al. (2002), we consider the deformation capacity of cantilever RC columns fixed at the base. The deformation capacity of the RC column is defined as the maximum drift ratio at the top of the column. Eleven inputs (capturing key geometric and material properties of the column) are considered, including  $x_1$  = dimension of column section  $D_{col}$ ,  $x_2$  = height of column  $H_c$ ,  $x_3$  = cover thickness of concrete  $c$ ,  $x_4$  = compressive strength of concrete  $f_c$ ,  $[x_5, x_6, x_7]$  = dimension, yield strength, and ratio of longitudinal reinforcement  $d_s, f_{ys}, \rho_s$ ,  $[x_8, x_9, x_{10}]$  = dimension, yield strength, ratio of volumetric transverse reinforcement  $d_{sv}, f_{yv}, \rho_{sv}$ , and  $x_{11}$  = axial load  $N$ . Therefore, in this example, we have  $n_x = 11$  inputs.

#### 4.5.1 High-fidelity database with censored data

The behavior of RC columns under repeated cyclic loading has been investigated using experimental tests by many researchers over the years. Many of the results have been collected at a database (see at <https://nisee.berkeley.edu/spd/index.html>). The high-fidelity data for this study will be established from this database. The latest database contains the experimental data of 253 rectangular columns and 163 circular columns under cyclic lateral loads. For this study, we focus on the circular columns. Some of the 163 circular columns were excluded from this study, including: (1) columns 13, 25, 34, 135 with varying axial loads; (2) column 38, 69, 75, 82 without spiral reinforcement; (3) columns 39, 40, 144-149 with 28-day concrete strength unreported; (4) columns 61-66 with square cross sections; (5) columns 141, 154 which are retrofitted or sliced columns.

Therefore, the experimental data of the remaining 139 columns are selected for this study. Among the 139 columns, 69 columns are removed in order to avoid ill-conditioned covariance matrices when building the Gaussian process model since some of these columns have close to identical inputs. The final high-fidelity database consists of 17 accurate responses (i.e., with direct measurement of the deformation capacity) and 53 censored responses (i.e., with measurement of the deformation capacity values before reaching failure). The detailed information about the final high-fidelity database is presented in Appendix B.

As can be seen, we only have a limited number of high-fidelity data. If relying on the high-fidelity data alone to establish a Gaussian process model, due to the small number of training data, the established Gaussian model is not expected to have high accuracy or generalize well to other inputs not in the database. Therefore, we use the multi-fidelity Gaussian process model to integrate this limited number of high-fidelity data with large number of low-fidelity data to inform a better Gaussian process model.

#### 4.5.2 *Low-fidelity database*

Though accurate, high-fidelity data from running physical experiments is expensive to establish. On the other hand, cheaper computational simulation or theoretical analysis requiring less resources can be applied to investigate the deformation capacity of RC columns; however, the results may not be as accurate as the high-fidelity experimental data. In this study, the low-fidelity data will be obtained from the combination of computational simulation and theoretical analysis. A theoretical model has been given by Gardoni et al. (2002) for calculating the deformation capacity of RC columns. Based on the model, the deformation capacity is decomposed into two parts, the elastic deformation  $\delta_y$  and the plastic deformation  $\delta_p$ , i.e.,  $\delta = \delta_y + \delta_p = (\Delta_y + \Delta_p)/H_c$ , where

$\Delta_y$  is the lateral displacement of the top when the column yields,  $\Delta_p$  is the lateral displacement of the top from when the column yields to a plastic hinge forms at the bottom.

The elastic displacement of the top can be obtained as (Gardoni et al. 2002; Priestley et al. 1996; Pujol 2002)  $\Delta_y = \phi_y l_{eff}^2 / 3 + (V_y H_c) / (G A_{ve}) + (\phi_y f_{ys} d_s H_c) / (8.64 \sqrt{f_c})$ , where  $\phi_y$  is the curvature of the bottom section when the column yields,  $l_{eff}$  is the effective height of the column which can be estimated by  $H_c + 0.022 f_{ys} d_s$ .  $V_y$  is the shear force when the column yields,  $G$  is the shear modulus of concrete and  $A_{ve}$  is the effective shear area which can be calculated by  $A_{ve} = (I_e / I) k_s A$ , where  $I_e$  and  $I$  are the effective moment of inertia and the gross moment of inertia of the cross section respectively,  $k_s$  is the cross-sectional shape factor, and  $A$  is the area of the gross section. The plastic deformation can be calculated as (Priestley et al. 1996)  $\Delta_p = (\phi_u - \phi_y) l_p H_c$ , where  $\phi_u$  is the curvature of the bottom section when the column fails,  $l_p$  the equivalent plastic hinge length which can be estimated by  $0.08 H_c + 0.022 f_{ys} d_s$ . Note that  $l_p$  should not be smaller than  $0.044 f_{ys} d_s$ .

The values of the parameters  $\phi_y$ ,  $\phi_u$ ,  $V_y$  and  $I_e$  need to be calculated from the moment-curvature relationship based on the cross section analysis for RC columns. In this study, the cross section analyses are done in OpenSees for all the columns in the training dataset. More specifically, the columns are modeled using fiber section elements, and pushover analyses are done for the columns to get the moment-curvature relationship of the bottom cross section. Among the parameters,  $\phi_y$ ,  $\phi_u$ ,  $V_y$  can be obtained directly according to the definition, while the effective moment of inertia  $I_e$  can be calculated by  $I_e = (M_{cr} / M_a)^3 I_g + [1 - (M_{cr} / M_a)^3] I_{cr} \leq I_g$ , where  $M_{cr} = (0.84 \sqrt{f_c}) I_g / (D_{col} / 2)$  is the cracking moment,  $M_a$  is the maximum moment and  $I_{cr}$  is the moment of inertia of cracked section.

#### 4.5.3 Implementation details

In this example, we select  $\ln(\delta)$  as the output, i.e.,  $y(\mathbf{x}) = \ln(\delta(\mathbf{x}))$ . When building the multi-fidelity Gaussian process model, for the high-fidelity data, we consider two cases. The first case uses all the high-fidelity data, i.e., including both the accurate responses and the censored responses. The corresponding high-fidelity data will be referred as “complete data”. The second case only uses the accurate responses. The corresponding high-fidelity data will be referred as “accurate data”. The performance of the multi-fidelity Gaussian process models built under these two cases will be investigated to illustrate the impact of including censored data. On the other hand, for the low-fidelity data  $\mathbf{x}_l$ , we know that  $\mathbf{x}_l$  is a superset of  $\mathbf{x}_h$ . When creating the low-fidelity data, except the inputs in  $\mathbf{x}_h$ , the rest of the low-fidelity data inputs are generated using Latin Hypercube Sampling (LHS). Linear basis functions are used in the Gaussian process models (e.g., for  $\mathbf{f}_l(\mathbf{x})$ ,  $\mathbf{f}_\rho(\mathbf{x})$  and  $\mathbf{f}_\delta(\mathbf{x})$ ). Take  $\mathbf{f}_l(\mathbf{x})$  for example, it corresponds to  $\mathbf{f}_l(\mathbf{x}) = [1, x_1, \dots, x_{11}]$ . Other basis functions can be used as well; alternatively, the basis functions can be optimally selected using optimization.

To evaluate the accuracy of the prediction model, error statistics are calculated using leave-one-out cross validation (LOOCV), which is performed as follows. First, the observations from the high-fidelity data set is sequentially removed, and the corresponding training data set will be composed of the remaining high-fidelity data and the original low-fidelity data. Then the Gaussian process model built using the corresponding training set is used to predict the response over the removed data point. The error between the prediction and the high-fidelity observation at the removed data point is then calculated. In the end, averaging errors over all the high-fidelity data points gives the overall LOOCV error. Here we use the mean error  $ME$  as the error statistics, and the LOOCV  $ME$  is given by:  $ME = \sum_{i=1}^{n_h} |y_h(\mathbf{x}_h^i) - E[\hat{y}_h(\mathbf{x}_h^i)]| / \sum_{i=1}^{n_h} |y_h(\mathbf{x}_h^i)|$ , where  $\hat{y}_h(\mathbf{x}_h^i)$

is the prediction from the established multi-fidelity Gaussian process model and  $E[\hat{y}_h(\mathbf{x}_h^i)]$  corresponds to the mean of the predictive distribution for  $\hat{y}_h(\mathbf{x}_h^i)$ . If the values for  $ME$  is small, the model fits the data well. For the “complete data” case, the responses of the censored data only correspond to lower bounds of the actual responses with the actual responses unknown. For calculating  $ME$  in this case, the censored responses will be artificially treated as the actual response so that we can establish the  $ME$  value for this case. The  $ME$  in this case needs to be interpreted accordingly.

It is expected that the number of low-fidelity data (i.e.,  $n_l$ ) will to some extent impact the accuracy of the established multi-fidelity Gaussian process model.  $n_l$  is selected to strike a balance between computational efficiency and model accuracy. While larger number of low-fidelity data might help improve the accuracy, the computational efficiency could reduce. To select  $n_l$ , we run the algorithm with different  $n_l$  values. The results indicated that when  $n_l$  is larger than 300, the  $ME$  has little change, and in the end  $n_l = 300$  is used. Regarding prior distributions for the unknown model parameters, we choose priors as mentioned earlier. More specifically, for Eq. (4.8), we determine  $\mathbf{m}_l$ ,  $\mathbf{m}_\rho$ ,  $\mathbf{m}_\delta$  using the least square regression approach discussed earlier. For the other parameters, we select  $[\alpha_l, \alpha_\delta, \alpha_\epsilon] = [2, 2, 2]$ ,  $[\gamma_l, \gamma_\delta, \gamma_\epsilon] = [1, 1, 1]$ ,  $[a_l, a_\delta] = [2, 2]$ , and  $[b_l, b_\delta] = [1, 1]$ . It should be noted that, the prior values are selected based on the consideration that the corresponding model parameters have reasonable mean values and also a large spread for the prior distributions. More specifically, to get a sense of the values for  $\sigma_l^2$ , we established a Gaussian process model for  $y_l(\mathbf{x})$  based on only low-fidelity data, which gives a point estimate of around 0.148. For  $\sigma_\delta^2 + \sigma_\epsilon^2$ , since we have  $\delta(\mathbf{x}) + \epsilon = y_h(\mathbf{x}) - \rho(\mathbf{x})y_l(\mathbf{x})$ , by assuming  $\rho(\mathbf{x})$  is zero we establish a Gaussian process model for  $\delta(\mathbf{x}) + \epsilon$  is estimated to be around 0.139. Then an inverse Gamma prior with shape parameter  $\alpha = 2$  and scale parameter  $\gamma = 1$  is selected for  $\sigma_l^2$ ,

$\sigma_\delta^2$ , and  $\sigma_\epsilon^2$ , which has mode of 1/3 (which is around twice of 0.148 and 0.139) and also ensures that the prior has large spread. For  $\theta_{l,i}$  and  $\theta_{\delta,i}$ , gamma distribution  $G(2,1)$  is selected as the prior. The mean and variance of this distribution are both 2 and the coefficient of variation is around 0.7 which means the distribution has a wide spread. In addition, the distribution covers a region from 0 to 10 and these values are typically used for  $\theta_i$  in Gaussian correlation functions. To establish point estimates for  $\theta_l$  and  $\theta_\delta$ ,  $K = 2000$  samples for  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  from their prior distributions are used in the optimization problem in Eq. (4.26). The proposed approach discussed in earlier sections is then used for sampling from the posterior distribution of the model parameters and establishing the predictive distributions for the outputs at new inputs.

#### 4.5.4 Posterior statistics of model parameters

Here we present the posterior statistics of model parameters established under the “complete data” case. By solving the optimization in Eq. (4.26) (i.e., maximizing the marginal posterior for  $\theta_l$  and  $\theta_\delta$ ), the MAP estimates of correlation parameters  $\theta_l$  and  $\theta_\delta$  are obtained, which are presented in Table 4.1. Using MAP estimates for  $\theta_l$  and  $\theta_\delta$ , we can then use Gibbs sampling discussed earlier to generate posterior samples for other parameters where samples for  $\tilde{y}_h^i, \beta_l, \beta_\rho, \beta_\delta, \sigma_l^2$  are directly generated from their conditional posterior distributions while samples for  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  are generated using Metropolis-Hastings algorithm. The posterior means, standard deviations and correlation matrices of  $\beta_l, \beta_\rho$  and  $\beta_\delta$  are then estimated based on the generated samples and are reported in Tables 4.2, 4.3 and 4.4, respectively.

Fig. 4.2 shows the prior and posterior distributions of parameters  $\sigma_l^2, \sigma_\delta^2$  and  $\sigma_\epsilon^2$ . Note that the posterior distributions are obtained using kernel density estimation based on the generated posterior samples for these parameters. The big differences between the priors and posteriors of  $\sigma_l^2, \sigma_\delta^2$  and  $\sigma_\epsilon^2$

**Table 4.1:** MAP estimates for  $\theta_l$  and  $\theta_\delta$ 

| Parameter       | Component           | Estimation |
|-----------------|---------------------|------------|
| $\theta_l$      | $\theta_{l1}$       | 0.536      |
|                 | $\theta_{l2}$       | 0.196      |
|                 | $\theta_{l3}$       | 0.046      |
|                 | $\theta_{l4}$       | 0.043      |
|                 | $\theta_{l5}$       | 0.084      |
|                 | $\theta_{l6}$       | 0.081      |
|                 | $\theta_{l7}$       | 0.118      |
|                 | $\theta_{l8}$       | 0.012      |
|                 | $\theta_{l9}$       | 1.444      |
|                 | $\theta_{l10}$      | 0.689      |
|                 | $\theta_{l11}$      | 4.802      |
| $\theta_\delta$ | $\theta_{\delta1}$  | 1.180      |
|                 | $\theta_{\delta2}$  | 1.270      |
|                 | $\theta_{\delta3}$  | 1.152      |
|                 | $\theta_{\delta4}$  | 0.180      |
|                 | $\theta_{\delta5}$  | 0.732      |
|                 | $\theta_{\delta6}$  | 0.317      |
|                 | $\theta_{\delta7}$  | 0.238      |
|                 | $\theta_{\delta8}$  | 0.413      |
|                 | $\theta_{\delta9}$  | 0.436      |
|                 | $\theta_{\delta10}$ | 0.971      |
|                 | $\theta_{\delta11}$ | 0.520      |

**Table 4.2:** Posterior statistics of  $\beta_l$ 

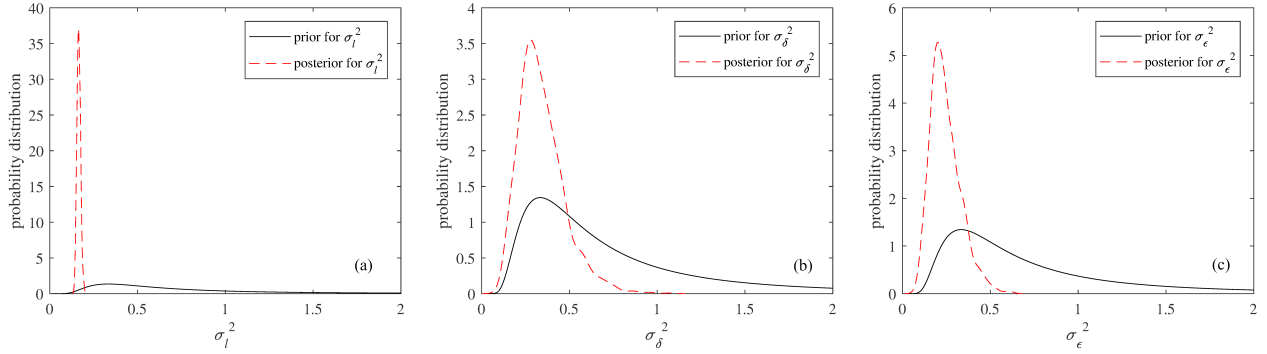
| Parameter     | Mean   | Standard deviation | Correlation Coefficient |              |              |              |              |              |              |              |              |               |               |               |
|---------------|--------|--------------------|-------------------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|---------------|---------------|---------------|
|               |        |                    | $\beta_{l1}$            | $\beta_{l2}$ | $\beta_{l3}$ | $\beta_{l4}$ | $\beta_{l5}$ | $\beta_{l6}$ | $\beta_{l7}$ | $\beta_{l8}$ | $\beta_{l9}$ | $\beta_{l10}$ | $\beta_{l11}$ | $\beta_{l12}$ |
| $\beta_{l1}$  | -0.002 | 0.001              | 1                       |              |              |              |              |              |              |              |              |               |               |               |
| $\beta_{l2}$  | -0.300 | 0.029              | 0.00                    | 1            |              |              |              |              |              |              |              |               |               |               |
| $\beta_{l3}$  | 0.381  | 0.026              | -0.01                   | -0.37        | 1            |              |              |              |              |              |              |               |               |               |
| $\beta_{l4}$  | 0.027  | 0.010              | 0.00                    | -0.16        | -0.01        | 1            |              |              |              |              |              |               |               |               |
| $\beta_{l5}$  | -0.178 | 0.020              | -0.04                   | 0.13         | 0.02         | -0.03        | 1            |              |              |              |              |               |               |               |
| $\beta_{l6}$  | 0.217  | 0.021              | -0.03                   | -0.20        | -0.03        | 0.08         | 0.01         | 1            |              |              |              |               |               |               |
| $\beta_{l7}$  | 0.158  | 0.020              | -0.03                   | -0.05        | 0.00         | 0.02         | -0.06        | 0.04         | 1            |              |              |               |               |               |
| $\beta_{l8}$  | -0.029 | 0.012              | 0.02                    | 0.05         | 0.01         | -0.01        | 0.01         | -0.01        | 0.01         | 1            |              |               |               |               |
| $\beta_{l9}$  | 0.006  | 0.003              | 0.04                    | -0.03        | -0.04        | -0.03        | 0.04         | -0.02        | -0.02        | 0.00         | 1            |               |               |               |
| $\beta_{l10}$ | 0.516  | 0.030              | -0.05                   | 0.08         | -0.08        | 0.00         | -0.03        | -0.03        | 0.02         | 0.03         | -0.01        | 1             |               |               |
| $\beta_{l11}$ | 0.484  | 0.028              | 0.03                    | 0.08         | -0.05        | -0.01        | -0.04        | -0.02        | -0.01        | -0.02        | -0.01        | 0.01          | 1             |               |
| $\beta_{l12}$ | 0.229  | 0.028              | 0.03                    | 0.37         | 0.06         | 0.00         | 0.29         | 0.03         | 0.02         | 0.11         | 0.01         | 0.04          | 0.01          | 1             |

**Table 4.3:** Posterior statistics of  $\beta_\rho$ 

| Parameter        | Mean   | Standard deviation | Correlation Coefficient |                 |                 |                 |                 |                 |                 |                 |                 |                  |                  |                  |
|------------------|--------|--------------------|-------------------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|-----------------|------------------|------------------|------------------|
|                  |        |                    | $\beta_{\rho1}$         | $\beta_{\rho2}$ | $\beta_{\rho3}$ | $\beta_{\rho4}$ | $\beta_{\rho5}$ | $\beta_{\rho6}$ | $\beta_{\rho7}$ | $\beta_{\rho8}$ | $\beta_{\rho9}$ | $\beta_{\rho10}$ | $\beta_{\rho11}$ | $\beta_{\rho12}$ |
| $\beta_{\rho1}$  | 0.725  | 0.198              | 1                       |                 |                 |                 |                 |                 |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho2}$  | 0.012  | 0.006              | -0.02                   | 1               |                 |                 |                 |                 |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho3}$  | 0.002  | 0.001              | 0.00                    | 0.03            | 1               |                 |                 |                 |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho4}$  | -0.013 | 0.006              | 0.05                    | 0.01            | 0.00            | 1               |                 |                 |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho5}$  | -0.001 | 0.001              | 0.01                    | 0.04            | -0.04           | -0.05           | 1               |                 |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho6}$  | -0.003 | 0.001              | 0.01                    | -0.03           | 0.02            | 0.03            | -0.01           | 1               |                 |                 |                 |                  |                  |                  |
| $\beta_{\rho7}$  | -0.011 | 0.005              | 0.00                    | 0.03            | -0.01           | -0.05           | 0.01            | -0.03           | 1               |                 |                 |                  |                  |                  |
| $\beta_{\rho8}$  | 0.005  | 0.002              | 0.00                    | 0.04            | -0.03           | -0.02           | 0.01            | 0.05            | -0.06           | 1               |                 |                  |                  |                  |
| $\beta_{\rho9}$  | -0.009 | 0.005              | 0.02                    | 0.05            | 0.01            | 0.04            | -0.02           | 0.02            | -0.01           | -0.01           | 1               |                  |                  |                  |
| $\beta_{\rho10}$ | -0.003 | 0.002              | 0.07                    | 0.01            | -0.01           | 0.00            | -0.01           | -0.04           | 0.03            | 0.00            | 0.04            | 1                |                  |                  |
| $\beta_{\rho11}$ | 0.011  | 0.006              | 0.03                    | 0.04            | 0.02            | -0.06           | 0.04            | -0.01           | 0.04            | 0.00            | 0.03            | 0.00             | 1                |                  |
| $\beta_{\rho12}$ | 0.022  | 0.011              | -0.06                   | -0.01           | -0.01           | 0.02            | -0.01           | 0.01            | 0.01            | 0.00            | 0.00            | 0.02             | -0.01            | 1                |

**Table 4.4:** Posterior statistics of  $\beta_\delta$ 

| Parameter           | Mean   | Standard deviation | Correlation Coefficient |                    |                    |                    |                    |                    |                    |                    |                    |                     |                     |                     |
|---------------------|--------|--------------------|-------------------------|--------------------|--------------------|--------------------|--------------------|--------------------|--------------------|--------------------|--------------------|---------------------|---------------------|---------------------|
|                     |        |                    | $\beta_{\delta 1}$      | $\beta_{\delta 2}$ | $\beta_{\delta 3}$ | $\beta_{\delta 4}$ | $\beta_{\delta 5}$ | $\beta_{\delta 6}$ | $\beta_{\delta 7}$ | $\beta_{\delta 8}$ | $\beta_{\delta 9}$ | $\beta_{\delta 10}$ | $\beta_{\delta 11}$ | $\beta_{\delta 12}$ |
| $\beta_{\delta 1}$  | 0.627  | 0.265              | 1                       |                    |                    |                    |                    |                    |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 2}$  | 0.012  | 0.006              | 0.02                    | 1                  |                    |                    |                    |                    |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 3}$  | -0.085 | 0.047              | -0.03                   | -0.02              | 1                  |                    |                    |                    |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 4}$  | 0.539  | 0.252              | 0.76                    | -0.01              | -0.11              | 1                  |                    |                    |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 5}$  | 0.159  | 0.134              | 0.00                    | -0.01              | -0.04              | -0.01              | 1                  |                    |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 6}$  | -0.229 | 0.128              | -0.61                   | -0.03              | -0.14              | -0.61              | -0.01              | 1                  |                    |                    |                    |                     |                     |                     |
| $\beta_{\delta 7}$  | 0.352  | 0.181              | 0.03                    | 0.02               | -0.02              | 0.09               | 0.05               | -0.16              | 1                  |                    |                    |                     |                     |                     |
| $\beta_{\delta 8}$  | -0.359 | 0.161              | 0.10                    | 0.01               | 0.10               | 0.06               | 0.07               | -0.03              | -0.04              | 1                  |                    |                     |                     |                     |
| $\beta_{\delta 9}$  | 0.139  | 0.072              | 0.06                    | 0.02               | -0.04              | -0.08              | 0.00               | -0.10              | -0.02              | -0.05              | 1                  |                     |                     |                     |
| $\beta_{\delta 10}$ | -0.298 | 0.141              | 0.14                    | 0.00               | 0.01               | 0.18               | -0.34              | -0.02              | -0.17              | -0.01              | 0.01               | 1                   |                     |                     |
| $\beta_{\delta 11}$ | 0.277  | 0.196              | 0.25                    | 0.04               | 0.08               | 0.11               | -0.40              | -0.03              | 0.15               | -0.30              | -0.12              | 0.35                | 1                   |                     |
| $\beta_{\delta 12}$ | 0.170  | 0.107              | -0.06                   | -0.02              | 0.00               | 0.14               | -0.07              | 0.09               | 0.04               | 0.00               | 0.00               | 0.14                | 0.19                | 1                   |

**Figure 4.2:** Prior and posterior distributions for  $\sigma_l^2$ ,  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$ .

indicate that the training data effectively updated the probability distributions of these parameters.

We can also observe that most proportions of the posterior distributions of these parameters are concentrated in small regions between 0 and 1, which indicates small variances. The posterior sample means of  $\sigma_l^2$ ,  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  are 0.165, 0.342 and 0.245 respectively. The posterior sample standard deviations of  $\sigma_l^2$ ,  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$  are 0.011, 0.131 and 0.087 respectively.

#### 4.5.5 Posterior predictive distributions for the outputs

Fig. 4.3 (a) shows a comparison between the mean predictions by the multi-fidelity Gaussian process model built using “complete data” and high-fidelity observations of the deformation capacity of the experimental columns. The censored observations and accurate observations are also labeled in the figure. Fig. 4.3 (b) shows the same comparison but with the multi-fidelity Gaussian process model built using only “accurate data”. From Fig. 4.3 (a), we can see that the predictions for most of the accurate data are above the 1:1 line which indicates over-prediction for those data. This is expected since the prediction model is constructed using the complete high-fidelity data set where the accurate data only accounts for one fourth of the total data. For censored data points, reasonable predictions should be larger than the observed values considering that the censored data correspond to lower bound data. In Fig. 4.3 (a), 47 out of 53 censored data points are distributed over the 1:1 line which indicates a relatively good fit.

In comparison, the prediction model in Fig. 4.3 (b) fits the accurate data well. However, the predictions over the censored data points have large errors. Unlike the case in Fig. 4.3 (a), only 33 out of 53 censored data points lie over 1:1 line. For those below the 1:1 line, the deformation capacities are largely under-predicted, since those observations are already lower bounds of the actual deformation capacities. The multi-fidelity Gaussian process model built using only “accurate data” leads to predictions even lower than the lower bounds for some of the data points, indicating large under-prediction of the actual deformation capacities.

To quantitatively compare the performance of the two models, the LOOCV  $ME$ 's are calculated. Note that we do not know the actual responses over the censored data points, but reasonable predictions over censored data points should be larger than the corresponding observations (which

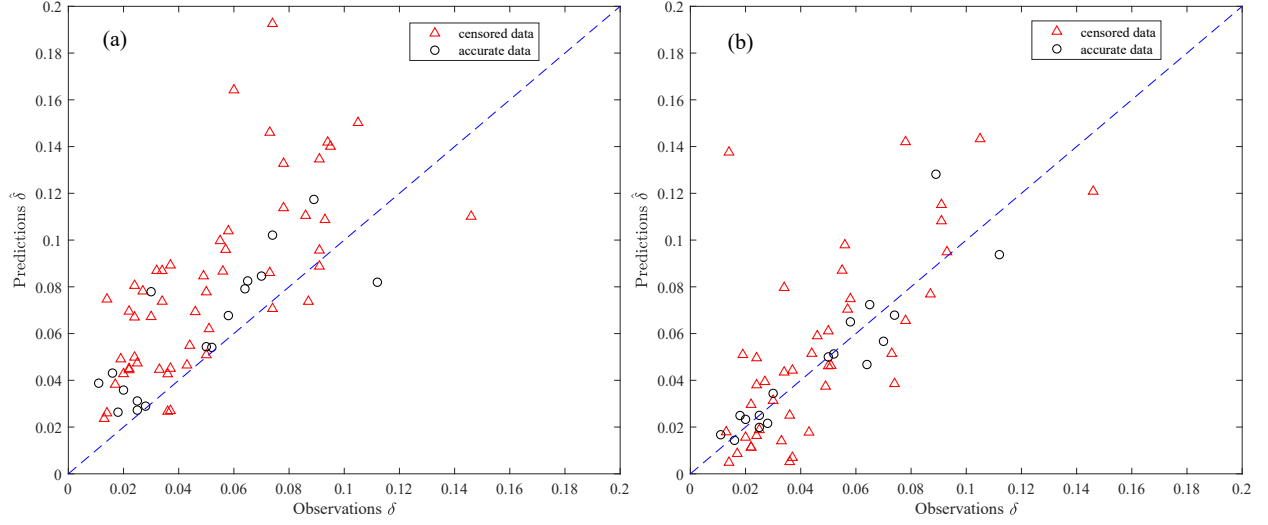
**Table 4.5:** Comparison of the mean errors

| Model type          | Mean errors over<br>censored data points | Mean errors over<br>accurate data points |
|---------------------|------------------------------------------|------------------------------------------|
| Complete data model | 0.0266 (0.0841)                          | 0.3547 (0.2772)                          |
| Accurate data model | 0.1072 (0.1264)                          | 0.1770 (0.2584)                          |

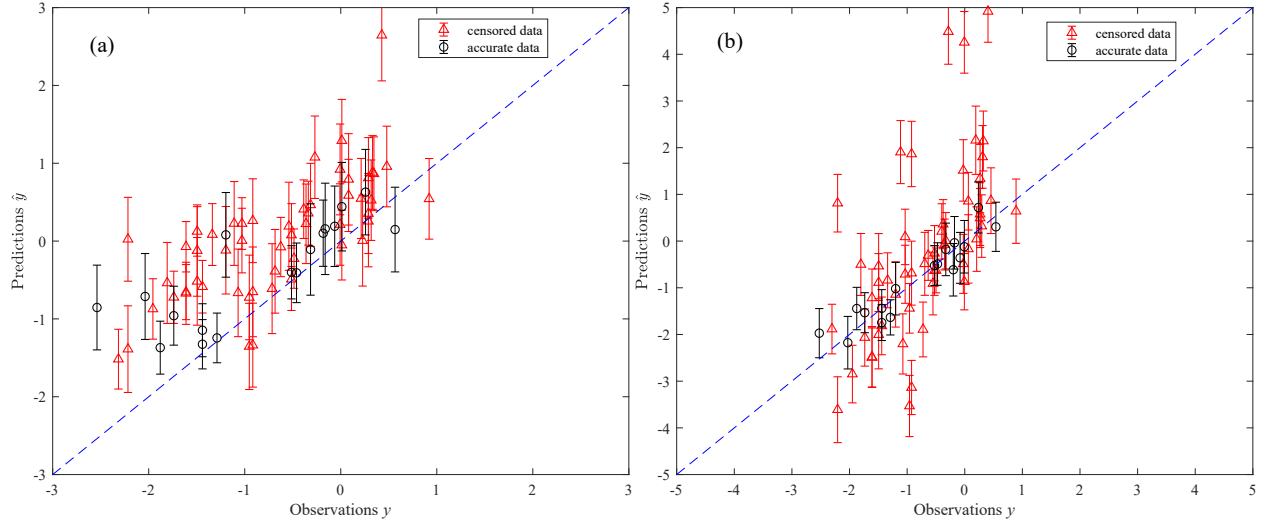
are low bounds for the actual responses). Therefore, in calculating the  $ME$ , for the censored data points, we assume the errors are zero for those whose mean predictions are larger than the observations. Table 4.5 gives the mean errors of the mean predictions over the high-fidelity censored data points and the high-fidelity accurate data points. As expected, the “complete data” model gives much smaller mean error over the censored data points than the “accurate data” model. In addition, as mentioned earlier, we have 69 column test results which were initially removed to avoid ill-conditioned covariance matrices when establishing the model. These data can actually be used as a further validation set for model predictions. The prediction errors over these 69 columns are also calculated and reported in Table 4.5 (i.e., inside the parenthesis). Similar trends for the  $ME$ s are observed though with slightly larger values for the censored data points.

By comparing Fig. 4.3 (a) and (b) and observing the results in Table 4.5, we can see that it is important to include the censored data in building the multi-fidelity Gaussian process model, since they also provide information regarding the underlying model. This is especially true when only a limited number of accurate data are available in which case it is critical to extract information also from the censored data. Relying only on limited accurate data would likely to lead to Gaussian process models that have low accuracy and cannot generalize well to inputs that are not in the training data.

In addition to the mean prediction, with samples from the posterior predictive distribution, we also calculated the standard deviation. Fig. 4.4 presents the mean predictions with  $\pm$  one standard

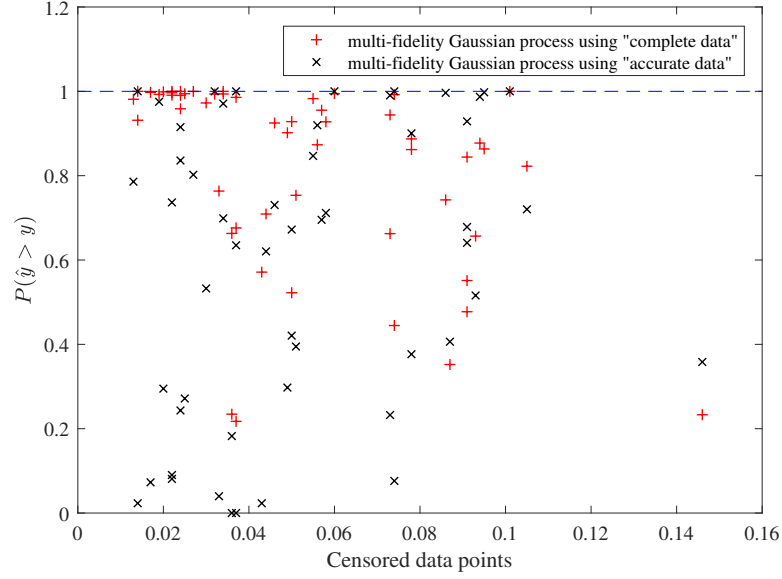


**Figure 4.3:** Comparison between the observations and the mean predictions by the multi-fidelity Gaussian process model established using (a) “complete data”, (b) “accurate data”.



**Figure 4.4:** Comparison between the observations and the predictive distributions (showing  $\mu \pm \sigma$ ) by the multi-fidelity Gaussian process model established using (a) “complete data”, (b) “accurate data”.

deviation (i.e.,  $\mu \pm \sigma$ ). Here the plotted predictions are the normalized natural logarithms of the original responses (i.e., the deformation capacity). Fig. 4.4 gives an indication of statistically how likely the model predictions will be larger than the observations for both censored data and accurate data. In order to quantify the probability that the Gaussian process model gives reasonable predictions over censored data points, the probabilities that the predictions will exceed observations over



**Figure 4.5:** Probabilities that predictions by the multi-fidelity Gaussian process models will exceed the observations for the censored data points.

all the censored data points are calculated (i.e., using samples from the posterior predictive distribution). The results are shown in Fig. 4.5. For the Gaussian process model built using only accurate data, almost half of the predictions have less than 50% probability to give predictions that exceed the corresponding lower bound observations. By contrast, for the Gaussian process model built using complete data, over half of the predictions have more than 80% probability to give predictions that exceed the corresponding lower bound observations. The comparison further demonstrates that Gaussian process model built using data including the censored data can give more reasonable predictions over censored data points than model built only using accurate data.

## 4.6 Conclusions

This chapter proposed a general multi-fidelity Gaussian process model to integrate data with different levels of accuracy. The idea is to leverage the trend information provided by a large number of low-fidelity data and the accuracy provided by a small number of expensive high-fidelity data to establish a model that is better than Gaussian process models built using either the high-

fidelity data or low-fidelity data alone. In calibration of the multi-fidelity Gaussian process model, censoring in the high-fidelity data, which is common in experimental testings and high-fidelity numerical simulations (e.g., with partially converged results), and the measurement errors were explicitly considered. Posterior distributions of the model parameters were established to explicitly take into account the uncertainties in the model parameters. To address the computational challenges in estimating the likelihood function for censored data, a data augmentation algorithm was proposed, which artificially treats the censored data as unknown parameters in addition to the original model parameters. Then closed form conditional posteriors were derived for these parameters, which were used in the context of Gibbs sampling to facilitate efficient sampling from the posterior distributions for the model parameters. In the end, the posterior samples for the model parameters were used to establish the posterior statistics for the output predictions at new inputs.

As an illustrative example, the proposed multi-fidelity Gaussian process model was applied to establish predictive model for the deformation capacity of RC columns. The results showed that it is important to explicitly integrate the information from censored data in establishing the multi-fidelity Gaussian process model, rather than only relying on accurate data, especially when censored data account for a significant portion of the high-fidelity data. Overall, the proposed multi-fidelity Gaussian process model is expected to be useful for problems when (i) only limited number of high-fidelity data are available, (ii) censored data account for a significant portion of the high-fidelity data, (iii) low-fidelity data can be established relatively efficiently, and (iv) uncertainties in the model parameters need to be explicitly considered.

## CHAPTER 5:

### DIMENSION REDUCTION ASSISTED GAUSSIAN PROCESS MODEL FOR HIGH-DIMENSIONAL INPUTS AND ITS APPLICATION IN OPTIMIZATION<sup>2</sup>

#### 5.1 Introduction

Building accurate Gaussian process models for problems with high-dimensional inputs typically faces significant challenges or is even computationally prohibitive due to the curse of dimensionality (Viana et al. 2010). As a result, surrogate model based optimization for problems with high-dimensional inputs (i.e., design variables) is still difficult, although the surrogate model can help reduce the computational cost in evaluating the high-fidelity model, not to mention the computational challenges due to large design space. To address the challenges in training the Gaussian process model stemming from high dimensionality of inputs, here we propose a dimension reduction assisted Gaussian process model. In this chapter, to demonstrate the performance of the proposed model, it is developed in the context of optimization for problems with high-dimensional design variables. To reduce the dimension of the high-dimensional design variables, first a set of reference inputs that have better (or more desirable) objective function values are established. Then principal component analysis (PCA), as a dimension reduction technique, is applied to this database to establish a low-dimensional representation of the design variables in latent design space. To reduce the computational effort in calculation of objective function values for each design, a Gaussian process model is established with respect to the low-dimensional latent design variables. This surrogate model is then used in place of the original high-fidelity model to facilitate efficient optimization. The low-dimensionality of the latent design variables also facilitates much easier

---

<sup>2</sup>This chapter is adapted from a published paper by the author (Li et al. 2019b).

optimization. To improve the accuracy of the surrogate model and facilitate efficient and effective search for the optimal, an adaptive approach that sequentially add more support points for the surrogate model is used. Once the optimal is established in the latent space, the original design variable can be easily established through the PCA transformation. Overall, the proposed approach facilitates efficient optimization for problems with high-dimensional design variables. Note that for the optimization problems, here the focus is on high-dimensional discrete/binary design optimization, which is typically more challenging than optimization problems with continuous design variables and more importantly there is no research yet on integrating PCA with Gaussian process model for high-dimensional discrete/binary optimization. But the proposed dimension reduction and surrogate based optimization framework is general and can also be applied to problems with continuous or other discrete design variables. The effectiveness and great efficiency of the proposed approach are verified through an example on topology optimization of 2D periodic structures to maximize the frequency bandgaps.

## 5.2 Problem Formulation: Optimization for Problems with High-dimensional Inputs

For a general optimization problem, assume the input is represented by a vector  $\mathbf{b} = [b_1, \dots, b_i, \dots, b_{n_b}]$  where  $b_i$  is the  $i^{th}$  design variable. The input might be continuous or discrete, i.e.,  $b_i$  could be continuous values or discrete values. A special case is that the design variables vector  $\mathbf{b}$  is a binary vector where  $b_i$  takes values of either 0 or 1, and this dissertation mainly focuses on the binary design variables. Typically, a general optimization problem can be defined as

$$\mathbf{b}^* = \operatorname{argmax} f(\mathbf{b}) \quad s.t. : C(\mathbf{b}) > 0 \quad (5.1)$$

where  $f(\mathbf{b})$  is the objective function, and  $C(\mathbf{b}) > 0$  is the constraint on the design variable  $\mathbf{b}$ . When the objective function is expensive to evaluate, the optimization problem is a computationally challenging problem since a large number of evaluations of the objective function are usually required for finding the optimum. Furthermore, if  $n_b$  is a large number (e.g.,  $> 100$ ),  $\mathbf{b}$  corresponds to high-dimensional design variables and the optimization becomes a high-dimensional optimization problem which is a more difficult task due to the extremely large design space.

To address the challenges stemming from high computational effort and high dimensionality of the design variables and to facilitate efficient optimization for problems with high-dimensional design variables, this chapter proposes an efficient dimension reduction and surrogate based optimization approach.

### 5.3 Low-dimensional Representation of High-dimensional Inputs

#### 5.3.1 Set of reference inputs

To reduce the dimensionality of the design problem, we propose to establish a low-dimensional representation of  $\mathbf{b}$ . For this purpose, we first collect a set of  $m$  reference inputs. We call this set the *reference set*, which can be represented through the reference set matrix  $\mathbf{B} = [\mathbf{b}^1, \dots, \mathbf{b}^m]^T$  with dimension  $m \times n_b$ . As to the selection of the reference set, since the goal of the optimization in Eq. (5.1) is find the design that maximizes the objective function, intuitively it makes sense to include design variables that have better performance (i.e., larger objective function values). With this idea in mind, we can first generate  $m_0$  basic inputs using Latin Hypercube Sampling (LHS), and then pick the top  $m$  inputs based on the objective function values. These  $m$  inputs will form the reference set. Then a low-dimensional representation of  $\mathbf{b}$  is established by exploring correlation

characteristics in the set of reference inputs  $\mathbf{B}$ , and here we use principal component analysis (PCA) for this purpose.

As a popular dimension reduction technique, PCA finds a low-dimensional representation, typically called *principal components* or *latent input/outputs*, of the high-dimensional data. PCA discovers the linear projections of the data (e.g.,  $\mathbf{B}$ ) with maximum variance, or equivalently, the lower dimensional subspace that yields the minimum squared reconstruction error (Schein et al. 2003; Jolliffe 2002). It has been used to reduce output dimensions to facilitate surrogate models for high-dimensional outputs (Jia and Taflanidis 2013; Jia et al. 2016). Here we propose to use PCA to reduce the dimension of design variables (inputs) by applying PCA to the reference set of inputs (i.e., the matrix  $\mathbf{B}$ ). For the case of binary input, since  $\mathbf{B}$  corresponds to a binary matrix, instead of linear PCA, we use logistic PCA, which has been found to be better suited to reconstruction of binary data than linear PCA. In the end, we can establish a transformation between high-dimensional binary design variables  $\mathbf{b}$  and the low-dimensional continuous latent design variables.

### 5.3.2 Dimension reduction by logistic PCA

Given binary data matrix  $\mathbf{B}$ , a low-dimensional representation that maximizes the log-likelihood of the data matrix  $\mathbf{B}$  can be established. More specifically, let  $\mathbf{Z}$  with dimension  $m \times n_b$  represent the corresponding log-odds matrix for  $\mathbf{B}$ . In the end, logistic PCA establishes a low-dimensional representation of  $\mathbf{B}$ ,

$$\mathbf{Z} = \mathbf{X}\mathbf{P} + \Delta \quad (5.2)$$

where  $\mathbf{X}$  is the  $m \times n_x$  coefficient matrix or the latent design variables matrix, and  $n_x$  is the number of latent design variables, where typically  $n_x \ll n_b$ ; and  $\mathbf{X}$  corresponds to a low-dimensional

representation of  $\mathbf{Z}$ ;  $\mathbf{P}$  is the  $n_x \times n_b$  projection matrix with each column vector corresponding to a basis vector; and  $\Delta$  is the  $m \times n_b$  bias matrix with each row corresponding to the same  $1 \times n_b$  bias vector  $\delta$ .

Note that the low-dimensional representation  $\mathbf{X}$  takes continuous real values, which means  $\mathbf{Z}$  also takes continuous real values. To convert  $\mathbf{Z}$  to the corresponding binary matrix  $\mathbf{B}$ , we can simply use a thresholding approach. More specifically, if the value of elements in  $\mathbf{Z}$  is larger than a threshold  $z_0$ , then the corresponding element in  $\mathbf{B}$  will take value of 1, otherwise, it will take value of 0. This threshold  $z_0$  can be established for a given reference set matrix and for selected value of  $n_x$  by choosing a  $z_0$  value that minimizes the overall error rates (i.e., misclassification rates between the reconstructed binary matrix and the original binary reference set matrix).

The relationship in Eq. (5.2) essentially defines a transformation between low-dimensional latent design variables  $\mathbf{x}$  and original high-dimensional design variables  $\mathbf{z}$  or  $\mathbf{b}$ . More specifically, we can write

$$\mathbf{z} = \mathbf{xP} + \delta \quad (5.3)$$

where  $\mathbf{x}$  is the  $1 \times n_x$  vector of latent design variables. Using the threshold  $z_0$ , we can convert  $\mathbf{z}$  to the corresponding vector of high-dimensional binary design variables  $\mathbf{b}$  through

$$b_i = I_F(z_i); \quad i = 1, \dots, n_b \quad (5.4)$$

where  $I_F(\cdot)$  is the indicator function,  $I_F(z_i) = 1$ , if  $z_i > z_0$ ; 0, otherwise. In the end, we can define a transformation

$$\mathbf{b} = T(\mathbf{z}) = T(\mathbf{x}\mathbf{P} + \boldsymbol{\delta}) \quad (5.5)$$

where  $T(\cdot)$  (i.e., defined by  $\mathbf{P}$  and  $\boldsymbol{\delta}$ ) is the transformation from latent design variables  $\mathbf{x}$  to original design variables  $\mathbf{b}$ .

Note that when more latent design variables are used (i.e., larger values for  $n_x$ ), Eq. (5.2) will create a more accurate reconstruction of the original reference set or  $\mathbf{Z}$ ; however, from a transformation perspective, once the number of latent design variables  $n_x$  is chosen, regardless it is a large number or small number, the transformation  $T(\cdot)$  in Eq. (5.5) will correspond to a specific transformation from latent space to the original design space. Later, we will use this transformation to convert from latent space to original space. Note that, in the current context, the reference set is only used to establish a certain transformation, therefore, maintaining a high reconstruction accuracy of the original reference set is desirable but not a big concern.

#### 5.4 Optimization in Low-dimensional Continuous Design Space

In the end, we can define an equivalent optimization problem with respect to the low-dimensional latent design variables  $\mathbf{x}$

$$\mathbf{x}^* = \operatorname{argmin} y(\mathbf{x}) \quad s.t. : C(\mathbf{x}) > 0 \quad (5.6)$$

Here, besides reducing the dimension of the design variables, for convenience we also write the maximization problem in Eq. (5.1) in its minimization form by setting  $y(\mathbf{x}) = -f(\mathbf{x})$ . In the end, we want to solve the equivalent minimization problem in Eq. (5.6).

Compared to the original optimization formulation in Eq. (5.1), the formulation in Eq. (5.6) has two attractive properties. First, the latent design variables  $\mathbf{x}$  take continuous real values, rather than discrete values. Second, the latent design variables have much lower dimension than the original design variables, i.e.,  $n_x \ll n_b$ . Both of these properties facilitate less challenging optimization compared to the original optimization in Eq. (5.1). In the end, it is much more efficient to explore different designs in the low-dimensional latent space, compared to the original high-dimensional design space. Surrogate based optimization will be used to substantially improve the optimization efficiency, which is discussed in the next section. In the end, once  $\mathbf{x}^*$  is established, the corresponding optimal design  $\mathbf{b}^*$  can be simply established as  $\mathbf{b}^* = T(\mathbf{x}^* \mathbf{P} + \boldsymbol{\delta})$  using Eq. (5.5).

## 5.5 Surrogate based Optimization with Adaptive Design of Experiment

For the optimization in Eq. (5.6), for given value of  $\mathbf{x}$ , we need to repeatedly calculate the corresponding objective function value using expensive high-fidelity models. To alleviate the computational effort in this calculation, Gaussian process model is used. It not only gives the prediction (see Eq. (2.8)) but also the associated uncertainty (see Eq. (2.9)), namely the local variance of the prediction error (Jia and Taflanidis 2013). This local variance will be explicitly incorporated in guiding adaptive design of experiments to adaptively improve the accuracy of the surrogate model and facilitating effective optimization, which will be discussed later. For notation simplicity purpose, the predictive mean and variance are represented by  $\hat{y}(\mathbf{x})$  and  $\hat{\sigma}^2(\mathbf{x})$  in this chapter.

To build/calibrate the surrogate model, we first calculate the objective function values for  $n$  different the latent design variables  $\{\mathbf{x}^h; h = 1, \dots, n\}$  and establish the corresponding outputs  $\{y(\mathbf{x}^h); h = 1, \dots, n\}$ . Using the training set  $\{\mathbf{x}^h, y(\mathbf{x}^h); h = 1, \dots, n\}$ , a Gaussian process model can be established following Chapter 2. Then intuitively it can be directly used for the optimization

in Eq. (5.6), i.e., replace  $y(\mathbf{x})$  with  $\hat{y}(\mathbf{x})$  from the predictor. Considering the high efficiency of the established surrogate model, any global optimization algorithm can be used to find the corresponding minimum. However, this could easily lead to a local minimum. This is because this way of search essentially assumes that the predictor is an accurate prediction of the actual output values, i.e., it puts too much emphasis on exploiting the predictor but puts no emphasis on exploring points where we are uncertain (Jones et al. 1998). To address this, we need to sample more at points where the uncertainty is high. The level of uncertainty can be measured by the local variance given in Eq. (2.9). To adaptively improve the accuracy of the surrogate model and also the performance of the optimization, the expected improvement (EI) function can be used, which balances exploitation and exploration (or local and global search) (Jones et al. 1998; Forrester et al. 2008; Forrester and Keane 2009).

Let  $y_{min}$  be the current minimum value out of the sample points that have been evaluated. Then the improvement of the objective function at  $\mathbf{x}$  can be written as

$$I(\mathbf{x}) = \max(y_{min} - y(\mathbf{x}), 0) \quad (5.7)$$

Since before we sample at  $\mathbf{x}$ , we do not know the value of  $y(\mathbf{x})$  or we are uncertain about this value. To model our uncertainty about  $y(\mathbf{x})$ , we can treat  $y(\mathbf{x})$  as the realization of a normally distributed random variable with mean given by the predictor  $\hat{y}(\mathbf{x})$  and standard deviation given by  $\hat{\sigma}(\mathbf{x})$ . Taking expectation of  $I(\mathbf{x})$  with respect to this uncertainty gives the expected improvement, which has the following closed form (Jones et al. 1998; Forrester et al. 2008; Forrester and Keane 2009)

$$E[I(\mathbf{x})] = (y_{min} - \hat{y}(\mathbf{x})) \Phi\left(\frac{y_{min} - \hat{y}(\mathbf{x})}{\hat{\sigma}(\mathbf{x})}\right) + \hat{\sigma}(\mathbf{x}) \phi\left(\frac{y_{min} - \hat{y}(\mathbf{x})}{\hat{\sigma}(\mathbf{x})}\right) \quad (5.8)$$

where  $\phi(\cdot)$  and  $\Phi(\cdot)$  are the probability distribution function (PDF) and cumulative distribution function (CDF) of the standard normal distribution.

Then the sample point  $\mathbf{x}$  that maximizes  $E[I(\mathbf{x})]$  can be added to the existing training set. After we evaluate the actual value  $y(\mathbf{x})$  at  $\mathbf{x}$  using the adopted numerical model, we can then update (i)  $y_{min} = \min(y_{min}, y(\mathbf{x}))$ , and (ii) the Gaussian process surrogate model. In addition to adding one sample point in each iteration, in general multiple sample points can be added in each iteration (e.g., make use of parallel computing) (Zhang et al. 2017). This adaptive addition of new sample points can be carried out iteratively until some convergence criterion is reached, e.g.,  $E[I(\mathbf{x})]/|y_{min}| < \epsilon$ , which means the algorithm will stop when the expected improvement is less than a certain percentage of the current best value. Other convergence criteria can be also used, e.g., set an upper limit on the number of iterations, or total number of model evaluations. The above described surrogate model based optimization is often referred to as *efficient global optimization* (EGO). Once the optimal is established in the latent space, the corresponding original design variable can be easily established through the PCA transformation.

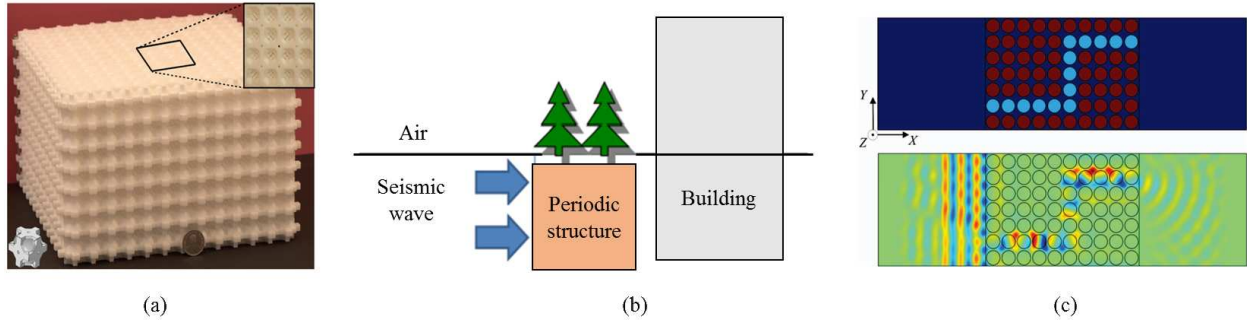
## 5.6 Illustrative Example: Topology Optimization of Periodic Structures

To illustrate the performance of the proposed algorithm, it is applied to topology optimization of 2D bi-component periodic structure with square lattice (i.e., shown in Figure 5.2(a),  $a_x = a_y = a$ ).

### 5.6.1 Motivation

Periodic structures, known as phononic crystals or elastic metamaterials, have the special dispersion frequency bandgap property (Li et al. 2017; Zhang et al. 2018; Wang and Kang 2019;

Cheng et al. 2018b,a). In some frequency regions (passing bands), waves can propagate in periodic structures freely; while in other frequency regions (stop bands), waves are forbidden. This filtering effect has many potential applications in civil and mechanical engineering, such as elastic filter, sound-proof materials, vibration isolation structures, seismic isolation systems and so on. Several applications of periodic structures are shown in Figure 5.1.



**Figure 5.1:** Periodic structures for (a) sound and vibration reduction (adapted from (Bilal et al. 2018)), (b) seismic isolation (adapted from (Kim and Das 2012)), and (c) wave guide (adapted from (Vasseur et al. 2011)).

For application purpose, periodic structures with special frequency bandgap characteristics are usually desired. To achieve this goal, topology optimization method has drawn a lot of research attention. Sigmund and Søndergaard Jensen (2003) first introduced the topology optimization method to optimize phononic bandgap crystals in 2003. After that, many works on the topic of topology optimization of frequency bandgap have been reported (Halkjær et al. 2006; Zhong et al. 2006; Dahl et al. 2008; Dong et al. 2014; Yi and Youn 2016; Xie et al. 2017). Recently, Bacigalupo et al. (2017, 2019) proposed a parametric optimization framework to solve bandgap maximization problems for periodic materials with beam lattice microstructure. Although it has the same objective with topology optimization, parametric optimization tries to find optimal parameters for cellular microstructure on the basis of an already fixed topology (Bacigalupo et al. 2017, 2019).

From this perspective, parametric optimization can be applied successively after topology optimization to further seek for the optimal material distribution of periodic cells, i.e., to tune a small set of parameters of an optimal topology obtained by topology optimization. Overall, methods for topology optimization can be classified into two groups: gradient-based topology optimization (GTO) method and non-gradient-based topology optimization (NGTO) method (Gazonas et al. 2006). For GTO methods, typically sensitivity calculation will be performed first to evaluate the influences of the system parameters on the responses, then the optimization algorithm is used to calculate the optimal solution. However, in many cases the gradient information is unavailable or unreliable to compute. To deal with these problems, NGTO methods have been developed. But, as pointed out in many works these NGTO methods usually encounter computational efficiency problems, especially when the dimension of the design variables is high. The advantages and the disadvantages of GTO and NGTO methods have been discussed in detail in Yi and Youn (2016). Overall, for both the GTO and NGTO methods, topology optimization typically involves many design variables corresponding to discretization of the unit cell into many pixels/elements with each pixel/element having specific material properties. This creates computational challenges for both optimization (i.e., corresponding to high-dimensional optimization) and the calculation of frequency bandgaps for a given design through finite element method or other numerical methods (i.e., finer discretization typically requires significantly higher computational effort).

To address the challenges stemming from high computational effort and high dimensionality of the design variables and to facilitate efficient topology optimization of periodic structures, this chapter applies the proposed efficient dimension reduction and surrogate based approach for topology optimization and demonstrates its performances. The approach in this application is named as Dimension REduction and Surrogate based Topology Optimization (DRESTO).

### 5.6.2 Frequency bandgap of periodic structures

For the 2D periodic structure shown in Figure 5.2(a), the corresponding basic unit cell is shown in Figure 5.2(b) with periodic constant vector  $\mathbf{A} = (a_x, a_y)$ . The periodic structure consists of two materials (i.e., material for the inclusion, shown in dark color, and material for the matrix, shown in gray color). Under the assumption of elastic, isotropic, small deformation and ignoring the effect of damping, the governing equation of the considered periodic structure can be given as

$$\rho(\mathbf{r}) \frac{\partial^2 \mathbf{u}}{\partial t^2} = \nabla [(\lambda(\mathbf{r}) + 2\mu(\mathbf{r}))(\nabla \cdot \mathbf{u})] - \nabla \times [\mu(\mathbf{r}) \nabla \times \mathbf{u}] \quad (5.9)$$

where  $\mathbf{u} = (u_x, u_y, u_z)$  is the displacement vector,  $\mathbf{r} = (x, y, z)$  is the coordinate vector;  $\lambda$  and  $\nu$  are the Lamé's constants,  $\rho$  is the mass density;  $\nabla$  is the Laplace operator,  $t$  is time. According to the Bloch-Floquet theorem, the solution of Eq. (5.9) can be written as

$$\mathbf{u}(\mathbf{r}, t) = \mathbf{u}_{\mathbf{K}}(\mathbf{r}) e^{i(\mathbf{K} \cdot \mathbf{r} - \omega t)} \quad (5.10)$$

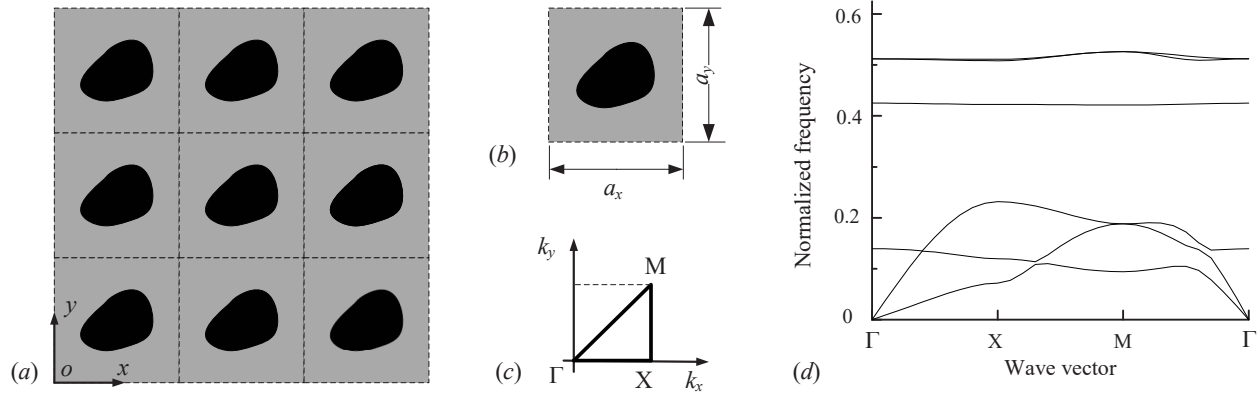
where  $\mathbf{K} = (k_x, k_y)$  is the wave vector and  $\mathbf{u}_{\mathbf{K}}(\mathbf{r})$  is the wave amplitude function. In particular, the wave amplitude function is a periodic function of the periodic constants, i.e.,  $\mathbf{u}_{\mathbf{K}}(\mathbf{r}) = \mathbf{u}_{\mathbf{K}}(\mathbf{r} + \mathbf{A})$ , which can be plugged into Eq. (5.10) to give

$$\mathbf{u}(\mathbf{r} + \mathbf{A}, t) = \mathbf{u}_{\mathbf{K}}(\mathbf{r} + \mathbf{A}) e^{i(\mathbf{K} \cdot (\mathbf{r} + \mathbf{A}) - \omega t)} = \mathbf{u}(\mathbf{r}, t) e^{i\mathbf{K} \cdot \mathbf{A}} \quad (5.11)$$

which is the so-called *periodic boundary condition*. Coupling the periodic boundary condition and the governing equation, the dispersion problem of an infinite periodic system will be transferred into an eigenvalue problem of a finite domain. Then, the dispersion equation can be obtained

$$(\Omega(\mathbf{K}) - \mathbf{M}) \cdot \mathbf{U} = \mathbf{0} \quad (5.12)$$

where  $\mathbf{M}$  is the mass matrix,  $\Omega(\mathbf{K})$  is the stiffness matrix of the considered system. Because of the periodic boundary condition, the stiffness matrix will be a function of the wave vector. For a given wave vector  $\mathbf{K}$ , the dispersion equation can be solved (e.g., by finite element numerical models (Jia and Shi 2010)). By varying the wave vector along the boundary of the first irreducible Brillouin zone (see Figure 5.2(c)), dispersion curves of the considered periodic system can be established. Figure 5.2(d) shows the dispersion curves of the in-plane wave propagating in a periodic structure. As can be seen, in some frequency regions, named frequency bandgap, no wave vector exists, which means waves in these frequency regions cannot propagate in the periodic structure. In other frequency regions, named passing bands, at least one wave vector exists, which means waves in these frequency regions can propagate in the periodic structure. This special filtering effect or frequency bandgap has many potential applications, drawing the attention of many researchers.



**Figure 5.2:** Illustration of the (a) two dimensional periodic structure, (b) the unit cell, (c) the first irreducible Brillouin zone, (d) frequency bandgap.

### 5.6.3 Topology optimization of unit cell of periodic structures

To design periodic structures with desirable frequency bandgap characteristics, topology optimization can be used. For topology optimization of the unit cell of periodic structures, typically the unit cell is discretized into large number of pixels/elements (Yi and Youn 2016; Xie et al. 2017). A pixel matrix is then defined to describe the topological distribution of materials within the unit cell. The topology optimization of unit cell can be defined as

$$\Sigma^* = \operatorname{argmax} f(\Sigma) \quad s.t. : C(\Sigma) > 0 \quad (5.13)$$

where  $\Sigma$  denotes the topological distribution of materials within the unit cell,  $f(\Sigma)$  is the objective function, and  $C(\Sigma) > 0$  is the constraint on the topological distribution  $\Sigma$  (e.g., symmetry, filling fraction of the inclusion). For unit cell with two types of materials, 0 and 1 can be used to represent whether a pixel is filled with one material or the other. The pixel-wise binary design variables are optimized to establish optimal distributions of two materials within the unit cell. As to  $f(\Sigma)$ , two typical objective functions are the absolute bandgap width and the relative bandgap width (Yi and Youn 2016; Dong et al. 2014). Absolute bandgap width is defined as the difference between two adjacent eigenfrequencies  $f(\Sigma) = \omega_{j+1}(\Sigma, \mathbf{K}) - \omega_j(\Sigma, \mathbf{K})$ . The relative bandgap width is defined as

$$f(\Sigma) = \frac{\min_{\mathbf{K}} : \omega_{j+1}(\Sigma, \mathbf{K}) - \max_{\mathbf{K}} : \omega_j(\Sigma, \mathbf{K})}{(\min_{\mathbf{K}} : \omega_{j+1}(\Sigma, \mathbf{K}) + \max_{\mathbf{K}} : \omega_j(\Sigma, \mathbf{K}))/2} \quad (5.14)$$

where  $\min_{\mathbf{K}} : \omega_j(\Sigma, \mathbf{K})$  and  $\max_{\mathbf{K}} : \omega_j(\Sigma, \mathbf{K})$  denote the minimum and maximum of the  $j^{th}$  eigenfrequency  $\omega_j$  over the entire set of discrete wave vectors along the boundaries of the irreducible

Brillouin zone for a given design  $\Sigma$  of the unit cell (Yi and Youn 2016; Dong et al. 2014). Besides the above two, which only consider one particular bandgap, the bandgaps in a certain frequency range (e.g., in  $[\omega_{lb}, \omega_{ub}]$ ) can be also considered. In this case, the sum of the bandgap widths can be used

$$f(\Sigma) = \sum_{j=j_l}^{j_u} (\omega_{j+1}(\Sigma, \mathbf{K}) - \omega_j(\Sigma, \mathbf{K})) \quad (5.15)$$

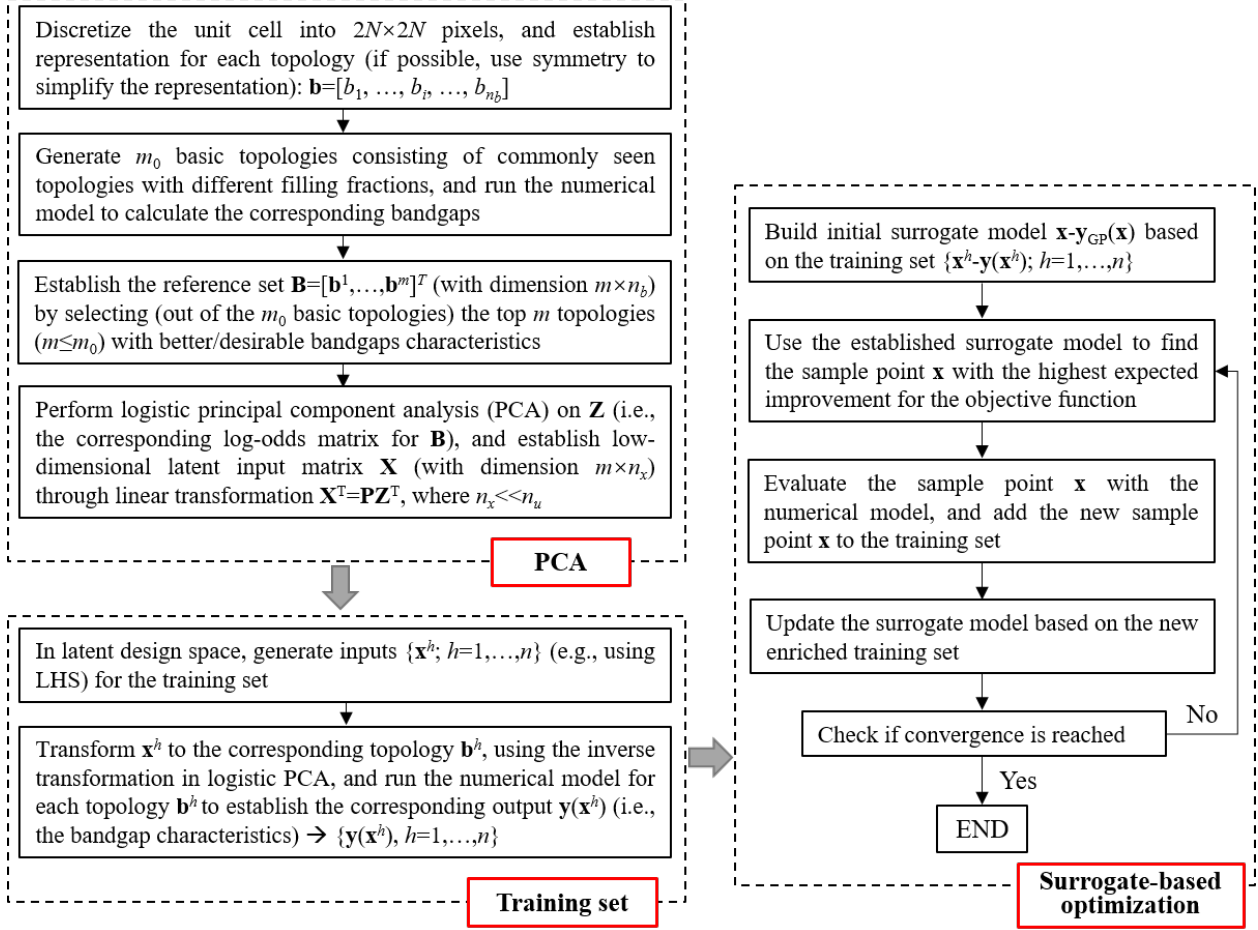
where  $\omega_{j_l} \geq \omega_{lb}$  and  $\omega_{j_u+1} \leq \omega_{ub}$ .

Considering the significant computational effort in calculation of the frequency band structures and solving the high-dimensional optimization problem, typically, symmetry is assumed to reduce the design domain, while the asymmetric unit cell has the largest design space. On the other hand, it has also been shown that, periodic structures with asymmetrical unit cells could provide larger relative bandgap width than those with symmetrical unit cells and changes in symmetry have large impact on the optimized structures (Yi and Youn 2016; Dong et al. 2015). Regardless of symmetry, topology optimization of periodic structures in general is a computationally challenging problem, especially when large number of pixels are used (i.e., corresponding to high-dimensional optimization) and the numerical method for calculating the bandgap for each topology is expensive.

#### 5.6.4 DRESTO: overview

To address the above computational challenges and facilitate efficient topology optimization of periodic structures, we applied the proposed DRESTO algorithm. The flowchart for the overall DRESTO approach is illustrated in Figure 5.3.

Suppose the unit cell can be discretized as  $2N \times 2N$  elements/pixels with each pixel having its corresponding material. For each topology  $\Sigma$ , it can be represented by a vector  $\mathbf{b} =$



**Figure 5.3:** Flowchart for the proposed DRESTO for topology optimization of periodic structures.

$[b_1, \dots, b_i, \dots, b_{n_b}]$  where  $n_b = 4N^2$  and  $b_i$  is the design variable for the  $i^{th}$  pixel. When symmetry is considered, the topology in the unit cell can be characterized by one-eighth of the unit cell. Then the material distribution  $\Sigma$  can be represented by  $\mathbf{b}$  with smaller  $n_b$ , i.e.,  $n_b = (N + 1)N/2$ . Consider the case of bi-component periodic structure,  $\Sigma$  can be represented by a binary matrix with 0 and 1 as element values. The corresponding design variables vector  $\mathbf{b}$  will be a binary vector where  $b_i$  takes values of either 0 or 1. Typically,  $n_b$  is a large number and  $\mathbf{b}$  corresponds to high-dimensional design variables.

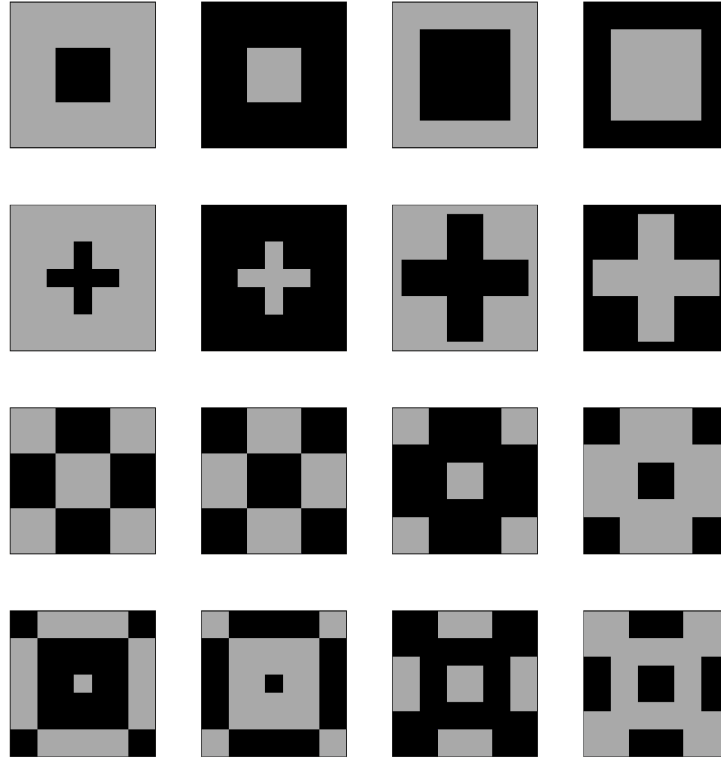
To reduce the dimensionality of the design problem, we establish a low-dimensional representation of  $\mathbf{b}$ . For this purpose, we first collect a set of  $m$  reference topologies (i.e., forming the

reference set), including some of the commonly seen topologies, represented by  $\mathbf{B} = [\mathbf{b}^1, \dots, \mathbf{b}^m]^T$ . As to the selection of the reference set, we can first generate  $m_0$  basic topologies, which will consist of commonly seen topologies (e.g., the ones shown in Figure 5.4) with different filling fractions as well as corresponding topologies with the materials switched (the latter is used to further enrich the set of basic topologies). Then we can pick the top  $m$  topologies based on the objective function values. These  $m$  topologies will form the reference set. Then as mentioned above logistic PCA (shown in Eq. (5.2)) is used to explore the correlations in these topologies to find a lower-dimensional representation of these topologies. In the end, we can establish a transformation between high-dimensional binary design variables  $\mathbf{b}$  and the low-dimensional continuous latent design variables. Later, we will use this transformation to convert from latent space to original space. Figure 5.5 shows the latent space representation of different topologies by three-dimensional latent inputs (i.e., with choice of  $n_x = 3$ ) and four example topologies corresponding to the four points in the latent space.

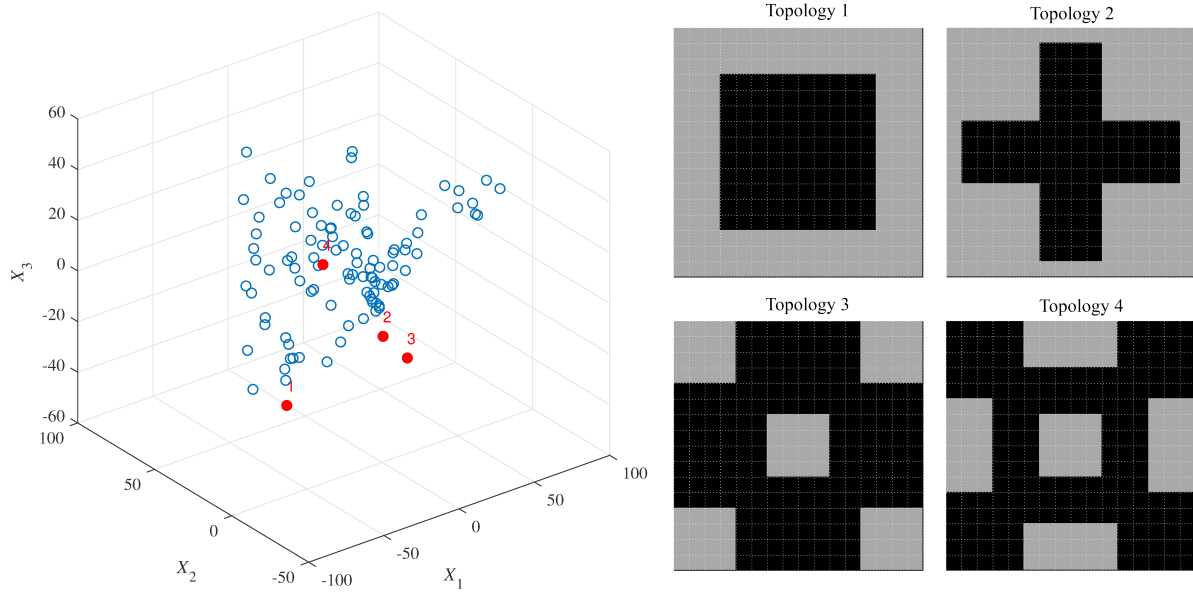
To reduce the computational effort in calculation of bandgap for each design (e.g., through finite element method or other numerical methods), a Gaussian process surrogate model is established with respect to the low-dimensional latent design variables. This surrogate model is then used in place of the original numerical model to facilitate efficient topology optimization. Once the optimal is established in the latent space, the original topology can be easily established through the PCA transformation.

#### 5.6.5 DRESTO: implementation details

In this example, the two materials in the periodic structure correspond to Pb and epoxy. For Pb, it has elastic modulus of  $E_{pb} = 2.433 \times 10^{10} Pa$ , Poisson's ratio of  $\nu_{pb} = 0.449$ , and density



**Figure 5.4:** Several example topologies in the set of basic topologies.



**Figure 5.5:** Illustration of latent space representation of different topologies  $\Sigma$  by three-dimensional latent inputs  $\mathbf{x}$  (i.e.,  $n_x = 3$ ), and four example topologies corresponding to the four points in the latent space (red colored ones).

of  $\rho_{pb} = 11350 \text{ kg/m}^3$ ; for epoxy,  $E_{ep} = 4.518 \times 10^9 \text{ Pa}$ ,  $\nu_{ep} = 0.399$ , and  $\rho_{ep} = 1200 \text{ kg/m}^3$ . Frequency bandgaps are investigated for the in-plane modes. Since it was found in (Dong et al. 2014) that the first and second bands do not exist (because the first two in-plane eigenmodes are both rigid ones at the  $\Gamma$  point) and it is impossible to open the first six bandgaps except the third one, therefore, for comparison purpose, we consider the same two cases as in (Dong et al. 2014), i.e., optimize the topology for the third and seventh bandgap with the relative frequency bandgap width as the objective function. The optimal solutions are validated and compared with the results in (Dong et al. 2014).

For the unit cell,  $N = 30$  is used, which leads to a total of  $4N^2 = 3600$  pixels. Pixels with  $b_i = 1$  correspond to the material Pb, and those with  $b_i = 0$  correspond to the material epoxy. Considering symmetry in the topology, each topology is represented by a design variables vector  $\mathbf{b}$  with dimension  $n_b = 465$ , still corresponding to high-dimensional design variables.

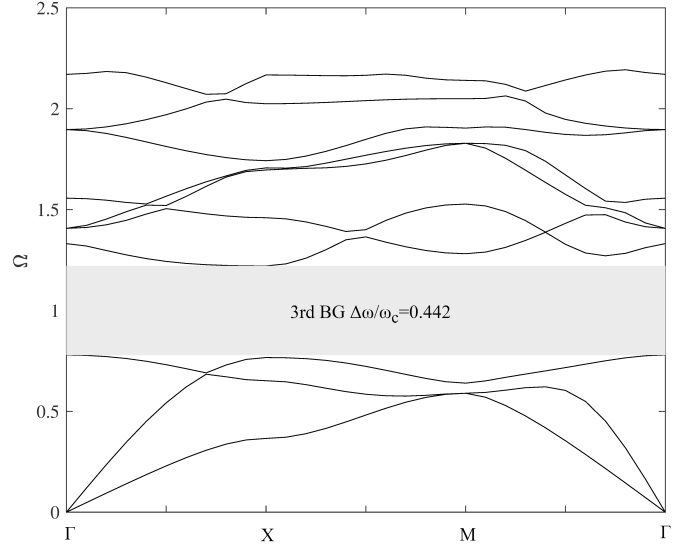
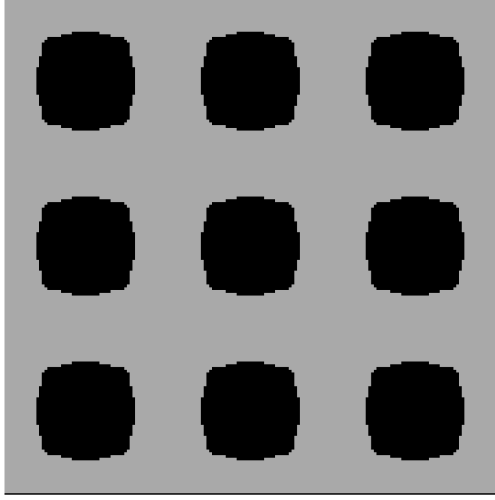
To establish the reference set, a set of basic topologies with different filling fractions are first generated. Some examples of the basic topologies are illustrated in Figure 5.4. For the current example,  $m_0 = 1500$  basic topologies are used. Based on the objective function values, in the end, the top  $m = 200$  topologies are selected to form the reference set. We then apply logistic PCA to this reference set to establish the transformation in Eq. (5.5). For the number of latent design variables  $n_x$ , it is chosen so that the misclassification rate (i.e., a measure of the reconstruction accuracy) is below 1%. This leads to  $n_x = 20$  when the objective function corresponds to the third bandgap and  $n_x = 12$  when the objective function corresponds to the seventh bandgap. It was found that further increasing  $n_x$  will not lead to significant improvement in the reconstruction accuracy. As can be seen, the dimension of the design variables is significantly reduced, i.e., from 465 to 20 and 12.

To build the surrogate model, the inputs  $\{\mathbf{x}^h; h = 1, \dots, n\}$  (i.e., samples points) for the training set is established in the latent design space using LHS. The initial number of samples is selected as  $n = 600$ . Another option is to select this number adaptively so that a certain accuracy level is achieved. Since here we are interested in using surrogate model in the context of optimization, EGO will adaptively select more samples to add in the training set so that the probability of finding a better solution is maximized. In other words, the adaptive selection mainly happens in the EGO. However, note that a initially large value for  $n$  typically helps improve the accuracy of the surrogate model, which might save computational efforts in the EGO optimization process (e.g., requiring fewer iterations to converge). Then the corresponding topologies are established using the transformation in Eq. (5.5), leading to a training set with 600 topologies. For each topology in the training set, we use finite element numerical model to calculate the corresponding frequency bandgap (Jia and Shi 2010). For EGO, the maximum number of iterations  $n_{iter,max}$  is set at 150.

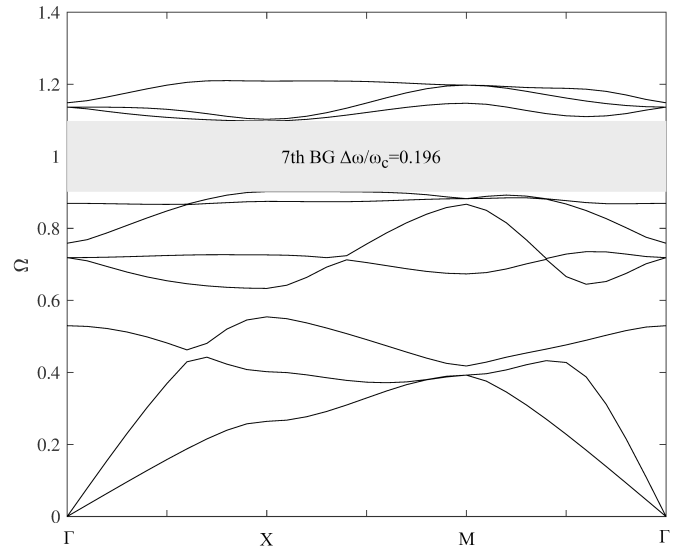
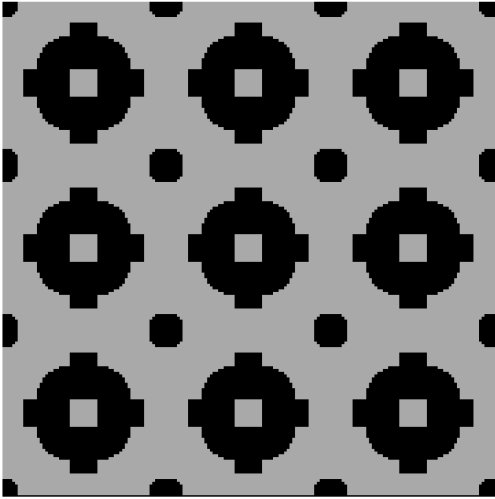
#### 5.6.6 Optimal topologies

Figure 5.6 illustrates the optimal topology for the third bandgap found by DRESTO and the corresponding band structures (with the bandgap labeled). For comparison purpose, Figure 5.8 shows the optimal topologies for the third and seventh bandgap obtained in (Dong et al. 2014). The maximum relative bandgap width found by DRESTO is 0.442, which is close to the value (i.e., 0.455) in (Dong et al. 2014). Here for the optimal topology, the Pb inclusion (embedded in Epoxy matrix) takes the form of a curved square with filling fraction of 0.312. This is slightly different from the one established in (Dong et al. 2014) where the inclusion takes the form of a square with round corners (see Figure 5.8) and has a filling fraction of 0.35. This can be explained by the fact

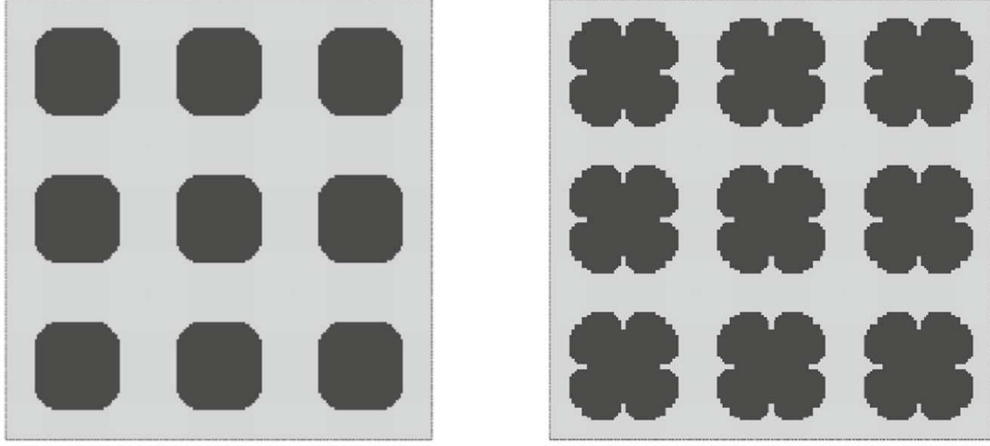
that multiple local minimum may exist and different topologies may give similar or close relative bandgap width values.



**Figure 5.6:** Optimal topology and corresponding band structures for the third bandgap of the in-plane mode.



**Figure 5.7:** Optimal topology and corresponding band structures for the seventh bandgap of the in-plane mode.



**Figure 5.8:** Optimal topologies for the third (left) and seventh (right) bandgap of the in-plane mode in Dong et al. (2014).

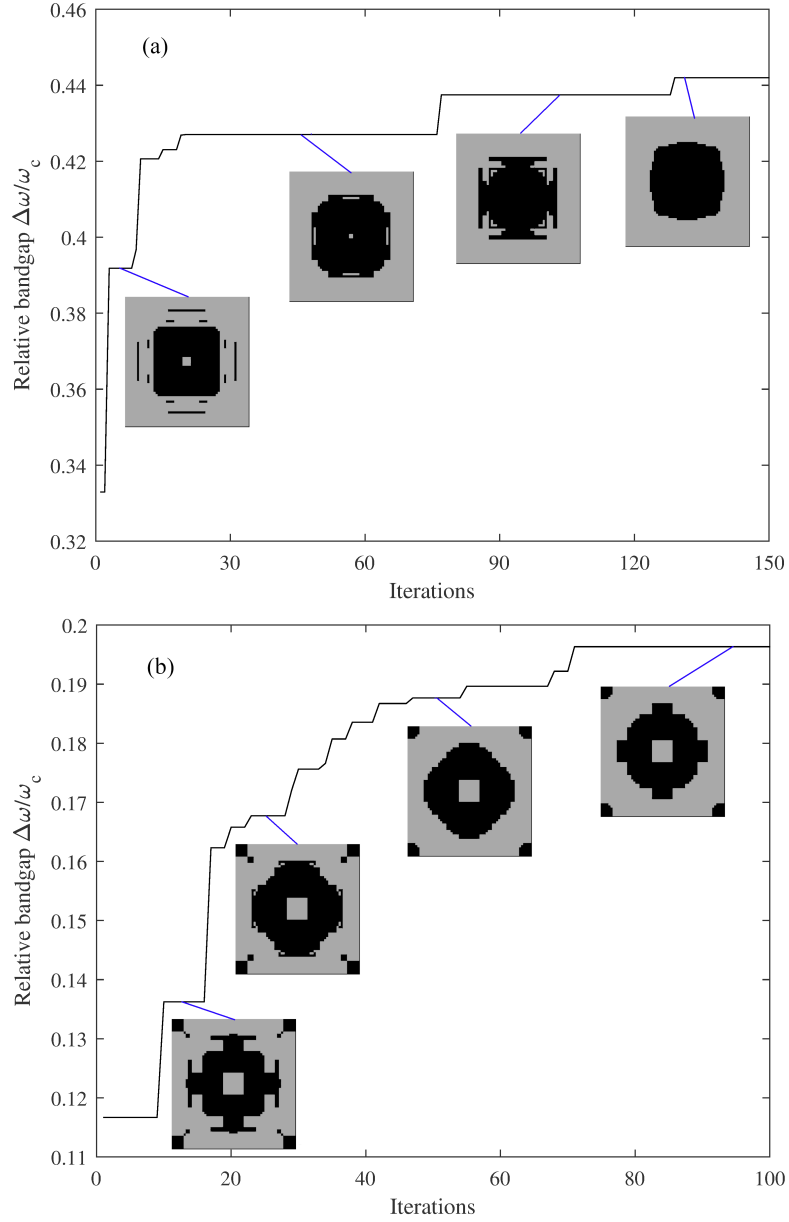
For the seventh bandgap of the in-plane mode, the optimal topology and corresponding band structures are presented in Figure 5.7. In this case, DRESTO finds a topology that gives a slightly larger (16% larger) maximum relative bandgap value (i.e., 0.196) than the optimal value (i.e., 0.169) obtained in (Dong et al. 2014). The optimal topology for the seventh bandgap has irregular shape. When looking at the unit cell, in addition to the main inclusion (which looks like the union of a cross and a round shape but with hollow square in the center), there are small one quarter of round-cornered squares at the four corners. When looking at the plot in Figure 5.7, overall, the periodic structure has some main inclusions and also some small round-cornered squares. The filling fraction is around 0.357. In comparison, for the optimal topology identified in (Dong et al. 2014), the inclusion looks like a pedal (see Figure 5.8), and the filling fraction is around 0.49. Overall, the optimal solutions obtained by DRESTO algorithm are consistent with those in (Dong et al. 2014) in terms of the values for the maximum relative bandgap width. For the seventh bandgap, DRESTO gives better objective function values and identifies a different optimal topology with much smaller filling fraction (i.e., 0.357 instead of 0.49). The above results verify the effectiveness of the proposed algorithm.

It should be noted that the optimal topologies might bring challenges to manufacturing in practice. To tackle the challenges, many research works on developing mathematical formulations of manufacturing constraints have been proposed, such as (Zhou et al. 2002; Xia et al. 2010; Gersborg and Andreasen 2011; Li et al. 2015). In addition, some advanced manufacturing techniques, e.g., additive manufacturing, can also be applied to address the difficulties in fabrication of complex designs produced by topology optimization. Also, before being applied in practice, the actual performance of the optimal topologies should be validated experimentally. Note that the focus of this example is on developing efficient topology optimization algorithms, and the manufacturing and experimental validation of the optimal topologies are out of the scope of this chapter, but may be areas for future research.

#### 5.6.7 *Computational efficiency*

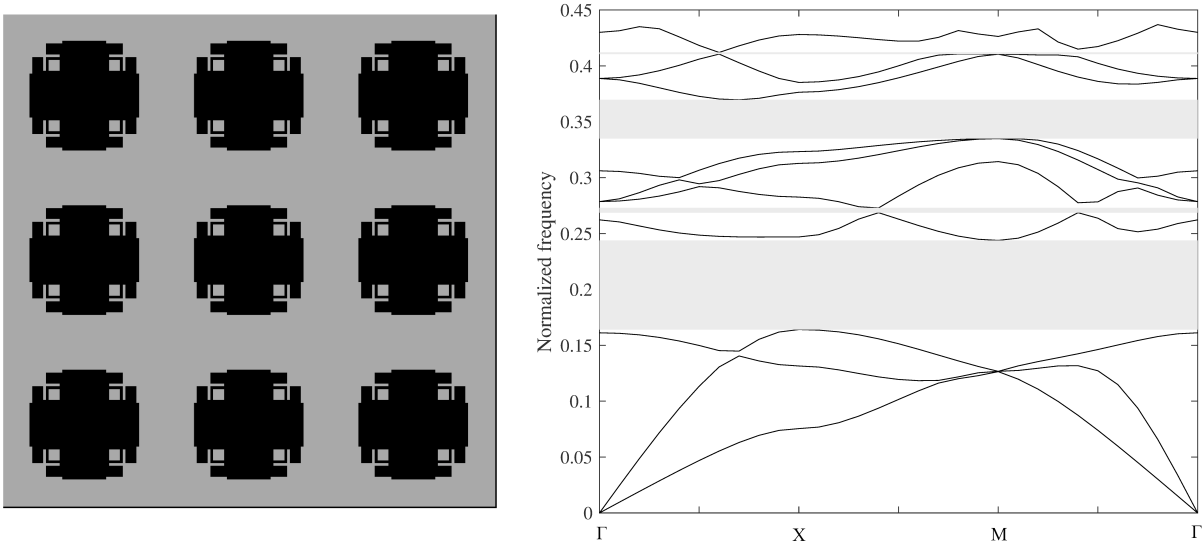
Figure 5.9 shows the variation of the objective function value over the iterations for both the third and seventh bandgap as well as some of the corresponding topologies identified at some of the iterations. For the third bandgap, the final topology has been identified by around 130 iterations, while for the seventh bandgap, the final topology has been identified by around 90 iterations. One interesting observation for the third bandgap is that the topology evolves from inclusions that are more spread to inclusions that are converging towards one single cluster (i.e., with fewer isolated or segments of pixels) and eventually to a curved square. For the seventh bandgap, the several topologies identified during the iterations all have small (round-cornered) squares and it seems that the small (round-cornered) squares help further widen the bandgap. Compared to the initial value of the objective function, i.e., the best value from the 600 topologies in the training set which is 0.333 and 0.116 for the third and seventh bandgap respectively, the EGO after less than 150 iterations

gives objective function values of 0.442 and 0.196, corresponding to improvements of 33% and 69%, respectively. This shows the effectiveness of the EGO with adaptive selection of new sample points.

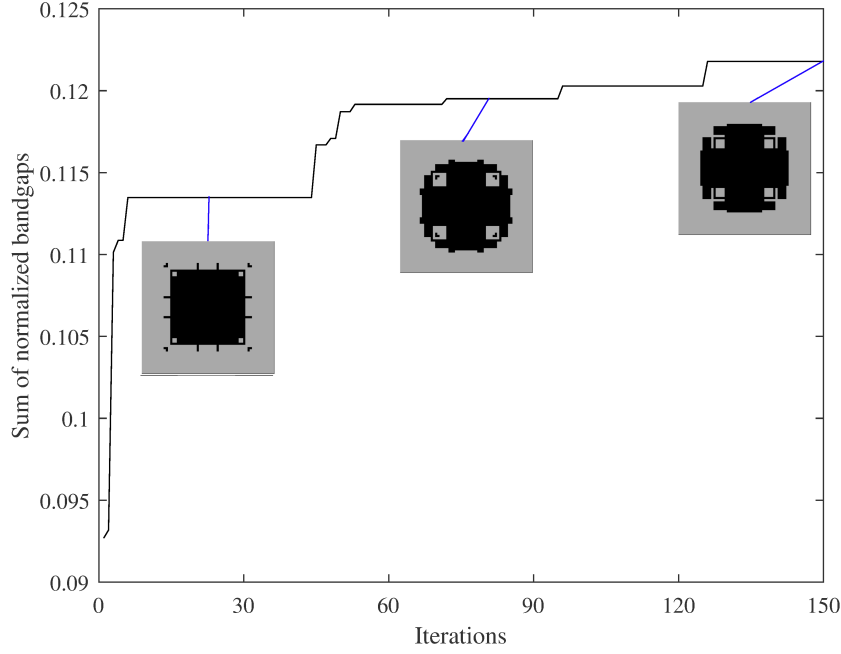


**Figure 5.9:** Topologies identified by DRESTO over the iterations for (a) the third bandgap and (b) the seventh bandgap.

In terms of computational effort, to establish the optimal topology, the total number of evaluations of the numerical model corresponds to  $n_{total} = m_0 + n + n_{iter}$ . For the third bandgap, this corresponds to  $n_{total} = 1500 + 600 + 150 = 2250$ . For the seventh bandgap, to establish the PCA transformation, the same set of basic topologies, (i.e.,  $m_0 = 1500$ ) and the corresponding eigenfrequency information can be directly used (i.e., to select the top  $m$  topologies). Only  $n + n_{iter} = 600 + 100 = 700$  additional evaluations of the numerical model is needed to establish the corresponding optimal topology. In the literature (Dong et al. 2014), for the same or similar problem, direct search using genetic algorithm typically requires around 1000-2000 generations to converge to the optimal with population size of at least around 20 for each generation, resulting to around 20,000-40,000 model evaluations. In comparison, overall the proposed algorithm only requires around 2250 model evaluations, corresponding to much better efficiency.



**Figure 5.10:** Optimal topology and corresponding band structures for the in-plane mode with the sum of the bandgap widths as objective function.



**Figure 5.11:** Topologies identified by DRESTO over the iterations for the sum of normalized bandgaps.

More importantly, as mentioned earlier, another advantage of the proposed algorithm is that the same set of  $m_0 = 1500$  basic topologies and the corresponding eigenfrequency information can be directly used for different definition of objective functions. To establish the optimal for the new objective function, only  $n + n_{iter}$  model evaluations are needed. For example, using the information in the same set of basic topologies, the algorithm is used to find the optimal for another objective function, i.e., the sum of the bandgap width within the first ten eigenfrequencies (see Eq. (5.15)). The optimal topology and the corresponding band structures are shown in Figure 5.10 while Figure 5.11 shows the variation of the objective function value over the iterations and topologies identified at some of the iterations. Note that the  $y$ -axis values in these two figures correspond to the normalized frequency and bandgap where the normalization  $(a\omega)/(2\pi c_t)$  is used where  $c_t$  is the shear velocity of the matrix material. All these results were established using only  $n + n_{iter} = 600 + 150 = 750$  additional model evaluations.

In addition, when more fine representation of the topology is needed, i.e., increase the value of  $N$ , the number of design variables in the original design space would increase quadratically with respect to  $N$ , making the resulting optimization more challenging with much larger search space. On the other hand, for the proposed algorithm, logistic PCA is used to exploit the patterns or correlations in the topologies, the number of latent design variables (i.e.,  $n_x$ ) will typically not increase much, which means that the search space (i.e., the latent design space) and the resulting computational efforts for the corresponding optimization problem (in the latent design space) will not increase much. This means the proposed algorithm is especially useful for topology optimization when very fine representation of the topology is needed.

## 5.7 Conclusions

This chapter proposed an efficient dimension reduction and surrogate based optimization approach for problems with expensive high-fidelity model and high-dimensional binary inputs. Using information from a set of reference inputs with better objective function values, dimension reduction technique (i.e., logistic PCA) was used to establish a low-dimensional representation of different inputs whose representation in the original design space corresponds to high-dimensional design variables. To reduce the computational effort in calculation of objective function values for each design, Gaussian process model was built with respect to the low-dimensional latent design variables and used within efficient global optimization to adaptively identify the optimal design. The low-dimensionality of the latent design variables facilitated much easier construction of the surrogate model and also the optimization. Once the optimal was established in the latent space, the original optimal design was easily established through the PCA transformation.

The illustrative example on the topology optimization of a 2D periodic structure verified the effectiveness and efficiency of the proposed algorithm. The proposed algorithm was able to identify the optimal topology within a small number of iterations. The proposed algorithm is especially useful for topology optimization when very fine representation of the topology is used since for the proposed algorithm the dimension of the latent space will not increase much while direct optimization in the original space would have much higher dimension for the design variables. The proposed algorithm can be easily applied to topology optimization for the out-of-plane modes or the joint modes. Future research will investigate how to address periodic structures with multiple components/materials and how to incorporate constraints in the proposed algorithm, e.g., for specific filling fraction, where the selection of the set of basic topologies would need to be properly modified. Lastly, as mentioned earlier, parametric optimization can be used as a successive refinement of the topology optimization, therefore topology optimization integrated with parametric optimization is also a future research topic of great interest.

## CHAPTER 6:

### SUMMARY OF DISSERTATION AND FUTURE STUDIES

#### 6.1 Conclusions

This dissertation aimed at efficiently constructing surrogate models to facilitate efficient analysis and design of complex engineering systems. The key novel contribution of this dissertation is on proposing different novel strategies by enriching the training data and enhancing model assumption to reduce the training cost and improve the prediction and generalization performance of Gaussian process models. More specifically, this dissertation (i) developed a physics-constrained Gaussian process model to efficiently incorporate our prior knowledge about physical constraints/characteristics of the input-output relationship by designing specific kernels, (ii) established a more general multi-fidelity Gaussian process model integrating a small number of expensive high-fidelity data and a large number of cheap low-fidelity data by developing a more general model form to enhance the existing multi-fidelity Gaussian process model and by developing a corresponding Bayesian calibration framework, (iii) proposed a multi-fidelity Gaussian process model capable of integrating training data with different level of accuracy (high-fidelity data and low-fidelity data) and completeness (i.e., accurate data and censored data), and (iv) developed an efficient surrogate modeling approach for design optimization problems with high-dimensional binary model inputs (or design variables) by integrating dimension reduction technique and Gaussian process model and by using adaptive sampling scheme. All the above developments are the novel contributions of this dissertation. These novel strategies can effectively reduce the required size of high-fidelity training data and meanwhile effectively boosts the prediction accuracy of the established Gaussian process model.

Chapter 2 introduced the construction of a standard Gaussian process model, including the model formulation, calibration, and prediction. Some related important topics, such as design of experiment and multidimensional outputs, were also discussed. This chapter laid the foundation for the proposed strategies to improve the performance of the Gaussian process model.

Chapter 3 proposed a physics-constrained Gaussian process model to efficiently and accurately predict responses that show invariance, symmetry, and additivity features. Instead of building a standard Gaussian process model and learning the physical constraints (i.e., invariance, symmetry, and additivity features) of the problem through a large training data set which was computationally inefficient or even prohibitive, the proposed physics-constrained Gaussian process model directly encoded these physical constraints/features into the model development. More specifically, the known physical constraints were encoded into the kernel by designing and integrating the invariant kernel and the additive kernel, and in this way a more “informative prior” was provided to the Gaussian process model construction. Once trained, the physics-constrained Gaussian process model was employed to directly and efficiently predict the responses. An application to the prediction of the hydrodynamic characteristics of arrays with different number of wave energy converters (WECs) demonstrates the high accuracy and efficiency of the proposed approach. The results showed that the designed integrated kernel was able to correctly capture the invariance, symmetry, and additivity features of the problem. The proposed physics-constrained Gaussian process model can accurately predict the hydrodynamic characteristics with a relatively small number of training data, especially when the wave frequency was low. More importantly, the proposed physics-constrained Gaussian process model was much less vulnerable to curse of dimensionality compared to standard Gaussian process model, and such good scalability was crucial for analyzing arrays with relatively large number of WECs.

Chapter 4 established a general multi-fidelity Gaussian process model to integrate data with different levels of accuracy. The idea is to leverage the trend information provided by a large number of low-fidelity (i.e., low accuracy) data and the accuracy provided by a small number of expensive high-fidelity data to establish a model that is better than Gaussian process models built using either the high-fidelity data or low-fidelity data alone. In calibration of the multi-fidelity Gaussian process model, censoring in the high-fidelity data, which is common in experimental testings and high-fidelity numerical simulations (e.g., with partially converged results), and the measurement errors were explicitly considered. Posterior distributions of the model parameters were established to explicitly take into account the uncertainties in the model parameters. To address the computational challenges in estimating the likelihood function for censored data, a data augmentation algorithm was proposed, which artificially treats the censored data as unknown parameters in addition to the original model parameters. Then closed form conditional posteriors were derived for these parameters, which were used in the context of Gibbs sampling to facilitate efficient sampling from the posterior distributions for the model parameters. In the end, the posterior samples for the model parameters were used to establish the posterior statistics for the output predictions at new inputs. The proposed multi-fidelity Gaussian process model was then applied to establish predictive model for the deformation capacity of RC columns. The results showed that it is important to explicitly integrate the information from censored data in establishing the multi-fidelity Gaussian process model, rather than only relying on accurate data, especially when censored data account for a significant portion of the high-fidelity data. Overall, the proposed multi-fidelity Gaussian process model is expected to be useful for problems when (i) only limited number of high-fidelity data are available, (ii) censored data account for a significant portion of the high-fidelity data, (iii) low-

fidelity data can be established relatively efficiently, and (iv) uncertainties in the model parameters need to be explicitly considered.

Chapter 5 proposed an efficient dimension reduction and surrogate based optimization approach for problems with high-dimensional binary inputs. Using information from a set of reference inputs with better objective function values, dimension reduction technique (i.e., logistic PCA) was used to establish a low-dimensional representation of different inputs whose representation in the original design space corresponds to high-dimensional design variables. To reduce the computational effort in calculation of objective function values for each design, Gaussian process model was built with respect to the low-dimensional latent design variables and used within efficient global optimization to adaptively identify the optimal topology. The low-dimensionality of the latent design variables facilitated much easier construction of the Gaussian process model and also optimization. Once the optimal was established in the latent space, the original optimal design was easily established through the PCA transformation. The illustrative example to the topology optimization of a 2D bi-component periodic structure verified the effectiveness and efficiency of the proposed algorithm. The proposed algorithm was able to identify the optimal topology within a small number of iterations. The proposed algorithm is especially useful for topology optimization when very fine representation of the topology is used since for the proposed algorithm the dimension of the latent space will not increase much while direct optimization in the original space would have much higher dimension for the design variables.

## 6.2 Future directions

Although this dissertation proposed several advanced strategies to develop efficient Gaussian process model for design and analysis of engineering systems, some possible improvements and extensions based on the current research still require investigation in the future.

First, the proposed strategies have several limitations and require further improvements.

- For the proposed physics-constrained Gaussian process model, the computational cost of evaluating the physics-constrained kernel (i.e., invariant kernel combined with additive kernel) depends on the complexity of features exhibited by the problem. Current kernel was developed based on the assumption that the symmetry/invariance group only includes discrete and a small/medium number of transformations. However, if the problem involves continuous or a large number of transformations in the symmetry/invariance group (e.g., problems about images which have thousands of pixels), the developed invariant kernel might require a significant computational cost or even become prohibitive to evaluate. Therefore, the invariant kernel needs some improvement to make it applicable to problems with more complex features. On the other hand, for the additive kernel part, we assumed the decomposition of the model responses converge fast and thus can obtain computational cost reduction by truncating the expansion. However, if the decomposition of the many-body system does not converge fast, e.g., some higher orders of interaction contribute more to the responses, this truncation method will not work effectively. One way to address this issue is to optimize the variances corresponding to each order of interaction and identify the important orders, however this may lead to significant computational issues especially when the input dimen-

sionality is high. Therefore, how to establish an efficient additive kernel for many-body problems which do not converge fast is also a topic of interest in the future.

- For the proposed multi-fidelity Gaussian process model calibrated by both accurate and censored data, this dissertation only considered censored responses in the high-fidelity data/responses. Nevertheless, we sometimes also have censored low-fidelity data (e.g., obtained by numerical models which are not fully converged). In such case, the likelihood function related to model parameters in low-fidelity model also needs to be modified to incorporate the censored low-fidelity data. As a result, how to efficiently calibrate the likelihood functions still requires further exploration. In addition, this dissertation only focused on constructing Gaussian process model using data with two levels of fidelity (i.e., high-fidelity and low-fidelity). For many problems, the simulations can be run at more fidelity levels with different costs (e.g., using different grid sizes in finite element model). This suggests another future research direction: establishing Gaussian process models using multiple (i.e., more than two) fidelity data when considering data censoring. The multi-fidelity model could be based on the recursive multi-fidelity Gaussian process model proposed in Le Gratiet (2013b) or the auto-regressive model developed in Kennedy and O'Hagan (2000). But note that the literature only deals with accurate training data, and needs to be extended to consider censored training data.
- The proposed dimension-reduction assisted Gaussian process model used the commonly used PCA technique which is a linear dimension reduction technique. In the future, research can be extended by applying nonlinear dimension reduction techniques in order to obtain a more realistic mapping between the high-dimensional inputs and the low-dimensional latent inputs.

One option of the nonlinear dimension reduction is the *autoencoder* (Hinton and Salakhutdinov 2006; Ng et al. 2011) which is a type of neural network designed for efficiently encoding and decoding features in the data. Compared with linear PCA, autoencoders is powerful in learning the nonlinear correlation between the model inputs, and restoring the inputs with much less information loss. Other nonlinear dimension reduction techniques which have attracted extensive research interests include *variational autoencoder* (Doersch 2016; Kingma and Welling 2019), *diffusion map* (Ferguson et al. 2011; Giannakis 2019), and etc. Additionally, this dissertation developed the dimension-reduction assisted Gaussian process model in the context of optimization. If we are interested in other applications such as sensitivity analysis, the applicability of the proposed model needs to be investigated.

Second, in terms of the applications, we also have some future research interests.

- For the application to the WECs in an array, the dissertation only investigated arrays with up to 10 WECs. Future research work will investigate the scalability of the proposed model to arrays with even larger sizes (e.g., with 20 to 50 WECs). Another future research topic of interest is to use the hydrodynamic characteristics predicted by the proposed physics-constrained Gaussian process model to calculate the total power generation of the WEC array and also optimize the layout of the WEC array.
- For the application to the topology optimization of 2D bi-component periodic structures, the proposed dimension reduction and surrogate based optimization can be easily applied to topology optimization for the out-of-plane vibration modes or the joint modes. Future research will investigate how to address periodic structures with multiple components/materials and how to incorporate constraints in the proposed algorithm, e.g., for spe-

cific filling fraction, where the selection of the set of basic topologies would need to be properly modified.

Finally, this dissertation developed the strategies in the context of Gaussian process models, but the ideas of these methods are general and can be applied to other surrogate models. Therefore, one potential research direction is to investigate the development of the strategies for other commonly used surrogate models, such as polynomial chaos expansions and support vector machines.

## BIBLIOGRAPHY

- Abrahamsen, P. and Benth, F. E. (2001). Kriging with inequality constraints. *Mathematical Geology*, 33(6):719–744.
- Albawi, S., Mohammed, T. A., and Al-Zawi, S. (2017). Understanding of a convolutional neural network. In *2017 International Conference on Engineering and Technology (ICET)*, pages 1–6. Ieee.
- Albert, C. G. (2019). Gaussian Processes for Data Fulfilling Linear Differential Equations. *Proceedings*, 33(1):5.
- Aversano, G., Parra-Alvarez, J. C., Isaac, B. J., Smith, S. T., Coussement, A., Gicquel, O., and Parente, A. (2019). Pca and kriging for the efficient exploration of consistency regions in uncertainty quantification. *Proceedings of the Combustion Institute*, 37(4):4461–4469.
- Bacigalupo, A., Gnecco, G., Lepidi, M., and Gambarotta, L. (2017). Optimal design of low-frequency band gaps in anti-tetrachiral lattice meta-materials. *Composites Part B: Engineering*, 115:341–359.
- Bacigalupo, A., Lepidi, M., Gnecco, G., Vadalà, F., and Gambarotta, L. (2019). Optimal design of the band structure for beam lattice metamaterials. *Frontiers in Materials*, 6:1–14.
- Bartók, A. P., Kondor, R., and Csányi, G. (2013). On representing chemical environments. *Physical Review B*, 87(18):184115.
- Beers, W. C. M. V. and Kleijnen, J. P. C. (2003). Kriging for interpolation in random simulation. *Journal of the Operational Research Society*, 54(3):255–262.

- Beymer, D. and Poggio, T. (1995). Face recognition from one example view. In *Proceedings of IEEE International Conference on Computer Vision*, pages 500–507. IEEE.
- Bilal, O. R., Ballagi, D., and Daraio, C. (2018). Architected lattices for simultaneous broadband attenuation of airborne sound and mechanical vibrations in all directions. *Physical Review Applied*, 10(5):054060.
- Blatman, G. and Sudret, B. (2010). Efficient computation of global sensitivity indices using sparse polynomial chaos expansions. *Reliability Engineering & System Safety*, 95(11):1216–1229.
- Borgarino, B., Babarit, A., and Ferrant, P. (2012). Impact of wave interactions effects on energy absorption in large arrays of wave energy converters. *Ocean Engineering*, 41:79–88.
- Bouhlef, M. A., Bartoli, N., Otsmane, A., and Morlier, J. (2016a). An improved approach for estimating the hyperparameters of the kriging model for high-dimensional problems through the partial least squares method. *Mathematical Problems in Engineering*, 2016.
- Bouhlef, M. A., Bartoli, N., Otsmane, A., and Morlier, J. (2016b). Improving kriging surrogates of high-dimensional design models by partial least squares dimension reduction. *Structural and Multidisciplinary Optimization*, 53(5):935–952.
- Bronstein, M. M., Bruna, J., LeCun, Y., Szlam, A., and Vandergheynst, P. (2017). Geometric deep learning: going beyond euclidean data. *IEEE Signal Processing Magazine*, 34(4):18–42.
- Bucher, C. G. and Bourgund, U. (1990). A fast and efficient response surface approach for structural reliability problems. *Structural safety*, 7(1):57–66.

- Budal, K. et al. (1977). Theory for absorption of wave power by a system of interacting bodies. *Journal of Ship Research*, 21(04):248–254.
- Cai, X., Qiu, H., Gao, L., Wei, L., and Shao, X. (2017). Adaptive radial-basis-function-based multifidelity metamodeling for expensive black-box problems. *AIAA journal*, 55(7):2424–2436.
- Cheng, M., Jiang, P., Hu, J., Shu, L., and Zhou, Q. (2021). A multi-fidelity surrogate modeling method based on variance-weighted sum for the fusion of multiple non-hierarchical low-fidelity data. *Structural and Multidisciplinary Optimization*, pages 1–22.
- Cheng, Z., Lin, W., and Shi, Z. (2018a). Wave dispersion analysis of multi-story frame building structures using the periodic structure theory. *Soil Dynamics and Earthquake Engineering*, 106:215–230.
- Cheng, Z., Shi, Z., and Mo, Y. (2018b). Complex dispersion relations and evanescent waves in periodic beams via the extended differential quadrature method. *Composite Structures*, 187:122 – 136.
- Chib, S. (1992). Bayes inference in the tobit censored regression model. *Journal of Econometrics*, 51(1-2):79–99.
- Choi, S.-K., Grandhi, R. V., Canfield, R. A., and Pettit, C. L. (2004). Polynomial chaos expansion with latin hypercube sampling for estimating response variability. *AIAA journal*, 42(6):1191–1198.
- Cohen, T., Weiler, M., Kicanaoglu, B., and Welling, M. (2019). Gauge equivariant convolutional networks and the icosahedral cnn. In *International Conference on Machine Learning*, pages 1321–1330. PMLR.

- Constantine, P. G. (2015). *Active subspaces: Emerging ideas for dimension reduction in parameter studies*. SIAM.
- Constantine, P. G., Dow, E., and Wang, Q. (2014). Active Subspace Methods in Theory and Practice: Applications to Kriging Surfaces. *SIAM Journal on Scientific Computing*, 36(6):A3030–A3031.
- Crestaux, T., Le Maître, O., and Martinez, J.-M. (2009). Polynomial chaos expansion for sensitivity analysis. *Reliability Engineering & System Safety*, 94(7):1161–1172.
- Dahl, J., Jensen, J. S., and Sigmund, O. (2008). Topology optimization for transient wave propagation problems in one dimension : Design of filters and pulse modulators. *Structural and Multidisciplinary Optimization*, 36(6):585–595.
- Dai, C., Xue, L., Zhang, D., and Guadagnini, A. (2016). Data-worth analysis through probabilistic collocation-based ensemble kalman filter. *Journal of Hydrology*, 540:488–503.
- Dao, T., Gu, A., Ratner, A., Smith, V., De Sa, C., and Ré, C. (2019). A kernel theory of modern data augmentation. In *International Conference on Machine Learning*, pages 1528–1537. PMLR.
- De Oliveira, V. (2005). Bayesian inference and prediction of gaussian random fields based on censored data. *Journal of Computational and Graphical Statistics*, 14(1):95–115.
- DeCoste, D. and Schölkopf, B. (2002). Training invariant support vector machines. *Machine learning*, 46(1):161–190.
- DeMaris, A. (2004). *Regression with social data: Modeling continuous and limited response variables*, volume 417. John Wiley & Sons.

- Deng, J., Gu, D., Li, X., and Yue, Z. Q. (2005). Structural reliability analysis for implicit performance functions using artificial neural network. *Structural Safety*, 27(1):25–48.
- Doersch, C. (2016). Tutorial on variational autoencoders. *arXiv preprint arXiv:1606.05908*.
- Dong, H.-W., Su, X.-X., Wang, Y.-S., and Zhang, C. (2014). Topological optimization of two-dimensional phononic crystals based on the finite element method and genetic algorithm. *Structural and Multidisciplinary Optimization*, 50:593–604.
- Dong, H. W., Wang, Y. S., Wang, Y. F., and Zhang, C. (2015). Reducing symmetry in topology optimization of two-dimensional porous phononic crystals. *AIP Advances*, 5(11):117149–1–16.
- Drew, B., Plummer, A. R., and Sahinkaya, M. N. (2009). A review of wave energy converter technology. *Proceedings of the Institution of Mechanical Engineers, Part A: Journal of Power and Energy*, 223(8):887–902.
- Dubrulle, O. and Kostov, C. (1986). An interpolation method taking into account inequality constraints: I. methodology. *Mathematical Geology*, 18(1):33–51.
- Durantín, C., Rouxel, J., Désidéri, J.-A., and Glière, A. (2017). Multifidelity surrogate modeling based on radial basis functions. *Structural and Multidisciplinary Optimization*, 56(5):1061–1075.
- Durrande, N., Ginsbourger, D., and Roustant, O. (2012). Additive covariance kernels for high-dimensional gaussian process modeling. In *Annales de la Faculté des sciences de Toulouse: Mathématiques*, volume 21, pages 481–499.

- Durrande, N. and Le Riche, R. (2017). Introduction to gaussian process surrogate models. Technical report, École thématique, France.
- Duvenaud, D. (2014). *Automatic model construction with Gaussian processes*. PhD thesis, University of Cambridge.
- Falnes, J. (1980). Radiation impedance matrix and optimum power absorption for interacting oscillators in surface waves. *Applied ocean research*, 2(2):75–80.
- Ferguson, A. L., Panagiotopoulos, A. Z., Kevrekidis, I. G., and Debenedetti, P. G. (2011). Nonlinear dimensionality reduction in molecular simulation: The diffusion map approach. *Chemical Physics Letters*, 509(1-3):1–11.
- Fernández-Godino, M. G., Park, C., Kim, N.-H., and Haftka, R. T. (2016). Review of multi-fidelity models. *arXiv preprint arXiv:1609.07196*.
- Forrester, A., Keane, A., et al. (2008). *Engineering design via surrogate modelling: a practical guide*. John Wiley & Sons.
- Forrester, A. I. and Keane, A. J. (2009). Recent advances in surrogate-based optimization. *Progress in Aerospace Sciences*, 45(1-3):50–79.
- Forrester, A. I., Sobester, A., and Keane, A. J. (2007). Multi-fidelity optimization via surrogate modelling. *Proceedings of the royal society a: mathematical, physical and engineering sciences*, 463(2088):3251–3269.
- Fox, J. and Weisberg, S. (2002). Cox proportional-hazards regression for survival data. *An R and S-PLUS companion to applied regression*, 2002.

- Frankel, A. L., Jones, R. E., and Swiler, L. P. (2020). Tensor basis gaussian process models of hyperelastic materials. *Journal of Machine Learning for Modeling and Computing*, 1(1).
- Gammelli, D., Peled, I., Rodrigues, F., Pacino, D., Kurtaran, H. A., and Pereira, F. C. (2020). Estimating latent demand of shared mobility through censored gaussian processes. *Transportation Research Part C: Emerging Technologies*, 120:102775.
- Gano, S., Sanders, B., and Renaud, J. (2004). Variable fidelity optimization using a kriging based scaling function. In *10th AIAA/ISSMO Multidisciplinary analysis and optimization conference*, page 4460.
- Gardoni, P., Der Kiureghian, A., and Mosalam, K. M. (2002). Probabilistic capacity models and fragility estimates for reinforced concrete columns based on experimental observations. *Journal of Engineering Mechanics*, 128(10):1024–1038.
- Gardoni, P., Mosalam, K. M., and der Kiureghian, A. (2003). Probabilistic seismic demand models and fragility estimates for rc bridges. *Journal of Earthquake Engineering*, 7(spec01):79–106.
- Gazonas, G. A., Weile, D. S., Wildman, R., and Mohan, A. (2006). Genetic algorithm optimization of phononic bandgap structures. *International Journal of Solids and Structures*, 43(18-19):5851–5866.
- Gelman, A., Stern, H. S., Carlin, J. B., Dunson, D. B., Vehtari, A., and Rubin, D. B. (2013). *Bayesian data analysis*. Chapman and Hall/CRC.
- Gensheimer, M. F. and Narasimhan, B. (2019). A scalable discrete-time survival model for neural networks. *PeerJ*, 7:e6257.

- Gersborg, A. R. and Andreasen, C. S. (2011). An explicit parameterization for casting constraints in gradient driven topology optimization. *Structural and Multidisciplinary Optimization*, 44(6):875–881.
- Ghanem, R. G. and Spanos, P. D. (2003). *Stochastic finite elements: a spectral approach*. Courier Corporation.
- Ghanizadeh, A. R., Heidarabadizadeh, N., and Heravi, F. (2021). Gaussian process regression (gpr) for auto-estimation of resilient modulus of stabilized base materials. *Journal of Soft Computing in Civil Engineering*, 5(1):80–94.
- Giannakis, D. (2019). Data-driven spectral decomposition and forecasting of ergodic dynamical systems. *Applied and Computational Harmonic Analysis*, 47(2):338–396.
- Ginsbourger, D., Bay, X., Roustant, O., and Carraro, L. (2012). Argumentwise invariant kernels for the approximation of invariant functions. In *Annales de la Faculté des sciences de Toulouse: Mathématiques*, volume 21, pages 501–527.
- Ginsbourger, D., Roustant, O., and Durrande, N. (2013). Invariances of random fields paths, with applications in gaussian process regression. *arXiv preprint arXiv:1308.1359*.
- Goldberg, Y. and Kosorok, M. R. (2017). Support vector regression for right censored data. *Electronic Journal of Statistics*, 11(1):532–569.
- Gong, J.-x., Yi, P., and Zhao, N. (2014). Non-gradient-based algorithm for structural reliability analysis. *Journal of Engineering Mechanics*, 140(6):04014029.

- Gordon, M. S., Fedorov, D. G., Pruitt, S. R., and Slipchenko, L. V. (2011). Fragmentation methods: A route to accurate calculations on large systems. *Chemical reviews*, 112(1):632–672.
- Götteman, M. (2020). Advances and challenges in wave energy park optimization-a review. *Frontiers in Energy Research*, 8(26).
- Grey, Z. J. and Constantine, P. G. (2018). Active subspaces of airfoil shape parameterizations. *AIAA Journal*, 56(5):2003–2017.
- Groot, P. and Lucas, P. (2012). Gaussian process regression with censored data using expectation propagation. In *Sixth European Workshop on Probabilistic Graphical Models*, pages 115–122. Granada: DECSAI.
- Gunn, S. R. et al. (1998). Support vector machines for classification and regression. *ISIS technical report*, 14(1):5–16.
- Haasdonk, B. and Burkhardt, H. (2007). Invariant kernel functions for pattern analysis and machine learning. *Machine learning*, 68(1):35–61.
- Halkjær, S., Sigmund, O., and Jensen, J. S. (2006). Maximizing band gaps in plate structures. *Structural and Multidisciplinary Optimization*, 32(4):263–275.
- Han, Z.-H., Görtz, S., and Zimmermann, R. (2013). Improving variable-fidelity surrogate modeling via gradient-enhanced kriging and a generalized hybrid bridge function. *Aerospace Science and technology*, 25(1):177–189.

- Hanuka, A., Duris, J., Shtalenkova, J., Kennedy, D., Edelen, A., Ratner, D., and Huang, X. (2019). Online tuning and light source control using a physics-informed Gaussian process. *arXiv*, (NeurIPS).
- Hao, P., Feng, S., Li, Y., Wang, B., and Chen, H. (2020). Adaptive infill sampling criterion for multi-fidelity gradient-enhanced kriging model. *Structural and Multidisciplinary Optimization*, 62(1):353–373.
- Harrell, F. E. (2015). Cox proportional hazards regression model. In *Regression modeling strategies*, pages 475–519. Springer.
- Hashimoto, E. M., Ortega, E. M., Cancho, V. G., and Cordeiro, G. M. (2010). The log-exponentiated weibull regression model for interval-censored data. *Computational statistics & data analysis*, 54(4):1017–1035.
- Hassoun, M. H. et al. (1995). *Fundamentals of artificial neural networks*. MIT press.
- Hearst, M. A., Dumais, S. T., Osuna, E., Platt, J., and Scholkopf, B. (1998). Support vector machines. *IEEE Intelligent Systems and their applications*, 13(4):18–28.
- Hinton, G. E. and Salakhutdinov, R. R. (2006). Reducing the dimensionality of data with neural networks. *science*, 313(5786):504–507.
- Hsu, H.-H. and Wu, Y.-C. (1997). The hydrodynamic coefficients for an oscillating rectangular structure on a free surface with sidewall. *Ocean Engineering*, 24(2):177–199.
- Huang, Z., Wang, C., Chen, J., and Tian, H. (2011). Optimal design of aeroengine turbine disc based on kriging surrogate models. *Computers & structures*, 89(1-2):27–37.

- Hurtado, J. E. (2004). An examination of methods for approximating implicit limit state functions from the viewpoint of statistical learning theory. *Structural Safety*, 26(3):271–293.
- Hurtado, J. E. and Alvarez, D. A. (2001). Neural-network-based reliability analysis: a comparative study. *Computer methods in applied mechanics and engineering*, 191(1-2):113–132.
- Jain, A. K., Mao, J., and Mohiuddin, K. M. (1996). Artificial neural networks: A tutorial. *Computer*, 29(3):31–44.
- Jakobsson, S., Patriksson, M., Rudholm, J., and Wojciechowski, A. (2010). A method for simulation based optimization using radial basis functions. *Optimization and Engineering*, 11(4):501–532.
- Jia, G. and Shi, Z. (2010). A new seismic isolation system and its feasibility study. *Earthquake Engineering and Engineering Vibration*, 9(1):75–82.
- Jia, G., Taflanidis, A., Scruggs, J., et al. (2015). Layout optimization of wave energy converters in a random sea. In *The Twenty-fifth International Ocean and Polar Engineering Conference*. International Society of Offshore and Polar Engineers.
- Jia, G. and Taflanidis, A. A. (2013). Kriging metamodeling for approximation of high-dimensional wave and surge responses in real-time storm/hurricane risk assessment. *Computer Methods in Applied Mechanics and Engineering*, 261:24–38.
- Jia, G., Taflanidis, A. A., Nadal-Caraballo, N. C., Melby, J. A., Kennedy, A. B., and Smith, J. M. (2016). Surrogate modeling for peak or time-dependent storm surge prediction over an extended coastal region using an existing database of synthetic storms. *Natural Hazards*, 81(2):909–938.

- Jin, R., Chen, W., and Simpson, T. W. (2001). Comparative studies of metamodeling techniques under multiple modelling criteria. *Structural and multidisciplinary optimization*, 23(1):1–13.
- Jin, R., Chen, W., and Sudjianto, A. (2002). On sequential sampling for global metamodeling in engineering design. In *International design engineering technical conferences and computers and information in engineering conference*, volume 36223, pages 539–548.
- Jin, R., Du, X., and Chen, W. (2003). The use of metamodeling techniques for optimization under uncertainty. *Structural and Multidisciplinary Optimization*, 25(2):99–116.
- Johnson, M. E., Moore, L. M., and Ylvisaker, D. (1990). Minimax and maximin distance designs. *Journal of statistical planning and inference*, 26(2):131–148.
- Jolliffe, I. T. (2002). *Principal Component Analysis*. Springer.
- Jones, D. (2001). A taxonomy of global optimization methods based on response surfaces. *Journal of Global Optimization*, 21:345–383.
- Jones, D. R., Schonlau, M., and Welch, W. J. (1998). Efficient global optimization of expensive black-box functions. *Journal of Global optimization*, 13:455–492.
- Joy, A. A., Hasan, M. A. M., and Hossain, M. A. (2019). A comparison of supervised and unsupervised dimension reduction methods for hyperspectral image classification. In *2019 International Conference on Electrical, Computer and Communication Engineering (ECCE)*, pages 1–6. IEEE.

- Kapsoulis, D., Tsiakas, K., Asouti, V., and Giannakoglou, K. (2016). The use of kernel pca in evolutionary optimization for computationally demanding engineering applications. In *2016 IEEE Symposium Series on Computational Intelligence (SSCI)*, pages 1–8. IEEE.
- Karniadakis, G. E., Kevrekidis, I. G., Lu, L., Perdikaris, P., Wang, S., and Yang, L. (2021). Physics-informed machine learning. *Nature Reviews Physics*, 3(6):422–440.
- Kaymaz, I. (2005). Application of kriging method to structural reliability problems. *Structural Safety*, 27(2):133–151.
- Keane, A. J. (2012). Cokriging for robust design optimization. *AIAA journal*, 50(11):2351–2364.
- Kennedy, M. C. and O’Hagan, A. (2000). Predicting the output from a complex computer code when fast approximations are available. *Biometrika*, 87(1):1–13.
- Khan, F. M. and Zubek, V. B. (2008). Support vector regression for censored data (svrc): a novel tool for survival analysis. In *2008 Eighth IEEE International Conference on Data Mining*, pages 863–868. IEEE.
- Kim, S.-H. and Das, M. P. (2012). Seismic waveguide of metamaterials. *Modern Physics Letters B*, 26(17):1250105.
- Kingma, D. P. and Welling, M. (2019). An introduction to variational autoencoders. *arXiv preprint arXiv:1906.02691*.
- Kokkinowrachos, K., Mavrakos, S., and Asorakos, S. (1986). Behaviour of vertical bodies of revolution in waves. *Ocean Engineering*, 13(6):505–538.
- Kondor, I. R. (2008). *Group theoretical methods in machine learning*. Columbia University.

- Kostov, C. and Dubrule, O. (1986). An interpolation method taking into account inequality constraints: II. practical approach. *Mathematical Geology*, 18(1):53–73.
- Kuya, Y., Takeda, K., Zhang, X., and J. Forrester, A. I. (2011). Multifidelity surrogate modeling of experimental and computational aerodynamic data sets. *AIAA journal*, 49(2):289–298.
- Lataniotis, C., Marelli, S., and Sudret, B. (2018). Extending classical surrogate modelling to high-dimensions through supervised dimensionality reduction: a data-driven approach. *arXiv preprint arXiv:1812.06309*.
- Lataniotis, C., Marelli, S., and Sudret, B. (2020). Extending classical surrogate modeling to high dimensions through supervised dimensionality reduction: a data-driven approach. *International Journal for Uncertainty Quantification*, 10(1).
- Lawrence, S., Giles, C. L., Tsoi, A. C., and Back, A. D. (1997). Face recognition: A convolutional neural-network approach. *IEEE transactions on neural networks*, 8(1):98–113.
- Le Gratiet, L. (2013a). Bayesian analysis of hierarchical multifidelity codes. *SIAM/ASA Journal on Uncertainty Quantification*, 1(1):244–269.
- Le Gratiet, L. (2013b). *Multi-fidelity Gaussian process regression for computer experiments*. PhD thesis, Université Paris-Diderot-Paris VII.
- Le Gratiet, L. L., Marelli, S., and Sudret, B. (2017). Metamodel-based sensitivity analysis: Polynomial chaos expansions and gaussian processes. In *Handbook of Uncertainty Quantification*, pages 1289–1325.

- LeCun, Y., Bengio, Y., et al. (1995). Convolutional networks for images, speech, and time series. *The handbook of brain theory and neural networks*, 3361(10):1995.
- Lee, J., Lee, Y., Kim, J., Kosiorek, A., Choi, S., and Teh, Y. W. (2019). Set transformer: A framework for attention-based permutation-invariant neural networks. In *International Conference on Machine Learning*, pages 3744–3753. PMLR.
- Li, G., Rosenthal, C., and Rabitz, H. (2001). High dimensional model representations. *The Journal of Physical Chemistry A*, 105(33):7765–7777.
- Li, G., Wang, S.-W., and Rabitz, H. (2000). High dimensional model representations (hdmr): Concepts and applications. In *Proceedings of the Institute of Mathematics and Its Applications Workshop on Atmospheric Modeling*, pages 15–19. Citeseer.
- Li, H., Li, P., Gao, L., Zhang, L., and Wu, T. (2015). A level set method for topological shape optimization of 3d structures with extrusion constraints. *Computer Methods in Applied Mechanics and Engineering*, 283:615–635.
- Li, J., Cai, J., and Qu, K. (2019a). Surrogate-based aerodynamic shape optimization with the active subspace method. *Structural and Multidisciplinary Optimization*, 59(2):403–419.
- Li, M., Cheng, Z., Jia, G., and Shi, Z. (2019b). Dimension reduction and surrogate based topology optimization of periodic structures. *Composite Structures*, 229:111385.
- Li, M. and Jia, G. (2020). Multifidelity gaussian process model integrating low-and high-fidelity data considering censoring. *Journal of Structural Engineering*, 146(3):04019215.

- Li, M., Wang, R.-Q., and Jia, G. (2020). Efficient dimension reduction and surrogate-based sensitivity analysis for expensive models with high-dimensional outputs. *Reliability Engineering & System Safety*, 195:106725.
- Li, Y., Zhou, X., Bian, Z., Xing, Y., and Song, J. (2017). Thermal tuning of the interfacial adhesive layer on the band gaps in a one-dimensional phononic crystal. *Composite Structures*, 172:311 – 318.
- Lin, D., Feuer, E., Etzioni, R., and Wax, Y. (1997). Estimating medical costs from incomplete follow-up data. *Biometrics*, pages 419–434.
- Liu, H., Ong, Y.-S., Cai, J., and Wang, Y. (2018a). Cope with diverse data structures in multi-fidelity modeling: a gaussian process method. *Engineering Applications of Artificial Intelligence*, 67:211–225.
- Liu, Y., Chen, S., Wang, F., and Xiong, F. (2018b). Sequential optimization using multi-level cokriging and extended expected improvement criterion. *Structural and Multidisciplinary Optimization*, 58(3):1155–1173.
- Lophaven, S. N., Nielsen, H. B., Søndergaard, J., et al. (2002). *DACE: a Matlab kriging toolbox*, volume 2. Citeseer.
- Lucia, D. J., Beran, P. S., and Silva, W. A. (2004). Reduced-order modeling: new approaches for computational physics. *Progress in aerospace sciences*, 40(1-2):51–117.
- Luck, M., Sylvain, T., Cardinal, H., Lodi, A., and Bengio, Y. (2017). Deep learning for patient-specific kidney graft survival analysis. *arXiv preprint arXiv:1705.10245*.

- Luo, J. and Lu, W. (2014). Comparison of surrogate models with different methods in groundwater remediation process. *Journal of Earth System Science*, 123(7):1579–1589.
- Marino, D. L. and Manic, M. (2019). Combining physicsbased domain knowledge and machine learning using variational gaussian processes with explicit linear prior. *arXiv preprint arXiv:1906.02160*.
- Mavrakos, S. (1991). Hydrodynamic coefficients for groups of interacting vertical axisymmetric bodies. *Ocean Engineering*, 18(5):485–515.
- Mavrakos, S. and Koumoutsakos, P. (1987). Hydrodynamic interaction among vertical axisymmetric bodies restrained in waves. *Applied Ocean Research*, 9(3):128–140.
- Mavrakos, S. and McIver, P. (1997). Comparison of methods for computing hydrodynamic characteristics of arrays of wave power devices. *Applied Ocean Research*, 19(5-6):283–291.
- McIver, P. (1994). Some hydrodynamic aspects of arrays of wave-energy devices. *Applied Ocean Research*, 16(2):61–69.
- McIver, P. (2002). Wave interaction with arrays of structures. *Applied Ocean Research*, 24(3):121–126.
- McKay, M. D., Beckman, R. J., and Conover, W. J. (2000). A comparison of three methods for selecting values of input variables in the analysis of output from a computer code. *Technometrics*, 42(1):55–61.
- Militino, A. F. and Ugarte, M. D. (1999). Analyzing censored spatial data. *Mathematical Geology*, 31(5):551–561.

- Minka, T. (2001). *A family of approximation methods for approximate Bayesian inference*. PhD thesis, PhD thesis, MIT.
- Minka, T. P. (2013). Expectation propagation for approximate bayesian inference. *arXiv preprint arXiv:1301.2294*.
- Muralitharan, K., Sakthivel, R., and Vishnuvarthan, R. (2018). Neural network based optimization approach for energy demand prediction in smart grid. *Neurocomputing*, 273:199–208.
- Ng, A. et al. (2011). Sparse autoencoder. *CS294A Lecture notes*, 72(2011):1–19.
- Ng, L. W.-T. and Eldred, M. (2012). Multifidelity uncertainty quantification using non-intrusive polynomial chaos and stochastic collocation. In *53rd AIAA/ASME/ASCE/AHS/ASC Structures, Structural Dynamics and Materials Conference 20th AIAA/ASME/AHS Adaptive Structures Conference 14th AIAA*, page 1852.
- Niyogi, P., Girosi, F., and Poggio, T. (1998). Incorporating prior information in machine learning by creating virtual examples. *Proceedings of the IEEE*, 86(11):2196–2209.
- Paiva, R. M., Carvalho, A. R., Crawford, C., and Suleman, A. (2010). Comparison of surrogate models in a multidisciplinary optimization framework for wing design. *AIAA journal*, 48(5):995–1006.
- Palar, P. S. and Shimoyama, K. (2018). On the accuracy of kriging model in active subspaces. In *2018 AIAA/ASCE/AHS/ASC Structures, Structural Dynamics, and Materials Conference*, page 0913.

- Palar, P. S., Zuhail, L. R., Shimoyama, K., and Tsuchiya, T. (2018). Global sensitivity analysis via multi-fidelity polynomial chaos expansion. *Reliability Engineering & System Safety*, 170:175–190.
- Papadopoulos, I., Ehre, M., and Straub, D. (2019). Pls-based adaptation for efficient pce representation in high dimensions. *Journal of Computational Physics*, 387:186–204.
- Picheny, V., Ginsbourger, D., Roustant, O., Haftka, R. T., and Kim, N.-H. (2010). Adaptive designs of experiments for accurate approximation of a target region. *Journal of Mechanical Design*, 132(7):071008.
- Prakash, S., Mamun, K., Islam, F., Mudliar, R., Pau'u, C., Kolivuso, M., and Cadralala, S. (2016). Wave energy converter: a review of wave energy conversion technology. In *2016 3rd Asia-Pacific World Congress on Computer Science and Engineering (APWC on CSE)*, pages 71–77. IEEE.
- Prataviera, F., Ortega, E. M., Cordeiro, G. M., and Braga, A. d. S. (2019). The heteroscedastic odd log-logistic generalized gamma regression model for censored data. *Communications in Statistics-Simulation and Computation*, 48(6):1815–1839.
- Priestley, M. J. N., Seible, F., and Calvi, G. M. (1996). *Seismic Design and Retrofit of Bridges*. John Wiley & Sons.
- Pujol, S. (2002). *Drift Capacity of Reinforced Concrete Columns Subjected to Displacement Reversals*. PhD thesis, Purdue University, West Lafayette, IN.
- Qian, P. Z. and Wu, C. J. (2008). Bayesian hierarchical modeling for integrating low-accuracy and high-accuracy experiments. *Technometrics*, 50(2):192–204.

- Quesenberry Jr, C. P., Fireman, B., Hiatt, R. A., and Selby, J. V. (1989). A survival analysis of hospitalization among patients with acquired immunodeficiency syndrome. *American journal of public health*, 79(12):1643–1647.
- Raissi, M. (2018). Deep hidden physics models: Deep learning of nonlinear partial differential equations. *The Journal of Machine Learning Research*, 19(1):932–955.
- Raissi, M., Perdikaris, P., and Karniadakis, G. E. (2019). Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378:686–707.
- Ramsay, J. O. and Dalzell, C. (1991). Some tools for functional data analysis. *Journal of the Royal Statistical Society: Series B (Methodological)*, 53(3):539–561.
- Raponi, E., Wang, H., Bujny, M., Boria, S., and Doerr, C. (2020). High dimensional bayesian optimization assisted by principal component analysis. In *International Conference on Parallel Problem Solving from Nature*, pages 169–183. Springer.
- Rasmussen, C. E. (2004). Gaussian processes in machine learning. In *Advanced lectures on machine learning*, pages 63–71. Springer.
- Regis, R. G. and Shoemaker, C. A. (2005). Constrained global optimization of expensive black box functions using radial basis functions. *Journal of Global Optimization*, 31(1):153–171.
- Roweis, S. T. and Saul, L. K. (2000). Nonlinear dimensionality reduction by locally linear embedding. *science*, 290(5500):2323–2326.

- Ruszczynski, A. and Shapiro, A. (2003). *Stochastic programming, handbooks in operations research and management science*. Elsevier.
- Sacks, J., Welch, W. J., Mitchell, T. J., and Wynn, H. P. (1989). Design and analysis of computer experiments. *Statistical science*, 4(4):409–423.
- Schein, A., Saul, L., and Ungar, L. (2003). A generalized linear model for principal component analysis of binary data. *Proceedings of the 9th International Workshop on Artificial Intelligence and Statistics*, page 546431.
- Scruggs, J., Lattanzio, S., Taflanidis, A., and Cassidy, I. (2013). Optimal causal control of a wave energy converter in a random sea. *Applied Ocean Research*, 42:1–15.
- Shewry, M. C. and Wynn, H. P. (1987). Maximum entropy sampling. *Journal of applied statistics*, 14(2):165–170.
- Shivaswamy, P. K., Chu, W., and Jansche, M. (2007). A support vector approach to censored targets. In *Seventh IEEE international conference on data mining (ICDM 2007)*, pages 655–660. IEEE.
- Sigmund, O. and Søndergaard Jensen, J. (2003). Systematic design of phononic band-gap materials and structures by topology optimization. *Philosophical Transactions of the Royal Society of London. Series A: Mathematical, Physical and Engineering Sciences*, 361(1806):1001–1019.
- Simpson, T. W., Mauery, T. M., Korte, J. J., and Mistree, F. (2001a). Kriging models for global approximation in simulation-based multidisciplinary design optimization. *AIAA journal*, 39(12):2233–2241.

- Simpson, T. W., Poplinski, J., Koch, P. N., and Allen, J. K. (2001b). Metamodels for computer-based engineering design: survey and recommendations. *Engineering with computers*, 17(2):129–150.
- Sobol', I. M. (2003). Theorems and examples on high dimensional model representation. *Reliability Engineering and System Safety*, 79(2):187–193.
- Sobol', I. (1993). Sensitivity estimates for nonlinear mathematical models. *Math. Model. Comput. Exp*, 1(4):407–414.
- Stein, M. L. (1992). Prediction and inference for truncated spatial data. *Journal of Computational and Graphical Statistics*, 1(1):91–110.
- Steinwart, I. and Christmann, A. (2008). *Support vector machines*. Springer Science & Business Media.
- Sudret, B. (2008). Global sensitivity analysis using polynomial chaos expansions. *Reliability Engineering & System Safety*, 93(7):964–979.
- Sudret, B. and Der Kiureghian, A. (2002). Comparison of finite element reliability methods. *Probabilistic Engineering Mechanics*, 17(4):337–348.
- Sudret, B., Lataniotis, C., Lüthen, N., Marelli, S., and Torre, E. (2019). Surrogate modelling meets machine learning. In *3rd International Conference on Uncertainty Quantification in Computational Sciences and Engineering (UNCECOMP 2019)*, pages U–19257. ETH Zurich, Chair of Risk, Safety and Uncertainty Quantification.

- Sun, L., Gao, H., Pan, S., and Wang, J.-X. (2020). Surrogate modeling for fluid flows based on physics-constrained deep learning without simulation data. *Computer Methods in Applied Mechanics and Engineering*, 361:112732.
- Taflanidis, A. A., Jia, G., Kennedy, A. B., and Smith, J. M. (2013). Implementation/optimization of moving least squares response surfaces for approximation of hurricane/storm surge and wave responses. *Natural Hazards*, 66(2):955–983.
- Tai, K. S., Bailis, P., and Valiant, G. (2019). Equivariant transformer networks. In *International Conference on Machine Learning*, pages 6086–6095. PMLR.
- Tao, S., Apley, D. W., Chen, W., Garbo, A., Pate, D. J., and German, B. J. (2019). Input mapping for model calibration with application to wing aerodynamics. *AIAA journal*, 57(7):2734–2745.
- Tao, S., Shintani, K., Bostanabad, R., Chan, Y.-C., Yang, G., Meingast, H., and Chen, W. (2017). Enhanced gaussian process metamodeling and collaborative optimization for vehicle suspension design optimization. In *International Design Engineering Technical Conferences and Computers and Information in Engineering Conference*, volume 58134, page V02BT03A039. American Society of Mechanical Engineers.
- Tipping, M. E. and Bishop, C. M. (1999). Probabilistic principal component analysis. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 61(3):611–622.
- Tong, M. A., Barajas-Solano, D. A., Tipireddy, R., and Tartakovsky, A. (2020). Physics-informed gaussian process regression for probabilistic states estimation and forecasting in power grids.

- Tripathy, R., Bilonis, I., and Gonzalez, M. (2016). Gaussian processes with built-in dimensionality reduction: Applications to high-dimensional uncertainty propagation. *Journal of Computational Physics*, 321:191–223.
- Twersky, V. (1952). Multiple scattering of radiation by an arbitrary configuration of parallel cylinders. *The Journal of the Acoustical Society of America*, 24(1):42–46.
- U.S. Department of Energy (2016). Water Power Program. Water power for a clean energy future. Technical Report March.
- Valsesia, D. and Magli, E. (2021). Permutation invariance and uncertainty in multitemporal image super-resolution. *arXiv preprint arXiv:2105.12409*.
- van der Wilk, M., Bauer, M., John, S., and Hensman, J. (2018). Learning invariances using the marginal likelihood. In *Advances in Neural Information Processing Systems*, pages 9938–9948.
- Vasseur, J., Matar, O. B., Robillard, J., Hladky-Hennion, A.-C., and Deymier, P. (2011). Band structures tunability of bulk 2d phononic crystals made of magneto-elastic materials. *AIP advances*, 1(4):041904.
- Verleysen, M. and François, D. (2005). The curse of dimensionality in data mining and time series prediction. In *International work-conference on artificial neural networks*, pages 758–770. Springer.
- Viana, F. A., Gogu, C., and Haftka, R. T. (2010). Making the most out of surrogate models: tricks of the trade. In *International design engineering technical conferences and computers and information in engineering conference*, volume 44090, pages 587–598.

- Wang, G. G. (2003). Adaptive response surface method using inherited latin hypercube design points. *Journal of Mechanical Design*, 125(2):210–220.
- Wang, Y. and Kang, Z. (2019). Concurrent two-scale topological design of multiple unit cells and structure using combined velocity field level set and density model. *Computer Methods in Applied Mechanics and Engineering*, 347:340 – 364.
- Willcox, K. and Megretski, A. (2005). Fourier series for accurate, stable, reduced-order models in large-scale linear applications. *SIAM Journal on Scientific Computing*, 26(3):944–962.
- Willcox, K. E. (2000). *Reduced-order aerodynamic models for aeroelastic control of turbomachines*. PhD thesis, Massachusetts Institute of Technology.
- Willmes, L., Back, T., Jin, Y., and Sendhoff, B. (2003). Comparing neural networks and kriging for fitness approximation in evolutionary optimization. In *The 2003 Congress on Evolutionary Computation, 2003. CEC'03.*, volume 1, pages 663–670. IEEE.
- Wu, J. L., Michelén-Ströfer, C., and Xiao, H. (2019). Physics-informed covariance kernel for model-form uncertainty quantification with application to turbulent flows. *Computers and Fluids*, 193:104292.
- Wu, W., Massart, D., and De Jong, S. (1997). The kernel pca algorithms for wide data. part i: theory and algorithms. *Chemometrics and Intelligent Laboratory Systems*, 36(2):165–172.
- Wu, Y.-F., Oehlers, D. J., and Griffith, M. C. (2004). Rational definition of the flexural deformation capacity of rc column sections. *Engineering Structures*, 26(5):641–650.

- Xia, Q. and Shi, T. (2018). A cascadic multilevel optimization algorithm for the design of composite structures with curvilinear fiber based on shepard interpolation. *Composite Structures*, 188:209–219.
- Xia, Q., Shi, T., Wang, M. Y., and Liu, S. (2010). A level set based method for the optimization of cast part. *Structural and Multidisciplinary Optimization*, 41(5):735–747.
- Xie, L., Xia, B., Huang, G., Lei, J., and Liu, J. (2017). Topology optimization of phononic crystals with uncertainties. *Structural and Multidisciplinary Optimization*, 56(6):1319–1339.
- Xiong, Y., Chen, W., Apley, D., and Ding, X. (2007). A non-stationary covariance-based kriging method for metamodeling in engineering design. *International Journal for Numerical Methods in Engineering*, 71(6):733–756.
- Yang, X., Tartakovsky, G., and Tartakovsky, A. (2018). Physics-informed kriging: A physics-informed gaussian process regression method for data-model convergence. *arXiv*.
- Yao, K., Herr, J. E., and Parkhill, J. (2017). The many-body expansion combined with neural networks. *The Journal of chemical physics*, 146(1):014106.
- Yi, G. and Youn, B. D. (2016). A comprehensive survey on topology optimization of phononic crystals. *Structural and Multidisciplinary Optimization*, 54(5):1315–1344.
- Zaheer, M., Kottur, S., Ravanbakhsh, S., Poczos, B., Salakhutdinov, R., and Smola, A. (2017). Deep sets. *arXiv preprint arXiv:1703.06114*.
- Zhang, H., Luo, Y., and Kang, Z. (2018). Bi-material microstructural design of chiral auxetic metamaterials using topology optimization. *Composite Structures*, 195:232 – 248.

- Zhang, J., Taflanidis, A. A., and Medina, J. C. (2017). Sequential approximate optimization for design under uncertainty problems utilizing Kriging metamodeling in augmented input space. *Computer Methods in Applied Mechanics and Engineering*, 315:369–395.
- Zhang, J., Taflanidis, A. A., and Scruggs, J. T. (2020a). Surrogate modeling of hydrodynamic forces between multiple floating bodies through a hierarchical interaction decomposition. *Journal of Computational Physics*, 408:109298.
- Zhang, J., Zeng, L., Chen, C., Chen, D., and Wu, L. (2015). Efficient bayesian experimental design for contaminant source identification. *Water Resources Research*, 51(1):576–598.
- Zhang, R., Liu, Y., and Sun, H. (2020b). Physics-guided convolutional neural network (phycnn) for data-driven seismic response modeling. *Engineering Structures*, 215:110704.
- Zhao, L. and Feng, D. (2020). Deep neural networks for survival analysis using pseudo values. *IEEE journal of biomedical and health informatics*, 24(11):3308–3314.
- Zhong, H.-L., Wu, F.-G., and Yao, L.-N. (2006). Application of genetic algorithm in optimization of band gap of two-dimensional phononic crystals. *Acta Physica Sinica*, 55(1):275–279.
- Zhou, M., Fleury, R., Shyy, Y.-K., Thomas, H., and Brennan, J. (2002). Progress in topology optimization with manufacturing constraints. In *9th AIAA/ISSMO Symposium on multidisciplinary analysis and optimization*, page 5614.
- Zhou, T. and Peng, Y. (2020). Kernel principal component analysis-based gaussian process regression modelling for high-dimensional reliability analysis. *Computers & Structures*, 241:106358.

- Zhu, X., Yao, J., and Huang, J. (2016). Deep convolutional neural network for survival analysis with pathological images. In *2016 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*, pages 544–547. IEEE.
- Zhu, Y., Zabaras, N., Koutsourelakis, P.-S., and Perdikaris, P. (2019). Physics-constrained deep learning for high-dimensional surrogate modeling and uncertainty quantification without labeled data. *Journal of Computational Physics*, 394:56–81.
- Zou, Q. and Chen, S. (2021). Resilience-based recovery scheduling of transportation network in mixed traffic environment: A deep-ensemble-assisted active learning approach. *Reliability Engineering & System Safety*, page 107800.

## APPENDIX A:

### DERIVATION OF MARGINAL POSTERIOR DISTRIBUTION FOR $\boldsymbol{\theta}$ IN MULTI-FIDELITY GAUSSIAN PROCESS MODEL

To facilitate understanding and also for completeness, here we show the detailed derivation of the marginal posterior distribution for  $\boldsymbol{\theta} = [\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta]$  in the multi-fidelity Gaussian process model in Chapter 4.

Based on the posterior distributions for  $\boldsymbol{\theta}_l$  and  $\boldsymbol{\theta}_\delta$  given in Eq. (4.20) and Eq. (4.21), the joint posterior distribution of  $\boldsymbol{\theta}_l$  and  $\boldsymbol{\theta}_\delta$  conditional on given values of  $\bar{\boldsymbol{\theta}} = [\boldsymbol{\beta}_l, \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_l^2, \sigma_\delta^2, \sigma_\epsilon^2]$  can be written as

$$\begin{aligned}
 & p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\boldsymbol{\theta}}) \\
 & \propto p(\boldsymbol{\theta}_l) p(\boldsymbol{\theta}_\delta) \frac{1}{(\sigma_l^2)^{n_l/2} |\mathbf{R}_l|^{1/2}} \exp \left[ -\frac{(\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l)^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l) \boldsymbol{\beta}_l)}{2\sigma_l^2} \right] \frac{1}{(\sigma_\delta^2)^{n_h/2} |\mathbf{R}_{\xi\delta}|^{1/2}} \\
 & \times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h) \boldsymbol{\beta}_\rho - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta)}{2\sigma_\delta^2} \right] \quad (\text{A.1})
 \end{aligned}$$

The marginal posterior distribution for  $\boldsymbol{\theta}$  can be calculated as

$$\begin{aligned}
 p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l) &= \int p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l, \bar{\boldsymbol{\theta}}) p(\bar{\boldsymbol{\theta}}) d\bar{\boldsymbol{\theta}} \\
 &= p(\boldsymbol{\theta}_l) p(\boldsymbol{\theta}_\delta) \int p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \bar{\boldsymbol{\theta}}) p(\bar{\boldsymbol{\theta}}) d\bar{\boldsymbol{\theta}} \quad (\text{A.2})
 \end{aligned}$$

To simplify the above integral, we will analytically integrate out the uncertain parameters  $[\boldsymbol{\beta}_l, \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta, \sigma_l^2]$  in  $\bar{\boldsymbol{\theta}}$ , which will reduce the integral to integration with respect to  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$ . More specifically, we analytically and sequentially integrate out  $\boldsymbol{\beta}_l, \boldsymbol{\beta}_\rho, \boldsymbol{\beta}_\delta$ , and  $\sigma_l^2$ .

(1) Integrate out  $\beta_l$

$$\begin{aligned}
& \int_{\beta_l} p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \bar{\boldsymbol{\theta}}) p(\beta_l) d\beta_l \\
& \propto \int_{\beta_l} (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-n_h/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{V}_l^{-1}|^{-1/2} \exp \left[ -\frac{(\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l)\beta_l)^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathbf{F}_l(\mathbf{x}_l)\beta_l)}{2\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)}{2\sigma_\delta^2} \right] \\
& \times \exp \left[ -\frac{(\beta_l - \mathbf{m}_l)^T (v_l \mathbf{I}) (\beta_l - \mathbf{m}_l)}{2\sigma_l^2} \right] d\beta_l \\
& \propto (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-n_h/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)}{2\sigma_\delta^2} \right]
\end{aligned} \tag{A.3}$$

where  $\mathbf{a}_l = \mathbf{F}_l(\mathbf{x}_l)^T \mathbf{R}_l^{-1} \mathbf{F}_l(\mathbf{x}_l) + v_l \mathbf{I}$ ,  $\mathbf{b}_l = -2\mathbf{F}_l(\mathbf{x}_l)^T \mathbf{R}_l^{-1} \mathbf{y}_l - 2(v_l \mathbf{I}) \mathbf{m}_l$ , and  $c_l = \mathbf{y}_l^T \mathbf{R}_l^{-1} \mathbf{y}_l + \mathbf{m}_l^T (v_l \mathbf{I}) \mathbf{m}_l$ .

(2) Integrate out  $\beta_\rho$

$$\begin{aligned}
& \int_{\beta_\rho} p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \bar{\boldsymbol{\theta}}) p(\beta_\rho) d\beta_\rho \\
& \propto \int_{\beta_\rho} (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-n_h/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{V}_\rho^{-1}|^{-1/2} |\mathbf{a}_l|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{(\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)^T \mathbf{R}_{\xi\delta}^{-1} (\mathbf{y}_h - \mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)\beta_\rho - \mathbf{F}_\delta(\mathbf{x}_h)\beta_\delta)}{2\sigma_\delta^2} \right] \\
& \times \exp \left[ -\frac{(\beta_\rho - \mathbf{m}_\rho)^T (\sigma_\delta^2 \mathbf{V}_\rho) (\beta_\rho - \mathbf{m}_\rho)}{2\sigma_\delta^2} \right] d\beta_\rho \\
& \propto (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{4c_\rho - \mathbf{b}_\rho^T \mathbf{a}_\rho^{-1} \mathbf{b}_\rho}{8\sigma_\delta^2} \right]
\end{aligned} \tag{A.4}$$

where

$$\begin{aligned}\mathbf{a}_\rho &= [\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} [\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)] + \sigma_\delta^2 \mathbf{V}_\rho \\ \mathbf{b}_\rho &= -2[\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} [\mathbf{y}_h - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta] - 2(\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho \\ c_\rho &= [\mathbf{y}_h - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta]^T \mathbf{R}_{\xi\delta}^{-1} [\mathbf{y}_h - \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta] + \mathbf{m}_\rho^T (\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho\end{aligned}$$

(3) Integrate out  $\boldsymbol{\beta}_\delta$ : Before integrating out  $\boldsymbol{\beta}_\delta$ , we first simplify our notations. The  $\mathbf{b}_\rho$  and  $c_\rho$  expressions above can be expanded as

$$\mathbf{b}_\rho = -2[\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h - 2(\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho + 2[\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta \quad (\text{A.5})$$

$$c_\rho = \mathbf{y}_h^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h + \mathbf{m}_\rho^T (\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho - 2\boldsymbol{\beta}_\delta^T \mathbf{F}_\delta(\mathbf{x}_h)^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h + \boldsymbol{\beta}_\delta^T \mathbf{F}_\delta(\mathbf{x}_h)^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{F}_\delta(\mathbf{x}_h) \boldsymbol{\beta}_\delta \quad (\text{A.6})$$

Let  $\mathbf{A} = -2[\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h - 2(\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho$  and  $\mathbf{B} = 2[\mathbf{A}_{y_l} \mathbf{F}_\rho(\mathbf{x}_h)]^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{F}_\delta(\mathbf{x}_h)$ ,  $C = \mathbf{y}_h^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h + \mathbf{m}_\rho^T (\sigma_\delta^2 \mathbf{V}_\rho) \mathbf{m}_\rho$ ,  $\mathbf{D} = -2\mathbf{F}_\delta(\mathbf{x}_h)^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{y}_h$ , and  $\mathbf{E} = \mathbf{F}_\delta(\mathbf{x}_h)^T \mathbf{R}_{\xi\delta}^{-1} \mathbf{F}_\delta(\mathbf{x}_h)$ , then we have  $\mathbf{b}_\rho = \mathbf{A} + \mathbf{B}\boldsymbol{\beta}_\delta$ ,  $c_\rho = C + \boldsymbol{\beta}_\delta^T \mathbf{D} + \boldsymbol{\beta}_\delta^T \mathbf{E} \boldsymbol{\beta}_\delta$ . Therefore, the  $4c_\rho - \mathbf{b}_\rho^T \mathbf{a}_\rho^{-1} \mathbf{b}_\rho$  term in Eq. (A.4) can be written as

$$\begin{aligned}4c_\rho - \mathbf{b}_\rho^T \mathbf{a}_\rho^{-1} \mathbf{b}_\rho &= 4(C + \boldsymbol{\beta}_\delta^T \mathbf{D} + \boldsymbol{\beta}_\delta^T \mathbf{E} \boldsymbol{\beta}_\delta) - (\mathbf{A} + \mathbf{B}\boldsymbol{\beta}_\delta)^T \mathbf{a}_\rho^{-1} (\mathbf{A} + \mathbf{B}\boldsymbol{\beta}_\delta) \\ &= (4C - \mathbf{A}^T \mathbf{a}_\rho^{-1} \mathbf{A}) + \boldsymbol{\beta}_\delta^T (4\mathbf{D} - 2\mathbf{B}^T \mathbf{a}_\rho^{-1} \mathbf{A}) + \boldsymbol{\beta}_\delta^T (4\mathbf{E} - \mathbf{B}^T \mathbf{a}_\rho^{-1} \mathbf{B}) \boldsymbol{\beta}_\delta\end{aligned} \quad (\text{A.7})$$

Further, let  $t_1 = 4C - \mathbf{A}^T \mathbf{a}_\rho^{-1} \mathbf{A}$ ,  $t_2 = 4\mathbf{D} - 2\mathbf{B}^T \mathbf{a}_\rho^{-1} \mathbf{A}$ , and  $t_3 = 4\mathbf{E} - \mathbf{B}^T \mathbf{a}_\rho^{-1} \mathbf{B}$ , we can integrate out  $\beta_\delta$ , which leads to the following expression

$$\begin{aligned}
& \int_{\beta_\delta} p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \bar{\boldsymbol{\theta}}) p(\beta_\delta) d\beta_\delta \\
& \propto \int_{\beta_\delta} (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{V}_\delta^{-1}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{t_1 + \beta_\delta^T \mathbf{t}_2 + \beta_\delta^T \mathbf{t}_3 \beta_\delta}{8\sigma_\delta^2} \right] \exp \left[ -\frac{(\beta_\delta - \mathbf{m}_\delta)^T (v_\delta \mathbf{I}) (\beta_\delta - \mathbf{m}_\delta)}{2\sigma_\delta^2} \right] d\beta_\delta \\
& \propto (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} \right]
\end{aligned} \tag{A.8}$$

where  $\mathbf{a}_\delta = \mathbf{t}_3 + v_\delta \mathbf{I}$ ,  $\mathbf{b}_\delta = \mathbf{t}_2 - 2(v_\delta \mathbf{I}) \mathbf{m}_\delta$ , and  $c_\delta = t_1 + \mathbf{m}_\delta^T (v_\delta \mathbf{I}) \mathbf{m}_\delta$ .

(4) Integrate out  $\sigma_l^2$ : Lastly, we integrate out  $\sigma_l^2$ , which leads to the following expression

$$\begin{aligned}
& \int_{\sigma_l^2} p(\mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l | \boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta, \bar{\boldsymbol{\theta}}) p(\sigma_l^2) d\sigma_l^2 \\
& \propto \int_{\sigma_l^2} (\sigma_l^2)^{-n_l/2} (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \exp \left[ -\frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8\sigma_l^2} \right] \\
& \times \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} \right] (\sigma_l^2)^{-\alpha_l - 1} \exp \left[ -\frac{\gamma_l}{\sigma_l^2} \right] d\sigma_l^2 \\
& \propto (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \left( \gamma_l + \frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8} \right)^{-\alpha_l - n_l/2} \\
& \times \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} \right]
\end{aligned} \tag{A.9}$$

In the end, the marginal posterior distribution in Eq. (A.2) can be simplified to

$$\begin{aligned}
& p(\boldsymbol{\theta}_l, \boldsymbol{\theta}_\delta | \mathbf{y}_{h,I}, \mathbf{y}_{h,J}, \mathbf{y}_l) \\
& \propto \int_{\sigma_\delta^2, \sigma_\epsilon^2} p(\boldsymbol{\theta}_l) p(\boldsymbol{\theta}_\delta) (\sigma_\delta^2)^{-(n_h - q_\rho)/2} |\mathbf{R}_l|^{-1/2} |\mathbf{R}_{\xi\delta}|^{-1/2} |\mathbf{a}_l|^{-1/2} |\mathbf{a}_\rho|^{-1/2} |\mathbf{a}_\delta|^{-1/2} \\
& \times \left( \gamma_l + \frac{4c_l - \mathbf{b}_l^T \mathbf{a}_l^{-1} \mathbf{b}_l}{8} \right)^{-\alpha_l - n_l/2} \exp \left[ -\frac{4c_\delta - \mathbf{b}_\delta^T \mathbf{a}_\delta^{-1} \mathbf{b}_\delta}{8\sigma_\delta^2} \right] d\sigma_\delta^2 d\sigma_\epsilon^2
\end{aligned} \tag{A.10}$$

which reduces the integral in Eq. (A.2) to integration with respect to only  $\sigma_\delta^2$  and  $\sigma_\epsilon^2$ .

APPENDIX B:

HIGH-FIDELITY DATABASE FOR DEVELOPING MULTI-FIDELITY GAUSSIAN  
PROCESS MODEL TO PREDICT DEFORMATION CAPACITY OF REINFORCED  
CONCRETE COLUMNS

This appendix presents the high-fidelity database for developing the multi-fidelity Gaussian process model to predict deformation capacity of reinforced concrete (RC) columns under repeated cyclic loading in Chapter 4. This database is established based on the existing experimental tests which have been collected at <https://nisee.berkeley.edu/spd/index.html>. As described in Chapter 4, the high-fidelity database after preprocessing consists of 17 accurate responses (i.e., with direct measurement of the deformation capacity) and 53 censored responses (i.e., with measurement of the deformation capacity values before reaching failure). Here Table B.1 and Table B.2 show the accurate data and censored data, respectively. Note that the second column to the twelfth column of both tables correspond to the model input  $x_1$  to  $x_{11}$ , while the last column is the model output, i.e., deformation capacity. In addition, Table B.3 shows the description of the involved notations and the corresponding units.

**Table B.1:** Accurate high-fidelity data

| Data ID | $D_{col}$ | $H_c$  | $c$  | $f_c$ | $d_s$ | $f_{ys}$ | $\rho_s$ | $d_{sv}$ | $f_{yv}$ | $\rho_{sv}$ | $N$     | $\delta$ |
|---------|-----------|--------|------|-------|-------|----------|----------|----------|----------|-------------|---------|----------|
| 1       | 400.0     | 1600.0 | 18.0 | 28.5  | 16.0  | 308.0    | 0.024    | 10.0     | 280.0    | 0.015       | -2111.0 | 0.031    |
| 2       | 400.0     | 800.0  | 18.0 | 29.5  | 16.0  | 448.0    | 0.032    | 6.0      | 372.0    | 0.004       | 0.0     | 0.043    |
| 3       | 400.0     | 800.0  | 18.0 | 29.9  | 16.0  | 448.0    | 0.032    | 6.0      | 372.0    | 0.005       | -751.0  | 0.035    |
| 4       | 400.0     | 800.0  | 18.0 | 36.2  | 16.0  | 436.0    | 0.032    | 6.0      | 326.0    | 0.010       | -455.0  | 0.046    |
| 5       | 400.0     | 1000.0 | 18.0 | 34.3  | 16.0  | 436.0    | 0.032    | 6.0      | 326.0    | 0.005       | -431.0  | 0.033    |
| 6       | 400.0     | 700.0  | 18.0 | 36.7  | 16.0  | 482.0    | 0.032    | 6.0      | 326.0    | 0.004       | -807.0  | 0.025    |
| 7       | 400.0     | 800.0  | 20.0 | 30.9  | 16.0  | 436.0    | 0.032    | 10.0     | 310.0    | 0.004       | 0.0     | 0.034    |
| 8       | 400.0     | 800.0  | 20.0 | 33.1  | 16.0  | 436.0    | 0.032    | 10.0     | 310.0    | 0.008       | 0.0     | 0.051    |
| 9       | 307.0     | 1910.0 | 36.0 | 38.8  | 12.0  | 240.0    | 0.018    | 6.0      | 240.0    | 0.006       | -145.0  | 0.042    |
| 10      | 152.0     | 1140.0 | 10.2 | 34.5  | 12.7  | 448.0    | 0.056    | 3.7      | 620.0    | 0.015       | -151.0  | 0.236    |
| 11      | 1520.0    | 9140.0 | 58.7 | 35.8  | 43.0  | 475.0    | 0.020    | 15.9     | 493.0    | 0.006       | -4450.0 | 0.081    |
| 12      | 250.0     | 750.0  | 9.9  | 24.1  | 7.0   | 446.0    | 0.020    | 3.1      | 441.0    | 0.014       | -120.0  | 0.094    |
| 13      | 250.0     | 1500.0 | 9.7  | 25.4  | 7.0   | 446.0    | 0.020    | 2.7      | 476.0    | 0.007       | -120.0  | 0.088    |
| 14      | 600.0     | 1800.0 | 30.2 | 31.4  | 22.2  | 448.0    | 0.019    | 9.5      | 431.0    | 0.005       | -400.0  | 0.067    |
| 15      | 609.6     | 2438.4 | 22.2 | 37.2  | 15.9  | 462.0    | 0.015    | 6.4      | 606.8    | 0.004       | -654.0  | 0.059    |
| 16      | 406.4     | 1047.8 | 10.4 | 34.7  | 12.7  | 458.5    | 0.014    | 4.5      | 691.5    | 0.001       | 0.0     | 0.045    |
| 17      | 406.4     | 1854.2 | 15.0 | 35.6  | 12.7  | 458.5    | 0.012    | 4.5      | 691.5    | 0.005       | 0.0     | 0.140    |

**Table B.2:** Censored high-fidelity data.

| Data ID | $D_{col}$ | $H_c$  | $c$  | $f_c$ | $d_s$ | $f_{ys}$ | $\rho_s$ | $d_{sv}$ | $f_{yv}$ | $\rho_{sv}$ | $N$     | $\delta$ |
|---------|-----------|--------|------|-------|-------|----------|----------|----------|----------|-------------|---------|----------|
| 1       | 500.0     | 2750.0 | 20.2 | 33.2  | 18.4  | 373.0    | 0.026    | 6.5      | 312.0    | 0.004       | -380.0  | 0.040    |
| 2       | 500.0     | 2730.0 | 20.2 | 40.0  | 18.4  | 305.0    | 0.026    | 8.0      | 389.0    | 0.013       | -26.4   | 0.095    |
| 3       | 250.0     | 1340.0 | 10.8 | 35.1  | 13.0  | 305.0    | 0.026    | 4.4      | 263.0    | 0.019       | -16.9   | 0.093    |
| 4       | 250.0     | 930.0  | 10.2 | 33.0  | 12.0  | 294.0    | 0.022    | 4.3      | 207.0    | 0.025       | -550.0  | 0.032    |
| 5       | 400.0     | 1600.0 | 16.0 | 26.0  | 16.0  | 308.0    | 0.024    | 6.0      | 308.0    | 0.008       | -680.0  | 0.034    |
| 6       | 600.0     | 1200.0 | 25.0 | 28.4  | 24.0  | 303.0    | 0.024    | 10.0     | 300.0    | 0.008       | -1920.0 | 0.025    |
| 7       | 600.0     | 1200.0 | 25.0 | 26.6  | 24.0  | 303.0    | 0.024    | 10.0     | 300.0    | 0.011       | -4300.0 | 0.023    |
| 8       | 600.0     | 1200.0 | 28.0 | 32.5  | 24.0  | 307.0    | 0.024    | 16.0     | 280.0    | 0.026       | -3385.0 | 0.041    |
| 9       | 400.0     | 800.0  | 21.0 | 31.2  | 16.0  | 448.0    | 0.032    | 12.0     | 332.0    | 0.010       | -784.0  | 0.046    |
| 10      | 400.0     | 800.0  | 18.0 | 33.7  | 24.0  | 424.0    | 0.032    | 6.0      | 326.0    | 0.005       | 0.0     | 0.060    |
| 11      | 400.0     | 800.0  | 18.0 | 34.8  | 16.0  | 436.0    | 0.019    | 6.0      | 326.0    | 0.005       | 0.0     | 0.047    |
| 12      | 400.0     | 1600.0 | 17.0 | 40.0  | 16.0  | 474.0    | 0.018    | 8.0      | 372.0    | 0.006       | -2652.0 | 0.029    |
| 13      | 400.0     | 1600.0 | 18.0 | 39.0  | 16.0  | 474.0    | 0.018    | 10.0     | 338.0    | 0.015       | -3620.0 | 0.040    |
| 14      | 400.0     | 800.0  | 20.0 | 38.0  | 16.0  | 423.0    | 0.032    | 10.0     | 300.0    | 0.014       | -907.0  | 0.042    |
| 15      | 400.0     | 800.0  | 18.0 | 37.0  | 16.0  | 475.0    | 0.032    | 6.0      | 340.0    | 0.005       | -1813.0 | 0.023    |
| 16      | 400.0     | 800.0  | 20.0 | 37.0  | 16.0  | 475.0    | 0.032    | 10.0     | 300.0    | 0.014       | -1813.0 | 0.037    |
| 17      | 307.0     | 900.0  | 36.0 | 35.9  | 12.0  | 240.0    | 0.018    | 6.0      | 240.0    | 0.006       | -145.0  | 0.026    |
| 18      | 1520.0    | 4570.0 | 60.3 | 34.3  | 43.0  | 475.0    | 0.020    | 19.1     | 435.0    | 0.015       | -4450.0 | 0.089    |
| 19      | 275.0     | 300.0  | 20.0 | 28.8  | 16.0  | 366.0    | 0.039    | 6.0      | 368.0    | 0.005       | 0.0     | 0.054    |
| 20      | 275.0     | 300.0  | 20.0 | 31.4  | 16.0  | 366.0    | 0.039    | 6.0      | 368.0    | 0.013       | -215.0  | 0.087    |
| 21      | 275.0     | 300.0  | 20.0 | 30.5  | 16.0  | 366.0    | 0.051    | 6.0      | 368.0    | 0.009       | -215.0  | 0.063    |
| 22      | 275.0     | 300.0  | 20.0 | 30.2  | 16.0  | 366.0    | 0.026    | 6.0      | 368.0    | 0.009       | -215.0  | 0.068    |
| 23      | 275.0     | 300.0  | 20.0 | 31.3  | 16.0  | 366.0    | 0.039    | 6.0      | 368.0    | 0.013       | -430.0  | 0.074    |
| 24      | 275.0     | 450.0  | 20.0 | 31.3  | 16.0  | 363.0    | 0.039    | 6.0      | 381.0    | 0.013       | 0.0     | 0.100    |
| 25      | 275.0     | 450.0  | 20.0 | 42.2  | 16.0  | 363.0    | 0.039    | 6.0      | 381.0    | 0.006       | -215.0  | 0.050    |
| 26      | 275.0     | 450.0  | 20.0 | 18.9  | 16.0  | 363.0    | 0.039    | 6.0      | 381.0    | 0.006       | -430.0  | 0.050    |
| 27      | 275.0     | 450.0  | 20.0 | 41.3  | 16.0  | 363.0    | 0.039    | 6.0      | 381.0    | 0.006       | -430.0  | 0.044    |
| 28      | 305.0     | 1372.0 | 14.5 | 29.0  | 9.5   | 448.0    | 0.020    | 4.0      | 434.0    | 0.009       | -200.0  | 0.076    |
| 29      | 610.0     | 914.5  | 15.9 | 30.0  | 12.7  | 462.0    | 0.005    | 6.4      | 361.0    | 0.003       | -503.0  | 0.029    |
| 30      | 610.0     | 3660.0 | 27.8 | 41.1  | 22.2  | 455.0    | 0.027    | 9.5      | 414.0    | 0.009       | -1780.0 | 0.062    |
| 31      | 457.0     | 910.0  | 24.8 | 38.3  | 15.9  | 427.5    | 0.024    | 9.5      | 430.2    | 0.011       | -1928.0 | 0.046    |
| 32      | 457.0     | 910.0  | 24.8 | 39.4  | 15.9  | 427.5    | 0.024    | 9.5      | 430.2    | 0.011       | -970.0  | 0.056    |
| 33      | 457.0     | 910.0  | 26.4 | 35.0  | 19.0  | 468.2    | 0.052    | 12.7     | 434.4    | 0.027       | -850.0  | 0.121    |
| 34      | 457.0     | 910.0  | 24.8 | 35.2  | 15.9  | 507.5    | 0.024    | 9.5      | 448.2    | 0.009       | -490.0  | 0.068    |
| 35      | 457.0     | 910.0  | 26.4 | 35.0  | 19.0  | 486.2    | 0.052    | 12.7     | 434.4    | 0.030       | -1914.0 | 0.112    |
| 36      | 457.0     | 3656.0 | 30.2 | 36.6  | 15.9  | 477.0    | 0.036    | 9.5      | 445.0    | 0.009       | -1780.0 | 0.075    |
| 37      | 609.6     | 4876.8 | 22.2 | 31.0  | 15.9  | 462.0    | 0.015    | 6.4      | 606.8    | 0.007       | -653.9  | 0.153    |
| 38      | 609.6     | 6096.0 | 22.2 | 31.0  | 15.9  | 462.0    | 0.015    | 6.4      | 606.8    | 0.007       | -653.9  | 0.180    |
| 39      | 609.6     | 2438.4 | 22.2 | 31.0  | 15.9  | 462.0    | 0.030    | 6.4      | 606.8    | 0.007       | -653.9  | 0.080    |
| 40      | 609.6     | 1828.8 | 28.6 | 34.5  | 19.0  | 441.3    | 0.027    | 6.4      | 606.8    | 0.009       | -911.8  | 0.090    |
| 41      | 609.6     | 4876.8 | 28.6 | 34.5  | 19.0  | 441.3    | 0.027    | 6.4      | 606.8    | 0.009       | -911.8  | 0.145    |
| 42      | 609.6     | 6096.0 | 28.6 | 34.5  | 19.0  | 441.3    | 0.027    | 6.4      | 606.8    | 0.009       | -911.8  | 0.171    |
| 43      | 250.0     | 1645.0 | 13.8 | 65.0  | 16.0  | 419.0    | 0.033    | 7.5      | 1000.0   | 0.015       | -1000.0 | 0.255    |
| 44      | 250.0     | 1645.0 | 15.6 | 65.0  | 16.0  | 419.0    | 0.033    | 11.3     | 420.0    | 0.035       | -1000.0 | 0.119    |
| 45      | 250.0     | 1645.0 | 14.0 | 90.0  | 16.0  | 419.0    | 0.033    | 8.0      | 580.0    | 0.018       | -1850.0 | 0.079    |
| 46      | 250.0     | 1645.0 | 15.6 | 90.0  | 16.0  | 419.0    | 0.033    | 11.3     | 420.0    | 0.017       | -1850.0 | 0.051    |
| 47      | 250.0     | 1645.0 | 13.8 | 90.0  | 16.0  | 419.0    | 0.033    | 7.5      | 1000.0   | 0.015       | -925.0  | 0.223    |
| 48      | 250.0     | 1645.0 | 13.8 | 90.0  | 16.0  | 419.0    | 0.033    | 11.3     | 420.0    | 0.034       | -1850.0 | 0.087    |
| 49      | 508.0     | 1524.0 | 21.3 | 56.3  | 16.0  | 455.0    | 0.010    | 4.5      | 455.0    | 0.001       | -1243.0 | 0.022    |
| 50      | 609.6     | 2438.4 | 22.2 | 37.2  | 15.9  | 462.0    | 0.015    | 6.4      | 606.8    | 0.007       | -1308.0 | 0.077    |
| 51      | 419.0     | 1968.5 | 55.6 | 60.6  | 22.2  | 429.5    | 0.021    | 9.5      | 413.7    | 0.018       | 0.0     | 0.182    |
| 52      | 457.2     | 2438.4 | 12.7 | 32.7  | 19.0  | 565.4    | 0.020    | 9.5      | 434.4    | 0.009       | -231.3  | 0.109    |
| 53      | 609.6     | 1219.2 | 18.6 | 29.8  | 15.9  | 454.0    | 0.014    | 4.9      | 200.0    | 0.001       | -18.8   | 0.021    |

**Table B.3:** Description and unit of notations

| Notation    | Description                                           | Unit |
|-------------|-------------------------------------------------------|------|
| $D_{col}$   | Dimension of column section                           | mm   |
| $H_c$       | Height of column                                      | mm   |
| $c$         | Cover thickness of concrete                           | mm   |
| $f_c$       | Compressive strength of concrete                      | MPa  |
| $d_s$       | Dimension of longitudinal reinforcement               | mm   |
| $f_{ys}$    | Yield strength of longitudinal reinforcement          | MPa  |
| $\rho_s$    | Longitudinal reinforcement ratio                      | —    |
| $d_{sv}$    | Dimension of volumetric transverse reinforcement      | mm   |
| $f_{yv}$    | Yield strength of volumetric transverse reinforcement | MPa  |
| $\rho_{sv}$ | Volumetric transverse reinforcement ratio             | —    |
| $N$         | Axial load                                            | kN   |
| $\delta$    | Deformation capacity                                  | —    |