



Delft University of Technology

## The Doctrine of Bayesian Statistics for Inverse Problems

Koers, G.J.

**DOI**

[10.4233/uuid:067c1954-6693-4bf3-9f1b-6fdcac31b59](https://doi.org/10.4233/uuid:067c1954-6693-4bf3-9f1b-6fdcac31b59)

**Publication date**

2023

**Document Version**

Final published version

**Citation (APA)**

Koers, G. J. (2023). *The Doctrine of Bayesian Statistics for Inverse Problems*. [Dissertation (TU Delft), Delft University of Technology]. <https://doi.org/10.4233/uuid:067c1954-6693-4bf3-9f1b-6fdcac31b59>

**Important note**

To cite this publication, please use the final published version (if applicable).  
Please check the document version above.

**Copyright**

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

**Takedown policy**

Please contact us and provide details if you believe this document breaches copyrights.  
We will remove access to the work immediately and investigate your claim.

# **The Doctrine of Bayesian Statistics for Inverse Problems**

Geerten Koers



# **The Doctrine of Bayesian Statistics for Inverse Problems**

Dissertation

for the purpose of obtaining the degree of doctor  
at Delft University of Technology  
by the authority of the Rector Magnificus  
Prof.dr.ir. T.H.J.J. van der Hagen;  
Chair of the Board for Doctorates  
to be defended publicly on  
Tuesday the 5<sup>th</sup> of December 2023 at 10:00 o'clock

by

Geerten Jens KOERS  
Master of Science in Mathematics, Universiteit Leiden,  
the Netherlands  
born in Groningen, the Netherlands



This dissertation has been approved by the promotor.

Composition of the doctoral committee:

- |                                 |                                     |
|---------------------------------|-------------------------------------|
| - Rector Magnificus,            | chairperson                         |
| - Prof. dr. A.W. van der Vaart, | Delft University of Technology,     |
|                                 | <i>promotor</i>                     |
| - Dr. B.T. Szabó,               | Università Bocconi, <i>promotor</i> |

Independent members:

- |                                  |                                |
|----------------------------------|--------------------------------|
| - Prof. dr. ir. G. Jongbloed,    | Delft University of Technology |
| - Prof. dr. R. Nickl,            | University of Cambridge        |
| - Prof. dr. A.J. Schmidt-Hieber, | University of Twente           |
| - Dr. H.N. Kekkonen,             | Delft University of Technology |

# Contents

|                                               |             |
|-----------------------------------------------|-------------|
| <b>Contents</b>                               | <b>v</b>    |
| <b>Samenvatting</b>                           | <b>ix</b>   |
| <b>Summary</b>                                | <b>xiii</b> |
| <b>1 Introduction</b>                         | <b>1</b>    |
| 1.1 Statistical models . . . . .              | 1           |
| 1.1.1 Frequentist statistics . . . . .        | 2           |
| 1.1.2 Bayesian statistics . . . . .           | 5           |
| 1.2 Asymptotic considerations . . . . .       | 7           |
| 1.2.1 Consistency . . . . .                   | 7           |
| 1.2.2 Contraction rates . . . . .             | 8           |
| 1.2.3 Minimax rates . . . . .                 | 8           |
| 1.2.4 Asymptotic normality . . . . .          | 9           |
| 1.2.5 Uncertainty quantification . . . . .    | 11          |
| 1.3 Statistical problems . . . . .            | 12          |
| 1.3.1 Infinite-dimensional problems . . . . . | 12          |
| 1.3.2 Priors on function spaces . . . . .     | 14          |
| 1.3.3 Regression . . . . .                    | 16          |
| 1.3.4 Inverse problems . . . . .              | 18          |
| 1.3.5 Non-linear inverse problems . . . . .   | 20          |
| 1.4 Computational considerations . . . . .    | 21          |
| 1.4.1 MCMC . . . . .                          | 22          |
| 1.4.2 Distributed methods . . . . .           | 22          |
| 1.5 Adaptation . . . . .                      | 24          |
| 1.5.1 Hierarchical Bayes . . . . .            | 25          |
| 1.5.2 Empirical Bayes . . . . .               | 25          |
| 1.6 Partial differential equations . . . . .  | 26          |
| 1.7 Overview . . . . .                        | 28          |

|          |                                                                                                    |           |
|----------|----------------------------------------------------------------------------------------------------|-----------|
| 1.8      | Notation . . . . .                                                                                 | 29        |
| <b>2</b> | <b>Linear methods for non-linear problems</b>                                                      | <b>31</b> |
| 2.1      | Introduction . . . . .                                                                             | 31        |
| 2.2      | General method . . . . .                                                                           | 36        |
| 2.2.1    | Posterior contraction rate . . . . .                                                               | 36        |
| 2.2.2    | Uncertainty quantification . . . . .                                                               | 37        |
| 2.3      | Bayesian analysis in the linear problem . . . . .                                                  | 39        |
| 2.3.1    | Smoothness scales . . . . .                                                                        | 39        |
| 2.3.2    | Eigenbasis . . . . .                                                                               | 40        |
| 2.3.3    | Sobolev spaces . . . . .                                                                           | 44        |
| 2.3.4    | Discrete observations . . . . .                                                                    | 45        |
| 2.4      | Schrödinger equation . . . . .                                                                     | 46        |
| 2.5      | Heat equation with absorption . . . . .                                                            | 50        |
| 2.6      | One-dimensional Darcy equation . . . . .                                                           | 52        |
| 2.7      | Exponentiated Volterra operator . . . . .                                                          | 55        |
| 2.8      | Numerical analysis . . . . .                                                                       | 56        |
| 2.9      | Discussion . . . . .                                                                               | 60        |
| 2.A      | Hilbert scales and Sobolev spaces . . . . .                                                        | 63        |
| 2.B      | Complements for Schrödinger equation . . . . .                                                     | 65        |
| 2.C      | Complements for heat equation with absorption . . . . .                                            | 71        |
| 2.D      | Smoothness norm contraction rates in the Gaussian white noise<br>linear inverse problems . . . . . | 74        |
| 2.D.1    | Smoothness scale . . . . .                                                                         | 75        |
| 2.D.2    | Inverse problem . . . . .                                                                          | 76        |
| 2.E      | Smoothness norm contraction rates in the discrete observation<br>linear inverse problems . . . . . | 83        |
| <b>3</b> | <b>Distributed Inverse Problems</b>                                                                | <b>89</b> |
| 3.1      | Introduction . . . . .                                                                             | 89        |
| 3.2      | Setup . . . . .                                                                                    | 91        |
| 3.2.1    | Distributed computing . . . . .                                                                    | 93        |
| 3.3      | Main results . . . . .                                                                             | 93        |
| 3.3.1    | Contraction rates . . . . .                                                                        | 93        |
| 3.3.2    | Uncertainty quantification . . . . .                                                               | 94        |
| 3.4      | Non-linear problems . . . . .                                                                      | 95        |
| 3.4.1    | Schrödinger equation . . . . .                                                                     | 96        |
| 3.5      | Simulations . . . . .                                                                              | 97        |
| 3.5.1    | Schrödinger equation . . . . .                                                                     | 98        |

|          |                                                                 |            |
|----------|-----------------------------------------------------------------|------------|
| 3.5.2    | Exponentiated Volterra operator . . . . .                       | 99         |
| 3.6      | Analysis score function . . . . .                               | 100        |
| 3.7      | Proofs . . . . .                                                | 104        |
| 3.7.1    | Proof of Theorem 3.3.1 . . . . .                                | 104        |
| 3.7.2    | Proof of Theorem 3.4.1 . . . . .                                | 106        |
| 3.7.3    | Proof of the lemmas . . . . .                                   | 107        |
| 3.8      | Technical lemmas . . . . .                                      | 116        |
| 3.9      | Local posterior variance . . . . .                              | 119        |
| 3.10     | Proof of Theorem 3.3.2 . . . . .                                | 128        |
| <b>4</b> | <b>Misspecified Bernstein-von Mises theorem</b>                 | <b>139</b> |
| 4.1      | Introduction . . . . .                                          | 139        |
| 4.2      | Description of the problem . . . . .                            | 141        |
| 4.3      | Misspecified Bernstein-von Mises Theorem . . . . .              | 143        |
| 4.3.1    | A general Bernstein-von Mises theorem . . . . .                 | 143        |
| 4.3.2    | Bernstein-von Mises in hierarchical models . . . . .            | 145        |
| 4.4      | Applications . . . . .                                          | 147        |
| 4.4.1    | Square integral operator . . . . .                              | 148        |
| 4.4.2    | Schrödinger equation . . . . .                                  | 150        |
| 4.4.3    | Parabolic PDE-constrained inverse problem . . . . .             | 151        |
| 4.5      | Numerical analysis . . . . .                                    | 153        |
| 4.5.1    | Square integral operator . . . . .                              | 153        |
| 4.5.2    | Schrödinger equation . . . . .                                  | 156        |
| 4.6      | Proof of the hierarchical Bernstein von-Mises theorem . . . . . | 157        |
| 4.6.1    | Proof of Lemma 4.3.2 . . . . .                                  | 157        |
| 4.6.2    | Proof of Theorem 4.3.3 . . . . .                                | 159        |
| 4.7      | Proofs for the applications . . . . .                           | 162        |
| 4.7.1    | Proof of Corollary 4.4.1 . . . . .                              | 163        |
| 4.7.2    | Proof of Corollary 4.4.2 . . . . .                              | 165        |
| 4.7.3    | Proof of Corollary 4.4.3 . . . . .                              | 166        |
| 4.8      | Technical lemmas . . . . .                                      | 168        |
| 4.9      | Proof of Proposition 4.3.1 . . . . .                            | 172        |
| <b>5</b> | <b>Conclusion</b>                                               | <b>175</b> |
|          | <b>Bibliography</b>                                             | <b>179</b> |
|          | <b>Acknowledgements</b>                                         | <b>191</b> |

|                         |            |
|-------------------------|------------|
| <b>Curriculum Vitae</b> | <b>193</b> |
|-------------------------|------------|

# Samenvatting

Stel je voor dat je jezelf moet navigeren door een compleet donkere grot. Een bekende manier om dit te bereiken, is echolocatie, wat werkt door het maken van geluid en het dat wordt gereflecteerd. Het probleem van het bepalen van de geometrische vorm van een ruimte uit een mengsel van reflecties van geluiden, is een voorbeeld van een invers probleem. In veel vakgebieden van de wetenschap, zijn er situaties waar het onmogelijk is om een parameter direct to meten, en in plaats daarvan, de enige beschikbare methode, is het meten van een ander object dat wordt beïnvloed door de parameter van belang. Deze statistische problemen kunnen uitdagend worden, wanneer het moeilijk, of zelfs onmogelijk is, om de observaties direct te inverteren naar de parameter waarin men is geïnteresseerd.

In veel gevallen heeft een statisticus een geloof over de echte waarde van de parameter voordat het experiment is gestart. Het Bayesiaanse paradigma is een aantrekkelijke methode om nieuwe informatie van observaties te combineren met dit voorafgaand geloof. De Bayesiaanse a-posteriori-verdeling geeft een degelijk mechanisme om het geloof over de waarheid te actualiseren.

In Hoofdstuk 2 beschouwen we het gebruik van deze Bayesiaanse technieken voor inverse problemen, als de prior niet noodzakelijk een uitdrukking van een a-priori geloof is over de ware parameter, maar eerder een mechanisme om de inferentie te regulariseren. In niet-lineaire inverse problemen, moet de analyse vaak gedaan worden per geval, omdat er geen algemene methode bestaat. In dit proefschrift presenteren we een methode om gebruik te maken van gevestigde Bayesiaanse methoden voor lineaire inverse problemen, met als doel het construeren van een goede hersteltechnieken voor niet-lineaire inverse problemen. Dit kan worden gedaan als het observationeel model geschreven kan worden als een lineair invers probleem voor een transformatie van de observatie, en deze transformatie uniek kan worden geïnverteerd naar de originele parameter van belang. We bewijzen dat de methode de parameter benadert met een optimale snelheid wanneer deze transformatie kan worden benaderd met dezelfde snelheid, en de afbeelding die de getransformeerde observatie inverteert

naar de parameter van belang Lipschitz is ten opzichte van geschikte normen. Bovendien, aangezien de lineaire methode adaptief is naar de ware gladheid van de functie, zal de a-posteriori verdeling voor het niet-lineaire probleem ook adaptief zijn. Dit betekent dat kennis over de gladheid van de ware functie niet noodzakelijk is voor het verkrijgen van minimax-snelheden. De methode wordt toegepast op verschillende voorbeelden voor inverse problemen met partiële differentiaalvergelijkingen, zoals de Schrödingervergelijking, de tijdsafhankelijke warmtevergelijking en de Volterra-operator. De resultaten worden geïllustreerd door een numerieke simulatiestudie voor de Schrödingervergelijking, gebruikmakend van synthetische data.

Om exacte statistische inferentie te doen in Gaussische modellen met  $N$  observaties, moeten Bayesiaanse methoden vaak een  $N$  bij  $N$  matrix inverteren. Het aantal berekeningen dat nodig is om zulke matrices te inverteren groeit met de derde macht van de dimensie van de matrix, dus deze methoden hebben last van een hoge berekeningskosten als er veel observaties beschikbaar zijn. Relatief weinig theoretische resultaten zijn bekend voor het verbeteren van de berekeningstijd voor deze statistische algoritmen. Een veel voorkomende manier om deze problemen aan te pakken is het toepassen van gedistribueerde methoden, dat de data opdeelt in enkele kleine verzamelingen, die individueel worden behandeld. Door het individueel behandelen van deze kleinere pakketjes op een lokale computer, kan men theoretisch een snelheidsverbetering krijgen die kwadratisch groeit in het aantal computers. In Hoofdstuk 4 bewijzen we dat deze methode dezelfde contractiesnelheid geeft als met toegang heeft tot discrete observaties. Ook wordt het aangetoond dat het een goede onzekerheidskwantificering biedt in dit geval, aangezien het eerlijke betrouwbaarheidsverzameling genereert. We vergelijken de methode met andere Bayesiaanse a-priori verdelingen in een numeriek onderzoek.

Het voordeel van niet-parametrische modellen is dat ze erg flexibel zijn en niet gevoelig voor misspecificatie, maar analytisch kunnen ze lastig zijn. In tegenstelling tot oneindig-dimensionale modellen, modellen met een eindig-dimensionale parameter zijn analytisch aantrekkelijk, maar ze lopen het risico om misgespecificeerd te zijn. Aangezien in veel gevallen het onmogelijk is om te bewijzen dat de ware data-genererende verdeling in het parametrische model zit, en in andere gevallen, modelleert men een probleem opzettelijk verkeerd om de resulterende berekeningen te versnellen. Daarom is het belangrijk om een goed begrip te hebben van het gedrag van de statistische methoden in misgespecificeerde situaties. In Hoofdstuk 4 behandelen we de Bayesiaanse a-posteriori verdelingen voor parametrische modellen die mogelijk misgespecificeerd zijn. We tonen aan dat in een Bayesiaanse situatie, een misgespecificeerd

---

model nog steeds leidt naar een asymptotische Gaussiaanse a-posteriori verdeling. Als het misgespecificeerde model Gaussiaans is, bewijzen we dat de a-posteriori verdeling krimpt rond de ware parameter, maar de *credible sets* van de a-posteriori-verdeling zijn mogelijk niet asymptotisch frequentistische betrouwbaarheidsverzamelingen. De resultaten worden ondersteund door een numeriek onderzoek met verscheidene voorbeelden.





# Summary

Imagine you need to navigate through a completely dark cave. A well-known way of achieving this is echolocation, which works by making sounds and listening to how they are reflected back. The problem of determining the geometric shape of a space from a mixture of reflections of emitted sounds, is an example of an inverse problem. In many fields of science, there are situations where it is impossible to measure a parameter of interest directly and instead, the only available method is to measure a different object that is affected by the parameter of interest. These statistical problems can become challenging when it is difficult, or even impossible, to invert the observations directly into the parameter that one is interested in.

In many cases, a statistician has a belief about the true value of the parameter before even starting the experiment. The Bayesian paradigm is an attractive method of combining the new information coming from observations with this prior belief. It gives a sound mechanism, namely the posterior distribution, to update the beliefs about the truth.

In Chapter 2 we consider the use of these Bayesian techniques for inverse problems, where the prior is not necessarily an expression of an a priori belief about the true parameter, but rather a mechanism to regularize the inference. In non-linear inverse problems, the analysis often needs to be done on a case-by-case basis, since a general method is lacking. In this thesis, we present a method of using established Bayesian methods for linear inverse problems, in order to construct proper recovery techniques for non-linear inverse problems. This can be done whenever the observational model can be written as a linear inverse problem for recovering a transformation of the observation, and this transformation can uniquely be inverted to the original parameter of interest. We prove that the method recovers the parameter at an optimal rate whenever this transformation can be recovered at the same rate, and the map inverting the transformed observation into the parameter of interest is Lipschitz with respect to appropriate norms. Furthermore, since the linear method is adaptive

to the true smoothness of the function, the posterior distribution for the non-linear problem will be adaptive as well. This means that knowledge about the smoothness of the true function is not necessary for achieving minimax rates. If the uncertainty quantification for the linear problem is good, then this method will also provide good uncertainty quantification for the non-linear problem. The method is applied to several examples regarding inverse problems involving partial differential equations, like the Schrödinger equation, the time-dependent heat equation, and the Volterra operator. The results are illustrated by a numerical simulation study in the Schrödinger equation, using synthetic data.

To perform exact statistical inference in Gaussian models with  $N$  observations, Bayesian methods often need to invert  $N$  by  $N$  matrices. The number of calculation that is needed to invert such matrices grows cubically in the dimension size of the matrix, so these methods suffer from a high computational cost when a large number of observations is available. Relatively few theoretical results are known regarding improving of the computational time of these statistical algorithms. A common way of tackling these issues is by applying distributed methods, which partition the data into several smaller sets that are handled individually. By handling these smaller packages individually on local computers in parallel, one can theoretically achieve a speed-up that grows cubically in the number of computers. In Chapter 3 we prove that this method gives the same rate of contraction when one has access to discrete observations. It is also shown that it provides appropriate uncertainty quantification in this case, since it generates honest confidence sets. We compare the method to other Bayesian prior distributions in a numerical study.

Non-parametric models have the advantage of being very flexible and are not prone to misspecification, but they can be analytically intractable. In contrast to infinite-dimensional models, models with a finite-dimensional parameter are analytically attractive, but are at risk of being misspecified. In many cases it is impossible to prove that the true data-generating distribution is included in a parametric model, and in other cases one might incorrectly specify the problem on purpose in order to speed up the resulting computations. Therefore it is important to have a good understanding of the behaviour of the statistical method in misspecified settings. In Chapter 4 we deal with the Bayesian posterior distributions for parametric models that are possibly misspecified. We show that in a Bayesian setting, a misspecified model still leads to an asymptotic Gaussian posterior under regularity conditions. We prove for a Gaussian surrogate likelihood in a hierarchical model that the posterior distribution still contracts around the true parameter, but credible sets from the posterior distribution might not

---

asymptotically be frequentist confidence sets. The results are supported by a numerical study with several examples.



# Chapter 1

## Introduction

The field of statistics concerns itself with the inference of an unknown truth from observations that carry partial information about this truth. The observations are assumed to come from some process with a random component, making it possible for two different data-generating mechanisms to produce exactly the same observed data. The inherent difficulty in a statistical problem is thus that there often is no guarantee which underlying data-generating mechanism is producing the observations. Statistical conclusions that are drawn from this stochastic data are thus fundamentally probabilistic statements about the truth, that offer no absolute certainty about the hidden mechanisms. From this, the two main points of statistics follow naturally: making an estimate of the underlying data-generating distribution, and subsequently quantifying the uncertainty of this inference. Within the field of statistics, there are two main paradigms on the nature of statistical models, namely the frequentist viewpoint and the Bayesian viewpoint. In this thesis, we focus on the Bayesian paradigm to derive methods for these statistical tasks for various problems, and the frequentist paradigm to analyse these methods.

### 1.1 Statistical models

In a mathematical framework, the set of all possible observations is denoted by the observational space  $\mathcal{X}$ . The collection of all possible ground truths is denoted by the parameter space  $\Theta$ , which is equipped with a  $\sigma$ -algebra  $\mathcal{B}$ . For each  $\theta \in \Theta$  there exists a corresponding probability distribution  $P_\theta$  on  $\mathcal{X}$ . The set of all probability distributions  $\mathcal{P} = \{P_\theta : \theta \in \Theta\}$  is called the *statistical model*. Now suppose that one is given some set  $\Omega$ , a  $\sigma$ -algebra  $\mathcal{F}$  on  $\Omega$  and some

probability distribution  $P$ . The *observation* or *data* is then a random variable on the probability space  $(\Omega, \mathcal{F}, P)$ , usually denoted by  $X : \Omega \rightarrow \mathcal{X}$ , that one observes. The random variable  $X$  thus induces some probability distribution on the space  $\mathcal{X}$ , and the goal of a statistician is to recover this induced probability distribution from having observed  $X$ . Observing a random variable from a mathematical point of view means that any function that is used for the statistical analysis should be measurable with respect to the  $\sigma$ -algebra generated by  $X$ . The model  $\mathcal{P}$  consists of all the descriptions of the random observations that are deemed possible, and the ultimate aim is to find the probability distribution within the model that describes the observed data most accurately. In many cases, the data  $X$  consists of a sequence of observations  $X_1, X_2, \dots, X_N$  all located in the same observational space  $\tilde{\mathcal{X}}$ , where each observation  $X_i \sim P_{\theta,i}$  is independent of the other observations, but does not necessarily follow the same distribution as the other observations. In this case the sample space is equal to  $\mathcal{X} = \tilde{\mathcal{X}}^N$  and one can write  $X = (X_1, \dots, X_N) \in \tilde{\mathcal{X}}^N$  and  $P_\theta = \otimes_{i=1}^N P_{\theta,i}$ . If the further assumption is made that the data is identically distributed, so that  $P_{\theta,i} = P_\theta$  for  $i = 1, \dots, N$ , one has independent and identically distributed random variables, commonly referred to as *i.i.d.* data. Such a set of i.i.d. random variables is a *sample* from the probability distribution  $P_\theta$ .

The prototypical statistical model is given by the observation of a variable  $X \sim P_\theta$  that satisfies

$$X = K(\theta) + \epsilon.$$

The mapping  $K : \Theta \rightarrow \mathcal{X}$  can be regarded as a transformation of the parameter of interest, into an object that can be observed by the statistician. Here  $\epsilon$  is the observational *noise*, a concept that is ubiquitous throughout the world of statistics. The noise can be assumed to be centred by modifying the map  $K$  with an additive term. By far the most common assumption on the noise is that it follows a Gaussian distribution, but other distributions have been studied as well.

### 1.1.1 Frequentist statistics

The frequentist paradigm assumes that there is a true probability distribution  $P_0$  that governs the observed data through  $X \sim P_0$ . In order to ensure that there is no ambiguity about which parameter is responsible for the observed data, it is common to assume additionally that the model is *identifiable*, meaning that the map  $\theta \mapsto P_\theta$  is injective. A priori, there are no restrictions for the statistician on the choice of model  $\mathcal{P}$  and thus it is a possibility that there does

not exist a  $\theta \in \Theta$  such that  $P_\theta = P_0$ . In the literature, often a distinction is made between *well-specified* models, where the true probability distribution  $P_0$  is an element of the model  $\mathcal{P}$ , and *misspecified* models, where  $P_0$  is not an element of  $\mathcal{P}$ . In well-specified models, the true probability distribution is an element of the model and often is denoted by  $P_0 = P_{\theta_0}$ , for some parameter  $\theta_0 \in \Theta$ .

### Point estimation

The problem of estimating the true parameter by a frequentist statistician can be formalized by the use of a *loss function*  $L : \Theta \times \Theta \rightarrow [0, \infty)$ . The goal is now to use the mathematical knowledge of the model  $(P_\theta)_{\theta \in \Theta}$ , to construct a mapping  $T : \mathcal{X} \rightarrow \Theta$  so that for all  $\theta_0 \in \Theta$ , with high probability, the quantity  $L(\theta_0, T(X))$  is small. The loss function quantifies the distance between the truth  $\theta_0$  and the point estimator  $T(X)$ , where a higher value of the loss function corresponds to a worse performance of  $T(X)$ . Thus  $L(\theta_0, T(X))$  measures how close  $T(X)$  is to  $\theta_0$  with respect to the loss function  $L$ . Common examples of loss functions are distance functions or (squared) norms defined on  $\Theta$ . Since from a frequentist perspective, there are no further assumptions made on the problem and the model, this mapping  $T$  has no restrictions on its form and can thus be modelled in any way the statistician deems appropriate for the problem at hand.

Since usually one can define many different estimators for the same statistical problem, it is necessary to compare the behaviour of different estimators. The *risk* (see (Le Cam and Yang 1990)) of an estimator  $T$  at  $\theta_0 \in \Theta$  with respect to a loss function  $L$  is defined as

$$R(\theta_0, T) := \mathbb{E}_{\theta_0} L(\theta_0, T(X)).$$

An estimator  $T$  is called *admissible* if there does not exist a measurable function  $T'$  such that  $R(\theta, T') \leq R(\theta, T)$  for all  $\theta \in \Theta$ , and for at least one parameter  $\theta_0 \in \Theta$  the inequality  $R(\theta_0, T') < R(\theta_0, T)$  holds. Such a function  $T$  can be regarded as an estimator that cannot be improved upon uniformly over the parameter space, and is usually favoured over other estimators.

Suppose there are two statisticians that have the same statistical problem with data  $X$  and model  $(P_\theta)_{\theta \in \Theta}$ . The first statistician will construct some map  $T_1$ , and the second will construct some map  $T_2$ . In principle, there is no reason for  $T_1$  and  $T_2$  to resemble each other in any way, as the statisticians might disagree on what is the best way to solve this problem. However, there are certain tools in frequentist statistics that are used very often. A common method for doing



frequentist inference is by the usage of the *maximum likelihood estimator* (MLE), which is defined as the function  $T_{\text{MLE}} : \mathcal{X} \rightarrow \Theta$  that satisfies

$$T_{\text{MLE}}(X) := \arg \max_{\theta \in \Theta} p_{\theta}(X). \quad (1.1)$$

The heuristic for this choice of mapping is clear, as presumably the estimator that has the best fit to the observed data is the parameter that maximizes the probability of observing this data. Since the derivation of the MLE is very intuitive, the principle has been proposed in the second half of the 18th century by many different authors in many different settings, and a case can be made for each of these scientists as the originator of the MLE (see (Stigler 1999) and (Stigler 2007)). A vast amount of research has been done on the statistical properties of this estimator since its introduction, and to this day it still receives widespread attention and is being used in many different applications.

### Confidence sets

In contrast to the problem of estimation, where there is a wide variety of possible methods and solutions, the problem of uncertainty quantification is often handled in the same way. In the frequentist paradigm, the idea is to construct a data-dependent set  $\mathcal{C} \in \mathcal{B}$ , with  $\mathcal{B}$  the  $\sigma$ -algebra of  $\Theta$ , that contains the true parameter  $\theta_0$  with a prescribed probability. Given a number  $\tau \in (0, 1)$ , the goal is thus to construct a function  $\mathcal{C} : \mathcal{X} \rightarrow \mathcal{B}$  such that

$$\inf_{\theta \in \Theta} P_{\theta}(\theta \in \mathcal{C}(X)) \geq 1 - \tau.$$

Such a set  $\mathcal{C}(X)$  is called a *confidence set*. One should interpret  $\mathcal{C}$  as a function such that the probability of the data  $X$  being generated such that the set  $\mathcal{C}(X)$  contains the underlying parameter  $\theta$ , is at least  $1 - \tau$ . As the true value of  $\theta_0$  is unknown, this needs to hold uniformly over the parameter space  $\Theta$ . In many cases, the definition of a confidence set does not uniquely determine  $\mathcal{C}(X)$ , which may enable the statistician to choose the set  $\mathcal{C}(X)$  such that it satisfies additional constraints. A common desirable property for a confidence set is to choose  $\mathcal{C}(X)$  such that it is of minimal size with respect to a measure  $\mu$  on  $\Theta$ , or has a certain geometrical shape, e.g. a ball with respect to a given distance function on  $\Theta$ . Such confidence sets of minimal size can in turn be regarded as being more informative compared to confidence sets  $\tilde{\mathcal{C}}(X)$  with a larger  $\mu$ -measure. Since the MLE by definition has the largest likelihood, statisticians often choose the confidence set equal to a ball centred at the MLE with a well-chosen radius. In practice, it can be difficult to explicitly calculate the value that this radius needs to have.

### 1.1.2 Bayesian statistics

In a Bayesian framework, the statistician assumes that the distribution  $P$  that generates the data is itself a random variable, and thus not fixed at a particular distribution. The probability distribution on the model  $\mathcal{P}$  is called the *prior distribution* and denoted by  $\Pi$ . Given a draw  $P \sim \Pi$  from the prior, the data itself follows the distribution  $X \mid P \sim P$ . Thus one can see the pair  $(P, X)$  as having a joint distribution on the space  $\mathcal{P} \times \mathcal{X}$ . Using this observation it is natural to condition the joint probability distribution on the data  $X$ , which gives the probability distribution  $P \mid X$ . This conditional probability distribution is called the *posterior distribution*, and it is the central point of inference for a Bayesian. The posterior distribution itself is thus a random probability measure, depending on the observed data  $X$ . The prior  $\Pi$  should be interpreted as the belief about the data-generating mechanism before any data has been observed; higher mass in the prior corresponds to a more likely explanation of the data, and similarly lower masses in the prior are for explanations that are deemed unlikely. The posterior in turn is the updated belief about the data-generating mechanism, given the observed data. It follows that a Bayesian statistician can use the posterior of a first experiment as the prior in a second trial of the same experiment.

On a more conceptual level, there exists a debate among frequentist statisticians and Bayesians over the correct viewpoint of statistics and the proper methodology. For a proponent's view of Bayesian statistics and responses to common criticisms of this paradigm, see (Berger 1980). One common argument comes from the common assumption that the order in which data is observed, does not matter. In many instances, it is reasonable to assume that the distribution of the data  $X_1, \dots, X_n$  is invariant under permutations of the individual observations. Such a sequence of random variables is said to be *exchangeable*. De Finetti's theorem says that a collection of random variables being exchangeable is equivalent to the existence of a random variable  $\theta$  such that  $X_1, \dots, X_n$  are i.i.d. given  $\theta$ , which is exactly the Bayesian setting (de Finetti 1931). Advocates of the Bayesian method use this equivalence as an argument in favour of this viewpoint.

The principled method of doing Bayesian inference is by computing the posterior distribution, as this distribution follows naturally from the conditional setup within the Bayesian framework. If a prior is an element of a family of priors  $(\Pi_\alpha)_{\alpha \in S}$  that depends continuously on  $\alpha \in S$  for some set  $S$ , then this parameter is called a *hyperparameter*.

A common assumption on the model  $\mathcal{P}$ , is that there exists a dominating

measure on  $\mathcal{X}$ , so that each probability distribution  $P_\theta \in \mathcal{P}$  has a density  $p_\theta$  relative to this dominating measure. The posterior probability for any measurable set  $A \subseteq \Theta$  can then be written as

$$\Pi(A \mid X) = \frac{\int_A p_\theta(X) d\Pi(\theta)}{\int_\Theta p_\theta(X) d\Pi(\theta)}. \quad (1.2)$$

This is *Bayes' theorem* and gets its name from (Bayes 1763). However, the question if Bayes indeed was the first to discover this theorem has no clear answer and is an interesting story on its own (Stigler 1999).

If the prior is part of a family of probability distributions, and the posterior distribution belongs to this same family, the family of priors is called *conjugate*. A well-known example of this is in an experiment where one tries to estimate the mean of a Gaussian, and this mean itself is equipped with a Gaussian prior. The posterior then belongs to the Gaussian family of probability distributions, making it conjugate.

Since the posterior distribution by nature is a probability distribution itself, it is perfectly suited to be used as a measure of uncertainty. The more concentrated the posterior is in a subset  $B$  of  $\Theta$ , the higher the belief should be that the parameter  $\theta$  of the data-generating probability measure  $P_\theta$  indeed belongs to  $B$ . If the statistician has managed to calculate the posterior  $\Pi(\cdot \mid X)$ , the second problem of quantifying the uncertainty about the parameter can be straightforwardly answered by using so-called *credible sets*. Given a  $\tau \in (0, 1)$ , a credible set  $B \in \mathcal{B}$  of posterior mass  $1 - \tau$  satisfies

$$\Pi(B \mid X) = 1 - \tau.$$

Since it might be impossible to construct a credible set with a credible level of exactly  $1 - \tau$ , a common alternative is to require the set  $B$  to be at least of mass  $1 - \tau$ , i.e. it satisfies  $\Pi(B \mid X) \geq 1 - \tau$ . A credible set can be seen as the Bayesian analogue of a frequentist confidence set, so naturally one can pose additional constraints on  $B$  whenever  $B$  is not uniquely determined. Similar to the confidence sets, often it is desirable to choose a credible set that is as small as possible with respect to a measure  $\mu$  on  $\Theta$ . In contrast to the frequentist procedure of finding a confidence set, if one has calculated the posterior distribution, credible sets are relatively straightforward to derive.

For a more comprehensive discussion of Bayesian statistics, see (Gelman et al. 2013) and (Ghosal and van der Vaart 2017).

## Frequentist Bayes

In a purely Bayesian setting, the posterior quantifies the full truth about the uncertainty of the statistical inference. Due to this viewpoint, it is impossible to make a precise statement about the quality of the method, as there is no ground truth to compare it to. In order to be able to derive quantitative results for the Bayesian methodology, one needs to embed the Bayesian method in a frequentist framework. Although the Bayesian procedure is derived from a setting that is incompatible with the frequentist assumption of the existence of a ground truth (except for trivial cases), one can still use the same methodology to derive the posterior distribution within the frequentist paradigm. One then assumes that there is a true data-generating distribution  $P_{\theta_0}$ , which is fixed. The prior  $\Pi$  does not necessarily reflect the prior belief about the true parameter  $\theta_0$ . Rather, it is a tool that can be freely adjusted to the liking of the statistician. In this case, the posterior is just another probability measure and its concentration around  $\theta_0$  can be quantified. A discussion of the advantages of working with both frequentist and Bayesian analysis and their symbiotic relationship can be found in (Bayarri and Berger 2004). In this thesis, we adapt the frequentist Bayes paradigm, as it allows us to mathematically analyse Bayesian methods.

## 1.2 Asymptotic considerations

Meaningful guarantees for statistical procedures where the sample size  $N$  is a small, fixed number can be difficult to obtain. It is for this reason that the most studied properties of frequentist and Bayesian techniques are asymptotic in nature; the number of observations  $N$  is assumed to grow arbitrarily large. Precise statements can be made about the inference from this viewpoint, as one expects that it should be possible to fully recover the unknown parameter  $\theta_0$  if one has access to an infinite amount of realisations of the random variables coming from  $P_{\theta_0}$ . For a general overview of asymptotic statistics, see (van der Vaart 1998).

### 1.2.1 Consistency

The most basic property of a posterior distribution is *consistency*, which means that for every neighbourhood of the parameter  $\theta_0$  in  $\Theta$ , in the limit of  $N$  going to infinity, the posterior puts all of its mass on this neighbourhood. The posterior

is weakly consistent at  $\theta_0$  if, for any neighbourhood  $B$  of  $\theta_0$ ,

$$\Pi(\theta \in B \mid X_1, \dots, X_n) \xrightarrow{P_{\theta_0}} 1.$$

Since the posterior is a random probability distribution, the left-hand-side of this display is a random variable in  $[0, 1]$ , and the convergence is in probability under the data-generating distribution  $P_{\theta_0}$ . If this convergence holds almost surely, the posterior is strongly consistent at  $\theta_0$ . As the ultimate goal of the statistical analysis is to derive the value of  $\theta_0$ , this property has to be satisfied for all  $\theta = \theta_0$  in  $\Theta$  to be of use in a statistical analysis. Since the neighbourhood  $B$  is fixed and does not depend on  $N$ , consistency does not provide any information about the speed at which the posterior distribution puts its mass on  $B$ . A refinement of the concept of consistency is needed in order to encapsulate this speed.

### 1.2.2 Contraction rates

A generalization of consistency is given by taking a sequence of neighbourhoods  $(B_N)_{N \in \mathbb{N}}$  that shrinks as the sample size  $N$  increases. For the sake of simplicity, we assume  $\Theta$  to be a metric space and we take the neighbourhoods to be balls centred at the truth  $\theta_0$  with radius  $\epsilon_N$ , where  $(\epsilon_N)_{N \in \mathbb{N}}$  is a sequence of non-increasing, nonnegative real numbers. If  $d : \Theta \times \Theta \rightarrow [0, \infty)$  is a semimetric on  $\Theta$ , then the posterior distribution  $\Pi(\cdot \mid X_1, \dots, X_N)$  has *contraction rate*  $\epsilon_N$  at  $\theta_0 \in \Theta$  with respect to the metric  $d$  if, for any real-valued sequence  $M_N \rightarrow \infty$ ,

$$\Pi(\theta : d(\theta, \theta_0) \leq M_N \epsilon_N \mid X_1, \dots, X_N) \xrightarrow{P_{\theta_0}} 1.$$

One should interpret this as that for large enough  $N$ , the posterior puts all of its mass around balls centred at  $\theta_0$  with radius  $\epsilon_N$  when the sample size is equal to  $N$ . Similarly to the consistency property, this contraction rate should hold for every  $\theta_0 \in \Theta$ . When the posterior contracts at a rate of  $\epsilon_N \rightarrow 0$  when  $N \rightarrow \infty$ , this immediately implies that the posterior is consistent as well.

### 1.2.3 Minimax rates

Usually, it is not possible to find an estimator  $T$  that has minimal risk for all  $\theta \in \Theta$ . Therefore, it is more natural to compare estimators uniformly over the set  $\Theta$ , by considering

$$\sup_{\theta_0 \in \Theta} R(\theta_0, T).$$

This quantity, known as the *maximum risk*, measures risk in the worst-case scenario. It is now natural to consider the lowest possible maximum risk by taking the infimum over all measurable functions  $T : \mathcal{X} \rightarrow \Theta$

$$\inf_T \sup_{\theta_0 \in \Theta} R(\theta_0, T).$$

When the data  $X = \{X_1, \dots, X_N\}$  is a sequence of  $N$  observations, the quantity on the left-hand-side is known as the *minimax rate*

$$r_N := \inf_T \sup_{\theta_0 \in \Theta} \mathbb{E}_{P_{\theta_0}^N} L(\theta_0, T(X_1, \dots, X_N)).$$

A sequence of estimator  $(T_N)_{N \in \mathbb{N}}$  that attains this minimax rate is (asymptotically) a *minimax estimator*.

### 1.2.4 Asymptotic normality

Within the fields of probability theory and statistics, perhaps the most common limiting distribution is that of a Gaussian distribution. The prototypical example in probability theory is given by the central limit theorem, stating that the average of an i.i.d. sample  $X_1, \dots, X_N$  of a distribution  $P$  with a finite second moment is, asymptotically, distributed as a normal distribution centred at  $\mathbb{E}_P X$  and with covariance  $\frac{1}{N} \text{Cov}(X)$ . In both frequentist and Bayesian statistics, there are analogous results concerning respectively the maximum likelihood estimator and the posterior.

#### Maximum likelihood estimator

Reconsider the maximum likelihood estimator as defined in Equation (1.1). For an i.i.d. sample  $X_1, \dots, X_N \sim P_{\theta_0}$ , by taking the logarithm of the likelihood and subtracting a constant term, the MLE can be equivalently defined as

$$T_{\text{MLE}} = \arg \max_{\theta \in \Theta} \frac{1}{N} \sum_{i=1}^N (\log p_{\theta}(X_i) - \log p_{\theta_0}(X_i)).$$

This can be seen as a particular *M-estimator*. It is known from the theory of *M*-estimation, that *M*-estimators are asymptotically normal around the maximizer of its expectation (van der Vaart 1998). The maximizer in this case is equal to

$$P_{\theta_0} \log p_{\theta}(X) - \log p_{\theta_0}(X).$$

The expression  $\theta \mapsto P_{\theta_0} \log p_{\theta}(X) - \log p_{\theta_0}(X)$  is known as the *Kullback-Leibler* divergence between the probability distributions  $P_{\theta}$  and  $P_{\theta_0}$ . It is always non-negative and can be seen as the degree of similarity between these distributions, although it should be noted that it is not symmetric and therefore also not a distance. For correctly specified models, the Kullback-Leibler divergence is maximized at  $\theta = \theta_0$ .

For smooth models, the  $M$ -estimator is the value at which the log-likelihood attains its maximum, so the derivatives of the log-likelihood vanish at that point. Thus the MLE can be written as the  $\theta$  that solves the equations

$$\sum_{i=1}^N \frac{\partial}{\partial \theta} \log p_{\theta}(X_i) = 0.$$

The term inside is called the *score-function* and is denoted by  $\dot{\ell}_{\theta}$ . The score function is used to define the *Fisher information* by  $I_{\theta} := P_{\theta} \dot{\ell}_{\theta} \dot{\ell}_{\theta}^{\top}$ . It has been observed by Fisher that the MLE asymptotically has a normal distribution, which is centred at  $\theta_0$  and has covariance  $\frac{1}{N} I_{\theta_0}^{-1}$  (Fisher 1925). A higher Fisher information corresponds to a smaller asymptotic variance of the MLE and consequently a better estimator. Heuristically, the score-function measures how distinct the model is in a neighbourhood around  $\theta_0$  as it is a derivative with respect to  $\theta$ . Hence the Fisher information, as an outer product of the score function, measures the difficulty of distinguishing models around the ground truth  $\theta_0$ .

### Bernstein-von Mises

Suppose that one has a *parametric model*, meaning that the parameter space  $\Theta$  is a subset of  $\mathbb{R}^d$  for some integer  $d \in \mathbb{N}$  and that the model  $(P_{\theta})_{\theta \in \Theta}$  in some sense continuously depends on the parameter  $\theta$ . Let the data  $X_1, \dots, X_N$  be a sample coming from some distribution  $P_{\theta_0}$  for a  $\theta_0 \in \Theta$ . A rescaled version of the posterior distribution converges, under mild regularity conditions, to a particular Gaussian distribution when the number of observations tends to infinity. We assume that the prior on  $\Theta$  is continuous and positive in a neighbourhood around  $\theta_0$ . When  $N$  goes to infinity, the distribution of  $\sqrt{N}(\theta - T_{\text{MLE}}) \mid X_1, \dots, X_N$  then asymptotically converges to the distribution  $N(0, I_{\theta_0}^{-1})$ . A result proving this Gaussian limit of a posterior distribution is known as a *Bernstein-von Mises phenomenon* (BvM). Since the asymptotic covariance matrix is equal to the inverse Fisher information  $I_{\theta_0}^{-1}$ , it purely depends on the true distribution  $P_{\theta_0}$  and the local behaviour of the model around  $\theta = \theta_0$ .

In particular, it should be noted that any influence by the prior vanishes as more and more data is observed. Although one might interpret this as meaning

that the choice of prior does not matter at all, one should note that for a finite number of observations, a careless chosen prior can significantly alter the quality of the posterior distribution.

Results proving behaviour similar to the Bernstein-von Mises phenomena date back to the work of Laplace (Laplace 1810). It gets its name from Bernstein (Bernstein 1917) and von Mises (von Mises 1931).

### 1.2.5 Uncertainty quantification

The natural way of expressing uncertainty from a Bayesian perspective is the use of credible sets. When one computes a posterior distribution from a frequentist point of view, one naturally asks the question if the computed credible sets are also (asymptotically) confidence sets. For a credible set  $\mathcal{C}(X)$  of credible level  $1 - \tau$ , the *frequentist coverage* is defined as

$$P_{\theta}(\theta \in \mathcal{C}(X)).$$

The *asymptotic coverage* of a sequence of credible sets  $(\mathcal{C}_N(X_N))_{N \in \mathbb{N}}$  is equal to

$$\liminf_{N \rightarrow \infty} \inf_{\theta \in \Theta} P_{\theta}(\theta \in \mathcal{C}_N(X_N)).$$

If the (asymptotic) coverage of a credible set is at least equal to the credible level,  $\mathcal{C}(X)$  is an *honest* confidence set. If the coverage converges to 1 when  $N$  goes to infinity, the credible set is also a *conservative* confidence set. As it turns out, determining if a credible set is also an honest confidence set is a rather difficult question for which the answer depends heavily on the particular statistical problem at hand. For parametric models, in view of the Bernstein-von Mises Theorem, the answer is affirmative when the model is sufficiently regular. Since in the finite dimensional case, the posterior distribution is approximately equal to a normal distribution centred at the maximum likelihood estimator, with a covariance equal to the covariance of the MLE, confidence sets constructed using the MLE are asymptotically equal to the credible sets centred at the posterior mean. In light of this result, the conclusions drawn by a frequentist statistician using the MLE, and a Bayesian statistician are the same when the number of observations goes to infinity.

The answer to this question in infinite-dimensional statistics is much harder, as there is no general Bernstein-von Mises theorem for non-parametric models. The coverage can either go to zero, any real number in  $(0, 1)$ , or one, depending on the parameter and the prior.



## 1.3 Statistical problems

### 1.3.1 Infinite-dimensional problems

Although parametric models admit to a relatively straightforward analysis, they are limited in their generality and are prone to misspecification if one does not have sufficient knowledge about the data-generating process. To overcome this problem, a significantly larger class of statistical models is given by *non-parametric models*, where the domain of the parameter is infinite-dimensional. It is commonly assumed that the parameter set  $\Theta$  is an infinite-dimensional vector space, and usually equals a function space defined on some domain.

If  $\theta$  is a function defined on a domain  $\mathcal{O} \subseteq \mathbb{R}^d$ , then often  $\theta$  is assumed to belong to  $\theta \in L^2(\mathcal{O})$ , the class of square-integrable functions:

$$L^2(\mathcal{O}) = \left\{ \theta : \mathcal{O} \rightarrow \mathbb{R} : \int_{\mathcal{O}} \theta(x)^2 dx < \infty \right\}.$$

If the space  $L^2(\mathcal{O})$  is equipped with the inner product

$$\langle \theta_1, \theta_2 \rangle_{L^2(\mathcal{O})} = \int_{\mathcal{O}} \theta_1(x) \theta_2(x) dx$$

for  $\theta_1, \theta_2 \in L^2(\mathcal{O})$ , then it becomes a Hilbert space with the squared norm given by  $\|\theta\|_{L^2(\mathcal{O})}^2 = \langle \theta, \theta \rangle_{L^2(\mathcal{O})}$ . This vector space has a countably infinite orthonormal basis, consisting of functions  $(\phi_i)_{i \in \mathbb{N}} \subset L^2(\mathcal{O})$  that satisfy the conditions  $\langle \phi_i, \phi_j \rangle_{L^2(\mathcal{O})} = \delta_{ij}$  and  $\|\phi_i\|_{L^2(\mathcal{O})} = 1$ . Using this basis, one can write the space  $L^2(\mathcal{O})$  as the sum of basis functions  $\phi_i$

$$L^2(\mathcal{O}) = \left\{ \sum_{i=1}^{\infty} \theta_i \phi_i : (\theta_i)_{i \in \mathbb{N}} \subset \mathbb{R}, \sum_{i=1}^{\infty} \theta_i^2 < \infty \right\}.$$

From this representation, it can be seen to be isomorphic to the Hilbert space  $\ell^2$ , consisting of real-valued sequences that have square-summable components. In particular, it should be noted that the coordinates of any element of  $\ell^2$  converge to zero when  $i$  goes to infinity.

A more general class of sets is given by the *Sobolev spaces*, also referred to as *SVD spaces*, that pose a restriction on the rate of decay of the coefficients  $\langle \theta, \phi_i \rangle$  for functions  $\theta \in L^2(\mathcal{O})$ . For any  $\beta \geq 0$ , the Sobolev space  $S^\beta(\mathcal{O})$  is defined as

$$S^\beta(\mathcal{O}) = \left\{ \sum_{i=1}^{\infty} \theta_i \phi_i : (\theta_i)_{i \in \mathbb{N}} \subset \mathbb{R}, \sum_{i=1}^{\infty} \theta_i^2 i^{\frac{2\beta}{d}} < \infty \right\}.$$

For the choice  $\beta = 0$ , one gets back  $S^0(\mathcal{O}) = L^2(\mathcal{O})$ .

Unlike parametric problems, in an infinite-dimensional setting, the choice of prior is paramount to good posterior behaviour. Consider the case where the data follows the law  $P_0$ , which can take a countably infinite number of distinct values on some space  $\mathcal{X}$ , and any probability distribution  $\hat{P}$  on  $\mathcal{X}$  not necessarily equal to  $P_0$ . Then one can construct a prior that puts positive mass in arbitrary small neighbourhoods of the  $P_0$ , but the posterior converges almost surely to  $\hat{P}$  (Freedman 1963).

For general settings, earlier results on posterior behaviour show that a prior  $\Pi$  is consistent on a set of  $\pi$ -measure one (Doob 1949). The theorem uses remarkably few assumptions on the parameter space and the distribution, making it applicable in a vast majority of cases. When a Bayesian is completely sure of their prior, this theorem is sufficient to get consistency of the posterior for almost all parameter values. Although from a measure-theoretic point of view, a set of prior probability one is as big as possible, the null set on which it might be inconsistent can be topologically very large. From a topological viewpoint, a set is considered small when it can be written as a countable union of sets that are nowhere dense. Such sets are said to be meagre. Then for a countably infinite observational space, all priors are consistent only on a meagre set, except for a meagre set on the space of priors itself (Freedman 1965). This result can be interpreted as showing that almost every prior is bad at almost every point in the parameter space. However, in practice, one does not arbitrarily choose a prior, and a careful choice of prior evades this problem entirely. Therefore this result only rules out the Bayesian methodology when it is applied without any consideration whatsoever to the problem. By choosing a good prior, one can achieve posterior consistency at every point of the parameter space. A theorem giving general conditions for posterior consistency at any point  $\theta_0 \in \Theta$  when observing i.i.d. data can be found in (Schwartz 1965). Furthermore, in (Freedman 1963) it was shown that posterior is asymptotically normal when the observations only take finitely many distinct values. More recent and general results on posterior contraction rates can be found in (Ghosal, Ghosh and van der Vaart 2000).

A general treatment of nonparametric models can be found in (Wasserman 2006) and (Giné and Nickl 2016). We discuss the general setup of nonparametric priors in the next section.

### 1.3.2 Priors on function spaces

Perhaps the most natural definition of a prior on a function space is by using a basis expansion. Let  $\mathcal{O} \subset \mathbb{R}^d$  be arbitrary domain, and take  $(\phi_i)_{i \in \mathbb{N}} \subset L^2(\mathcal{O})$  to be functions defined on this domain. Then the *random basis expansion* is given by the probability distribution induced by

$$\theta = \sum_{j=1}^{\infty} c_j \phi_j, \quad (1.3)$$

where the  $c_j \in \mathbb{R}$  are coefficients which themselves are equipped with a prior on  $\mathbb{R}$ . Whenever for all but a finite number of indices  $j$ , the term  $c_j \phi_j$  is identically zero almost surely, this prior reduces to a finite-dimensional parameter space. A typical choice is to take the  $\phi_j$  equal to an orthonormal basis of  $L^2(\mathcal{O})$ , where the specific type of basis is determined by the data-generating process. Common choices are polynomial, trigonometric, or wavelet bases. For an orthonormal basis, one usually chooses a prior for which the coefficients  $c_j$  are independent.

#### Gaussian priors

Possibly the most popular choice of prior is a Gaussian prior, as it can be analytically tractable and is easily modified to possess desirable properties. To start, a multivariate Gaussian distribution with mean  $\mu \in \mathbb{R}^n$  and covariance matrix  $\Sigma \in \mathbb{R}^{n \times n}$  is a random variable  $X \in \mathbb{R}^n$  that is distributed as  $\mu + LZ$ , where  $L$  is the matrix such that  $\Sigma = LL^\top$  from the Cholesky decomposition of  $\Sigma$ , and  $Z$  is a vector in  $\mathbb{R}^n$  of which the components are i.i.d. standard normal random variables. Such a vector is denoted to follow the distribution  $X \sim N(\mu, \Sigma)$  and this random variable satisfies  $\mathbb{E}X = \mu$  and  $\text{Cov}(X) = \Sigma$ .

It is possible to generalize the notion of a multivariate Gaussian distribution from a finite-dimensional space to an infinite-dimensional index set  $T$ . A *Gaussian process*  $W$  on an index set  $T$  is a set of random variables  $(W(t) : t \in T)$ , for which the finite-dimensional marginals have a Gaussian distribution. Specifically, for each set of indices  $t_1, \dots, t_n \in T$ , the  $\mathbb{R}^n$ -valued random variable  $(W(t_1), \dots, W(t_n))$  has a multivariate normal distribution. Equivalently, one can consider the map  $t \mapsto W(t)$ , which is called a *sample path*. A sample path is a realisation of a random variable in the space of functions  $\mathbb{R}^T = \{f : T \rightarrow \mathbb{R}\}$ . Since the marginals of a Gaussian process are multivariate normally distributed, they are fully determined by their *mean function*  $\mu : T \rightarrow \mathbb{R}$  and *covariance*

kernel  $C : T \times T \rightarrow \mathbb{R}$ , which are defined as

$$\begin{aligned}\mu(t) &= \mathbb{E}W(t) \\ C(t_1, t_2) &= \text{Cov}(W(t_1), W(t_2)).\end{aligned}$$

Every Gaussian process defined on  $T$  induces a covariance function  $C$  on  $T \times T$  which is positive-semidefinite, meaning that a matrix  $(A_{ij})_{1 \leq i, j \leq n}$ , with  $A_{ij} = C(t_i, t_j)$ , is positive-semidefinite for all choices of  $t_1, \dots, t_n \in T$ . Conversely by Kolmogorov's consistency theorem, given a positive-semidefinite function  $\tilde{C} : T \times T \rightarrow \mathbb{R}$ , there exists a centred Gaussian process with covariance kernel  $\tilde{C}$ . In many instances of a Gaussian process prior, the mean is taken to be identically zero. The characteristics of a Gaussian process are then fully determined by the choice of kernel function, which regulates properties like the smoothness of the sample paths.

Now suppose that the index set is equal to some spatial domain  $T = \mathcal{O} \subset \mathbb{R}^d$ . If the coefficients  $c_j$  in Equation (1.3) follow a normal distribution, then  $\theta$  is distributed as a Gaussian process. On the other hand, for a given Gaussian process  $W$ , there exists a basis of functions  $(\phi_j)_{j \in \mathbb{N}} \subset \mathbb{R}^{\mathcal{O}}$  and independently distributed coefficients  $c_j$  such that  $\sum_{j=1}^{\infty} c_j \phi_j$  has the same distribution as  $W$  (see (Ghosal and van der Vaart 2017)). Now assume that the covariance kernel is square integrable  $\int_{\mathcal{O}} \int_{\mathcal{O}} C(x, y)^2 dx dy < \infty$  and strictly positive. We can now define the linear operator  $V : L^2(\mathcal{O}) \rightarrow L^2(\mathcal{O})$  as

$$(Vg)(x) = \int_{\mathcal{O}} g(y) C(x, y) dy,$$

which is then compact and positive. The latter two properties imply that there exists an orthonormal basis of eigenfunctions  $(v_j)_{j \in \mathbb{N}}$  of  $V$  with corresponding eigenvalues  $\lambda_j > 0$ , which can be chosen so that  $\lambda_j \searrow 0$  when  $j$  goes to infinity. The kernel can then be written as

$$C(x, y) = \sum_{j=1}^{\infty} \lambda_j \phi_j(x) \phi_j(y).$$

The distribution of a Gaussian process with such a kernel is then equal to the law of the random variables  $W$  satisfying for any  $x \in \mathcal{O}$ ,

$$W(x) = \sum_{j=1}^{\infty} \sqrt{\lambda_j} Z_j \phi_j(x), \quad (1.4)$$

with  $Z_j \stackrel{i.i.d.}{\sim} N(0, 1)$ . This form of the Gaussian process is known as the *Karhunen-Loève expansion* (see (Ghosal and van der Vaart 2017)).

## Reproducing Kernel Hilbert Space

Any Gaussian process defines a Hilbert Space which describes intrinsic properties of the process, like the support and the shape of the distribution. We define the *first chaos* of a centred Gaussian process  $W$  as the closure in  $L^2(\text{Pr})$  of the set of all linear combinations of its finite-dimensional marginals, i.e.

$$\overline{\text{lin}}(W) := \text{cl} \left\{ \sum_{j=1}^n c_j W(t_j) : c_1, \dots, c_n \in \mathbb{R}, t_1, \dots, t_j \in T, n \in \mathbb{N} \right\}.$$

As the first chaos is defined as a closure of a subset in a Hilbert space,  $\overline{\text{lin}}(W)$  itself is a Hilbert space. Now the collection  $\mathbb{H}$  consisting of functions  $z_L : T \mapsto \mathbb{R}$  defined as  $t \mapsto \mathbb{E}(W_t L)$ , for  $L \in \overline{\text{lin}}(W)$  is known as the *reproducing kernel Hilbert space* (RKHS). On the RKHS one can define the inner product  $\langle z_{L_1}, z_{L_2} \rangle_{\mathbb{H}} := \mathbb{E}(L_1 L_2)$ . Thus the isometry  $\Psi : \mathbb{H} \rightarrow \overline{\text{lin}}(W)$  given by  $\Psi(z_L) = L$  shows that the RKHS  $\mathbb{H}$  is also naturally a Hilbert space. Note that for every  $t \in \mathbb{T}$ , the function  $s \mapsto C(t, s)$  is contained in the RKHS by taking  $L = W_t$ . Furthermore, the RKHS has the *reproducing property*: for any  $t \in T$  and  $h \in \mathbb{H}$ , it satisfies  $\langle C(t, \cdot), h \rangle_{\mathbb{H}} = h(t)$ .

For a series prior as defined in Equation (1.4), the RKHS can be described as the set of functions  $\sum_{j=1}^{\infty} w_j \phi_j$  such that  $\sum_{j=1}^{\infty} \frac{w_j^2}{\lambda_j} < \infty$ . The values of  $\lambda_j$  determine the ellipsoid-like shape of the Gaussian distribution, which has vanishing axes when  $j$  grows to infinity.

The RKHSs of many different Gaussian processes have been studied, and their relation to posterior contraction rates in general infinite-dimensional models (e.g. upper bounds in (van der Vaart and van Zanten 2008) and corresponding lower bounds in (Castillo 2008)).

### 1.3.3 Regression

In general, in a regression problem, the goal is to recover a parameter that is only corrupted with additive noise. One observes pairs of random variables  $(X_i, Y_i) \in \mathcal{O} \times \mathbb{R}$ , that satisfy the relationship

$$Y_i = \theta_0(X_i) + \epsilon_i, \quad i = 1, \dots, N, \quad (1.5)$$

for some unknown function  $\theta_0 \in L^2(\mathcal{O})$ . The function  $\theta_0$  is observed at the points  $X_1, \dots, X_N \in \mathcal{O}$  up to some centred noise  $\epsilon_i \stackrel{i.i.d.}{\sim} G$  coming from a known distribution  $G$  on  $\mathbb{R}$ . A distinction is made between the case where the

$X_i$  are chosen deterministically, which is referred to as *fixed design*, and the case where the  $X_i$  themselves are an i.i.d. sample from a probability distribution  $F$  on  $\mathcal{O}$ , which is a regression problem with *random design*. In the latter case, it follows that the  $(X_i, Y_i)$  are an i.i.d. sample of a probability distribution.

The mathematical analysis of the model from Equation (1.5) can be very cumbersome and often different observational models are used in practice. It is known (see (Reiß 2008)) that when the noise  $\epsilon_i$  comes from a Gaussian distribution and the design points are approximately equidistant, the model is asymptotically equivalent (in the Le Cam Distance, see (Giné and Nickl 2016)) to having observed the equation

$$Y = \theta_0 + \frac{1}{\sqrt{N}} \mathbb{W}. \quad (1.6)$$

This statistical model is known as the *signal-in-white-noise* model. In the above display,  $\mathbb{W}$  is centred Gaussian white noise indexed by  $L^2(\mathcal{O})$ . Therefore it is a Gaussian process such that for each  $f, g \in L^2(\mathcal{O})$ ,  $\mathbb{E}\mathbb{W}(f) = 0$  and  $\mathbb{E}\mathbb{W}(f)\mathbb{W}(g) = \langle f, g \rangle_{L^2(\mathcal{O})}$ . Observing Equation (1.6) is defined as seeing the Gaussian process  $(\langle Y, f \rangle)_{f \in L^2(\mathcal{O})}$ , such that for each  $f, g \in L^2(\mathcal{O})$ ,  $Y(f)$  has a marginal distribution

$$Y(f) \sim N(\langle \theta_0, f \rangle_{L^2(\mathcal{O})}, \|f\|_{L^2(\mathcal{O})}^2)$$

and covariance  $\text{Cov}(Y(f), Y(g)) = \langle f, g \rangle_{L^2(\mathcal{O})}$ .

As a mathematical model, working directly with the white noise model still can be difficult. However, if one picks any orthonormal basis  $(\phi_i)_{i \in \mathbb{N}}$  of  $L^2(\mathcal{O})$ , one can choose to observe the Gaussian process at the  $\phi_i$ . Let us define  $\theta_{0,i} := \langle \theta_0, \phi_i \rangle$  and  $Z_i := \mathbb{W}(\phi_i)$ , so that we observe

$$Y_i := \theta_{0,i} + \frac{1}{\sqrt{N}} Z_i, \quad i \in \mathbb{N}. \quad (1.7)$$

The  $Z_i$  are i.i.d. random variables from a standard normal distribution, as can be derived from the properties of the Gaussian white noise process and the orthogonality of the functions  $\phi_i$ . Since the  $Z_i$  are all independent random variables, this model is much more suited to mathematical analysis, as it reduces to recovering the mean of a countably infinite number of independent Gaussian random variables. Therefore perhaps the most used prior for signal-in-white-noise models are Gaussian processes, as they are conjugate to this model.

The minimax rate in the signal-in-white-noise model depends on the subspace to which  $\theta_0$  is assumed to belong. A common choice is to consider subspaces

that are defined by the smoothness parameter of  $\theta$ . A higher smoothness of  $\theta_0$  means that its tails are easier to estimate, and consequently the minimax rate is faster. We can define the subsets  $\Theta^\beta := S^\beta(\mathcal{O})$  that increase to  $L^2(\mathcal{O})$  for  $\beta \rightarrow 0$ . The most common loss function that is used in model (1.7) is the  $L^2$ -loss, defined as  $L(f, g) = \|f - g\|_{L^2(\mathcal{O})}$  for  $f, g \in L^2(\mathcal{O})$ . With respect to this loss function, when  $\theta_0 \in \Theta^\beta$ , the minimax rate of this problem is given by  $cN^{-\frac{\beta}{d+2\beta}}$  for some fixed constant  $c > 0$ , see (Tsybakov 2008).

In the last decades many statistical analyses have been performed on these equivalent models, see for example (Kimeldorf and Wahba 1970), (Pinsker 1980), (Ibragimov and Khas'minskii 1985), (Cox 1993) and (Belitser and Ghosal 2003).

### 1.3.4 Inverse problems

It is not always possible to have access to noisy measurements of the parameter of interest. Instead, often the parameter is transformed through a known operator, which is then corrupted by observational noise. This is a broad class of problems that may arise in many different scientific fields. A common application of inverse problems can be found in tomography, where boundary measurements of waves travelling through a medium are used to make inference on the interior of the material. Examples can be found in geophysics where measurements of the propagation speed of waves on the surface are used to make inference on the structural composition inside the Earth (Woodhouse and Dziewonski 1984), or medical imaging, to detect malignant breast tumours (Zou and Guo 2003).

Mathematically, in an inverse problem, one supposes that the observations consist of

$$Y_i = K(\theta_0)(X_i) + \epsilon_i, \quad i = 1, \dots, N,$$

for some  $X_1, \dots, X_n \in \mathcal{O}$ . Here  $K : \Theta \rightarrow L^2(\mathcal{O})$  is a known, injective map called the *forward map*, and  $\epsilon_i \sim G$  is i.i.d. centred noise from a known distribution  $G$  on  $\mathbb{R}$ . The naive approach would be to apply some version of an inverse operator  $K^{-1}$  to the observations. However, for a large number of inverse problems, this is not a feasible option, since the inverse operator might not be continuous, or even be undefined on the sample space due to the addition of the noise. A common example of such a phenomenon is when the operator  $K$  is smoothing, and the noise itself is rougher than  $K(\theta)$ .

When the true parameter  $\theta_0$  belongs to an SVD space  $S^\beta(\mathcal{O})$  for some  $\beta > 0$ , the maximum likelihood estimator may behave very poorly, as it does not take into account the smoothness of  $\theta_0$  and overfits to the data. A common solution to these problems is by making use of *regularization*, a process that

penalises functions that are deemed too complex. In the frequentist literature, a standard example of such regularization is ridge regression, which adds a squared norm to the objective function that is minimized. From a Bayesian point of view, this amounts to endowing the parameter space  $\Theta$  with a prior that places higher mass on functions with higher smoothness. A class of inverse problems that has been studied extensively is the collection of linear inverse problems, where the operator  $K : \Theta \rightarrow L^2(\mathcal{O})$  is assumed to be a continuous, linear map. Generally, there is a good understanding of the behaviour of these models, as the asymptotic contraction rate and posterior limit distribution for linear functionals have been derived in a number of different situations. We now make the additional assumption that  $\Theta \subseteq L^2(\mathcal{O})$  and that the operator  $K^\top K$  is compact. Under these conditions,  $K^\top K$  has an orthonormal eigenbasis  $(\phi_i)_{i \in \mathbb{N}}$  with corresponding eigenvalues  $(\kappa_i^2)_{i \in \mathbb{N}}$  that are strictly positive and are non-increasing. Then analogous to the model defined in Equation (1.7), one can define the statistically equivalent *singular value decomposition* model as

$$Y_i = \kappa_i \theta_{0,i} + \frac{1}{\sqrt{N}} Z_i, \quad i \in \mathbb{N}.$$

The signal-in-white-noise model is recovered for  $\kappa_i = 1$  for all  $i \in \mathbb{N}$ , which corresponds to choosing  $K$  equal to the identity operator, which is non-compact as an operator on  $L^2(\mathcal{O})$ . The difficulty of reconstructing  $\theta_0$  from these random variables depends on the rate of decay of  $\theta_{0,i}$  as  $i$  goes to infinity. For eigenvalues that decay polynomially as  $\kappa_i \asymp i^{-\frac{p}{d}}$  for some  $p > 0$ , the problem is a *mildly ill-posed* inverse problem. The larger  $p$  is, the more difficult it becomes to recover  $\theta_0$ , as the coefficients signal-to-noise ratio decreases at a faster rate. This increased difficulty is reflected in the minimax rate, which for a function  $\theta_0 \in S^\beta(\mathcal{O})$  is proportional to  $N^{-\frac{\beta}{d+2\beta+p}}$  (see (Cavalier 2008)) and is slower for larger  $p$ . For  $\beta$  going to infinity, the minimax rate asymptotically approaches  $N^{-\frac{1}{2}}$ , which is the typical contraction rate of a finite-dimensional problem. When the eigenvalues decay at an exponential rate, so that  $\kappa_i \asymp \exp(-pi)$  for some  $p > 0$ , the problem is referred to as a *severely ill-posed* problem. In both the polynomial and the exponential case, the value  $p$  with which the eigenvalues decay is called the *degree of ill-posedness*. A popular choice of prior for  $\theta$  is a product of Gaussian priors  $\Pi = \otimes_{i=1}^\infty N(\mu_i, \sigma_i^2)$ . Since the noise follows a normal distribution as well, this is a conjugate model with a posterior consisting of a countable product of univariate Gaussian random variables, which simplifies the analysis significantly.

The typical choice is to take  $\mu_i = 0$  and  $\sigma_i^2 = i^{-1-\frac{2\alpha}{d}}$ . The draws of this prior almost surely lie in  $S^{\alpha'}(\mathcal{O})$ , for all  $\alpha' < \alpha$ . It is known that this prior achieves a



contraction rate proportional to  $N^{-\frac{\alpha \wedge \beta}{d+2\beta+2p}}$ . This contraction rate is proportional to the optimal rate for the choice  $\alpha = \beta$ ; choosing an  $\alpha$  unequal to  $\beta$  leads to sub-optimal rates of recovery. If one's goal is to construct credible sets with suitable coverage,  $\alpha$  needs to be chosen slightly below  $\beta$ , which corresponds to a slightly undersmoothing of the prior. For a thorough analysis of this choice of prior see (Knapik, van der Vaart and van Zanten 2011) and (Knapik, Szabó et al. 2016).

More recently, for general linear inverse problems with additive Gaussian noise, a Bernstein-von Mises theorem has been proven for a class of linear functionals of the unknown parameter (Giordano and Kekkonen 2020).

For an overview of noiseless inverse problems in partial differential equations, we refer to (Isakov 2017). In the last two decades statistical methods for (linear) inverse problems have been studied extensively, see (Evans and Stark 2002). For a Bayesian analysis of such inverse problems, we refer to (Stuart 2010), (Florens and Simoni 2012), (Agapiou, Stuart and Zhang 2014), (Ray and Schmidt-Hieber 2016) and (Gugushvili, van der Vaart and Yan 2020).

### 1.3.5 Non-linear inverse problems

Much harder models are non-linear inverse problems, where the map  $K : \Theta \rightarrow L^2(\mathcal{O})$  no longer behaves linearly. Whenever it is assumed that the model is identifiable, it can be seen that  $K$  has to be injective. However, the inverse of  $K$  might not have an explicit expression or be analytically tractable. In recent years contraction rates for non-linear inverse problems have slowly been developed. This has been done by establishing bounds of the form  $\|\theta - \theta_0\| \leq C\|K(\theta) - K(\theta_0)\|^\eta$  for some  $C > 0$ , whenever  $\|K(\theta) - K(\theta_0)\|$  is small enough. These so-called stability estimates are especially prevalent in the field of PDEs, and a detailed treatment of these methods can be found in (Nickl 2023). We give one example of an inverse problem within the field of PDEs coming from the time-independent Schrödinger equation.

Consider a spatial  $C^\infty$  domain  $\mathcal{O} \subset \mathbb{R}^d$  for some integer  $d \geq 2$  and an unknown function  $\theta_0 \in C^s(\mathcal{O}) = \Theta$  that is strictly positive. Let  $g \in C^{s+2}(\overline{\mathcal{O}})$  be an arbitrary known function, and consider  $u_{\theta_0} \in L^2(\mathcal{O})$  to be the solution to

$$\begin{cases} \frac{\Delta}{2} u_{\theta_0} - \theta_0 u_{\theta_0} &= 0 & \text{on } \mathcal{O}, \\ u_{\theta_0} &= g & \text{on } \partial\mathcal{O}. \end{cases}$$

One now observes

$$Y_i = u_{\theta_0}(X_i) + \epsilon_i, \quad i = 1, \dots, N$$

with  $\epsilon_i \stackrel{i.i.d.}{\sim} N(0, 1)$  and  $X_i$  fixed, approximately equidistant points in  $\mathcal{O}$ . The mapping  $K : \Theta \rightarrow L^2(\mathcal{O})$  defined by  $\theta \mapsto u_\theta$  is non-linear, making inference much more difficult. In this example, it has been shown that for specific wavelet priors the posterior contracts at the minimax rate, and that a version of the Bernstein-von Mises theorem holds for smooth enough linear functionals of the parameter  $\theta$  (Nickl 2018). Similar consistency results have been proven in other elliptic PDEs (Giordano and Nickl 2020) and parabolic PDEs (Kekkonen 2022).

## 1.4 Computational considerations

One of the fundamental challenges of Bayesian statistics is the computation of the posterior distribution. Although in some cases, the prior is conjugate and the derivation and analysis of the posterior is straightforward, this cannot be guaranteed. Even when the likelihood admits a tractable expression that can be analysed, the posterior defined in Equation (1.2) might be difficult to compute, as it involves integrating over the parameter space. In infinite-dimensional models, or even high-dimensional models, the computation of this integral can become very difficult and time consuming.

Even in the case of conjugate priors, sampling from the posterior distribution may take up a lot of computational time. Consider  $N$  observations in the regression model with noise that follows a Gaussian distribution, for which the function of interest  $\theta_0$  has been equipped with a Gaussian prior. Although the posterior distribution has an explicit expression, the covariance matrix contains the inverse of an  $N \times N$  matrix. If  $N$  becomes too large, computation of this covariance matrix and subsequently taking draws from the posterior distribution becomes impractical due to limits on processing power or memory. Given the covariance matrix, in order to draw a sample from a multivariate Gaussian distribution with this covariance matrix, one needs to perform a Cholesky decomposition. In MATLAB, a popular computational tool in applied sciences, the function `mvnrnd`, which generates draws from a multivariate Gaussian distribution, usually chooses the QZ algorithm. This algorithm requires  $O(N^3)$  computations for an  $N \times N$  matrix (Moler and Stewart 1973). Thus drawing from the posterior distribution can quickly become very time-consuming whenever the sample size exceeds the computational limits of the computer.

### 1.4.1 MCMC

In order to approximate the posterior distribution, tools have been developed to circumvent the calculation of this integral entirely. An extremely popular algorithm for posterior sampling is Markov chain Monte Carlo (MCMC), which aims to generate draws from the posterior distribution without explicit computation of the posterior. The goal is to define a Markov chain that is easy to sample from and that has the desired posterior distribution as its stationary distribution. If one finds such a Markov chain, the procedure is to draw enough samples from the Markov chain, so that the samples come from the stationary distribution and so that far-enough spaced draws behave (nearly) independently. Finding a suitable MCMC method for a particular statistical model, and fine-tuning possible parameters like burn-in time and run length can be a non-trivial problem on its own. Popular MCMC schemes include Metropolis-Hastings algorithms, Gibbs samplers, and more recently Piecewise Deterministic Markov Processes. We refer to (Robert and Casella 1999), (Brooks et al. 2011), and (Fearnhead et al. 2016) for an introduction to MCMC methods and their applications.

### 1.4.2 Distributed methods

An intuitively clear method to reduce the computational costs of statistical methods is to decrease the size of the datasets that a computer handles at once. Because the statistician still wants to use all the data that has been collected, the dataset is partitioned into multiple smaller packages with fewer data points. Following that, on each of these datasets, a statistical algorithm is performed by a *local computer*, and those results are combined at a central computer for the final statistical product. This whole scheme of computation is known as *distributed computation*. For a general overview of these techniques within a Bayesian setting see (Szabó and van Zanten 2019). We will illustrate this technique by applying it to the discrete observational model coming from a regression problem.

### Random distribution

A straightforward way of dividing the data into smaller sets is by randomly assigning each data point to a partition. We assume for simplicity that  $N = mn$ , where  $m \in \mathbb{N}$  is the number of local computers and  $n \in \mathbb{N}$  is the size of each part of the partition. So for a data set  $\mathbb{D}_N = \{(X_i, Y_i) : i = 1, \dots, N\}$ , one can

define for  $k = 1, \dots, m$ , the sets

$$\mathbb{D}_n^{(k)} := \{(X_i^{(k)}, Y_i^{(k)}) : i = 1, \dots, n\}.$$

These sets satisfy that for all  $(X_{i_1}, Y_{i_1}), \dots, (X_{i_n}, Y_{i_n}) \in \mathbb{D}_N$ , with distinct indices  $i_1, \dots, i_n \in \{1, \dots, N\}$ , we have

$$\mathbb{P}((X_{i_1}, Y_{i_1}), \dots, (X_{i_n}, Y_{i_n}) \in \mathbb{D}_n^{(k)}) = \frac{1}{\binom{N}{n}}.$$

Asymptotically, this method will guarantee an approximately uniform spread of the data points in each part of the partition. Since each part of the partition has access to data distributed on the entire domain  $\mathcal{O}$ , inference from a local machine can be considered to be informative of the parameter of interest on the whole domain. The partition is uniformly random, so within a dataset  $\mathbb{D}_n^{(k)}$ , the data is distributed i.i.d. and follows the same distribution as  $\mathbb{D}_n$ .

On each local machine, the posterior is computed, and  $r \in \mathbb{N}$  samples  $\theta_1^{(k)}, \dots, \theta_r^{(k)} \in L^2(\mathcal{O})$  are drawn from the posterior, which are sent to the global machine. A simple method of combining these draws is averaging them to get one posterior draw

$$\theta_i = \frac{1}{m} \sum_{k=1}^m \theta_i^{(k)}, \quad i = 1, \dots, r.$$

### Spatial distribution

The second method of dividing the data is not by grouping them randomly, but by assigning a point  $(X_i, Y_i)$  to a set  $\mathbb{D}_N^{(k)}$  based on their spatial location  $X_i$ . The domain  $\mathcal{O}$  is partitioned into  $m$  sets  $\mathcal{O}^{(k)} \subset \mathcal{O}$ ,  $k = 1, \dots, m$ , such that  $\bigcup_{k=1}^m \mathcal{O}^{(k)} = \mathcal{O}$  and  $\mathcal{O}^{(k_1)} \cap \mathcal{O}^{(k_2)} = \emptyset$  if and only if  $k_1 \neq k_2$ . For the case that the spatial domain  $\mathcal{O}$  equals an interval  $[a, b] \subset \mathbb{R}$  on the real line, it is straightforward to define the intervals  $\mathcal{O}^{(k)} := [a + (k-1)(b-a)/m, a + k(b-a)/m)$  for  $k = 1, \dots, m$ . Then one defines the part  $\mathbb{D}_N^{(k)}$  as all the data points that lie in the  $k$ th interval of  $[a, b]$ :

$$\mathbb{D}_N^{(k)} := \{(X_i, Y_i) : X_i \in \mathcal{O}^{(k)}\}, \quad k = 1, \dots, m.$$

For spatial domains that have a more intricate geometrical shape, in particular domains with dimension  $d \geq 2$ , the construction of the partition requires more care.

Since the parts are dependent on the spatial component of the data, the observations are not identically distributed. Aggregating draws from a local posterior to a draw from the global posterior requires more care, as the data from  $\mathcal{O}^{(k)}$  might not be as informative about the value of  $\theta$  on  $\mathcal{O}^{(k')}$ ,  $k' \neq k$ . A common method is to stitch the draws from the local posterior with the help of a set of functions  $f^{(1)}, \dots, f^{(m)} \in L^2(\mathcal{O})$  such that  $\sum_{k=1}^m f^{(k)} \equiv 1$  on  $\mathcal{O}$ . Then the global posterior is defined as

$$\theta_i = \sum_{k=1}^m f^{(k)} \theta_i^{(k)}, \quad i = 1, \dots, m.$$

The straightforward method is to define the functions  $f^{(k)}$  to be only supported on  $\mathcal{O}^{(k)}$ , leading to the choice  $f^{(k)} = 1_{\mathcal{O}^{(k)}}$ . Since the draws  $\theta_i^{(k_1)}$  and  $\theta_i^{(k_2)}$  do not necessarily agree on the boundary of  $\mathcal{O}^{(k_1)}$  and  $\mathcal{O}^{(k_2)}$ , this might lead to discontinuous functions. If smoother draws are desired, a partition of unity on  $\mathcal{O}$  can be considered, where each function is continuous or differentiable up to an order chosen in advance.

## 1.5 Adaptation

Estimators in regression or inverse problems that achieve minimax rates often make use of the smoothness parameter  $\beta$  of the unknown parameter  $\theta_0 \in S^\beta(\mathcal{O})$ . Take for instance the prior  $\Pi_\alpha = \otimes_{i=1}^\infty N(0, i^{-1-\frac{2\alpha}{d}})$  that is common in signal-in-white-noise models and inverse problems. It is only for the choice  $\alpha = \beta$  that this prior achieves the minimax rate. In practice, however, one does not always have access to the value of  $\beta$ , so estimators using this are of little use. Estimators that make no use of the smoothness parameter  $\beta$ , but still achieve the minimax rate for each value of  $\beta > 0$  are said to be *adaptive*. Mathematically, this means that an estimator  $T$  adapts to the smoothness if for every  $\beta \in B$ , the risk of  $T$  over  $\Theta^\beta$  is bounded by

$$\sup_{\theta_0 \in \Theta^\beta} \mathbb{E}_{P_{\theta_0}^{(N)}} L(\theta_0, T(X)) \leq C(\beta) r_{N,\beta},$$

with  $r_{N,\beta}$  the minimax estimation rate for  $\Theta^\beta$ . Such estimators are highly desirable, as they are applicable over a wide range of problems at the same time, without the need to know the smoothness parameter. We will discuss two principled ways of constructing adaptive procedures in a Bayesian way.

### 1.5.1 Hierarchical Bayes

In many Bayesian models, the prior belongs to a family  $(\Pi_\alpha)_{\alpha \in (0, \infty)}$ , where choosing  $\alpha$  is crucial for attaining optimal rates. The most natural way of incorporating the unknown smoothness parameter into this Bayesian model is by endowing this hyperparameter itself with a prior. We take  $G$  to be some distribution on  $(0, \infty)$  and we define the *hierarchical model* as

$$\begin{aligned}\alpha &\sim G, \\ \theta &| \alpha \sim \Pi_\alpha, \\ X &| \alpha, \theta \sim P_\theta.\end{aligned}$$

The distribution  $G$  is referred to as a *hyperprior*. The hierarchical prior is equivalent to directly putting the prior  $\Pi = \int \Pi_\alpha dG(\alpha)$  on the parameter  $\theta$ , and thus is still a fully Bayesian method. The downside of putting an additional prior on the hyperparameter is that the posterior might now be intractable, even if for fixed  $\alpha$  the posterior has an explicit expression. Hierarchical Bayesian methods have received much attention (see e.g. (Hjort et al. 2010)). When the hyperprior  $G$  is chosen carefully, the model recovers the parameter  $\theta_0$  at the optimal rate.

### 1.5.2 Empirical Bayes

Instead of putting a prior on the hyperparameter, it is also possible to make a frequentist estimate of the parameter based on the observations. Consider an estimator  $\hat{\beta}_N$  of  $\beta$  based on the observations  $X_1, \dots, X_N$ . Then the empirical Bayes posterior is given by  $\theta | X_1, \dots, X_N$  with prior  $\theta \sim (\Pi_\beta)_{|\beta=\hat{\beta}_N}$ . Although this technique is not considered to be fully Bayesian, it offers computational advantages whenever the estimator  $\hat{\beta}_N$  is relatively easy to calculate and the posterior is tractable given a fixed  $\beta$ . There are many ways of choosing the function  $\hat{\beta}_N(X_1, \dots, X_N)$ , but one of the most common choices is an analogue of the MLE. Consider a model where  $\theta$  is drawn from the prior  $\Pi_\alpha$  and the data  $X_1, \dots, X_N$  given  $\theta$  is generated by  $P_\theta$  with density  $p_\theta$ . Then define the estimator  $\hat{\beta}_N$  as

$$\hat{\beta}_N := \arg \max_{\beta > 0} \int_{\Theta} P_\theta(X_1, \dots, X_N) d\Pi_\beta(\theta).$$

This estimator maximizes the marginal likelihood of the observations given that they are generated by the Bayesian model. It is common for technical reasons to restrict the maximization to a set  $B_N$  that slowly grows to the full parameter

space  $(0, \infty)$  as  $N$  goes to infinity. Empirical Bayes estimators have been found to be adaptive and achieve minimax rates both in direct problems for signal-in-white-noise models (Szabó, van der Vaart and van Zanten 2013) and inverse problems in white noise models (see (Szabó, van der Vaart and van Zanten 2015) and (Knapik, Szabó et al. 2016)). For more results on empirical Bayes methods, we refer to (Petrone et al. 2014), (Rousseau and Szabo 2017), and (Donnet et al. 2018).

## 1.6 Partial differential equations

A *partial differential equation* (PDE) is a system of equations for an unknown function, which involves its partial derivatives. If the unknown function is only a function of one argument, it is said to be an *ordinary differential equation* (ODE). The goal of the study of PDEs is to establish the existence and uniqueness of solutions of certain partial differential equations, explicit expressions for these solutions, and the behaviour and properties of the solutions.

For an integer  $\ell \in \mathbb{N}$ , let  $D^{(\ell)}$  denote all the partial derivative operators of order  $\ell$ , i.e. the collection of partial derivatives  $\frac{\partial^\ell f(x)}{\partial x_1^{\ell_1} \dots \partial x_d^{\ell_d}}$  for  $\ell_1, \dots, \ell_d \in \mathbb{N}$  such that  $\sum_{j=1}^d \ell_j = \ell$ . A general version of a PDE for a function  $f$  on an open domain  $\mathcal{O} \subset \mathbb{R}^d$  for  $d \in \mathbb{N}$  is an equation of the form

$$G(x, f(x), D^{(1)}f(x), \dots, D^{(m)}f(x)) = 0, \quad x \in \mathcal{O},$$

for a function  $G : \mathcal{O} \times \mathbb{R} \times \dots \times \mathbb{R}^{n^m} \rightarrow \mathbb{R}$ . The highest integer  $k$  such that  $G$  depends on  $D^{(k)}$  is called the *order* of the PDE. Usually, solutions to a PDE are required to satisfy certain conditions on subsets  $\Gamma \subset \mathcal{O}$  of the domain on which they are defined. Such restrictions commonly describe boundary conditions or initial values and are often important to impose the uniqueness of the solution.

The PDE is said to be *linear* if  $G$  is a linear combination of partial derivative operators, i.e.  $G = \sum_{\ell=1}^m h_\ell(x) D^{(\ell)}f(x) - h(x)$ , for known functions  $h, h_\ell : \mathcal{O} \rightarrow \mathbb{R}$ . An operator  $G$  is *elliptic* if one can write as

$$Gf = - \sum_{i,j=1}^d \frac{\partial}{\partial x_j} \left( a_{ij}(x) \frac{\partial}{\partial x_i} f \right) + \sum_{j=1}^d b_j(x) \frac{\partial}{\partial x_j} f + c(x)f,$$

for functions  $a_{ij}, b_j, c : \mathcal{O} \rightarrow \mathbb{R}$  and there exists a constant  $c > 0$  such that the

functions  $a_{ij}$  are symmetric, i.e.  $a_{ij} = a_{ji}$  and satisfy

$$\sum_{i,j=1}^d a_{ij}(x) y_i y_j \geq c \|y\|^2$$

for all  $y \in \mathbb{R}^d$  and almost every  $x \in \mathcal{O}$ .

A common assumption in the field of PDE is to consider functions that are not classically differentiable, but differentiable in a less strict sense. Let  $C_c^\infty(\mathcal{O})$  denote the space of functions  $v : \mathcal{O} \rightarrow \mathbb{R}$  that are infinitely differentiable, and supported on a compact subset of  $\mathcal{O}$ . For a given  $\alpha \in \mathbb{N}^k$ , a function  $f : \mathcal{O} \rightarrow \mathbb{R}$  is said to have *weak derivative*  $D^\alpha f$  if for all functions  $\psi \in C_c^\infty(\mathcal{O})$  it satisfies

$$\int_{\mathcal{O}} f D^\alpha \psi dx = (-1)^{|\alpha|} \int_{\mathcal{O}} (D^\alpha f) \psi dx.$$

If  $f$  has a classical derivative  $D^\alpha$ , one sees that the weak derivative as defined here equals the classical derivative due to a repeated application of integration by parts (the functions  $\psi$  are chosen to be compactly supported in order to remove any boundary terms).

For  $k \in \mathbb{N}$ , the space of functions

$$H^k(\mathcal{O}) := \{f : \mathcal{O} \rightarrow \mathbb{R} : D^\alpha f \in L^2(\mathcal{O}), |\alpha| \leq k\}$$

is called a *Sobolev space*. It is a Banach space when it is equipped with the squared norm

$$\|f\|_{H^k(\mathcal{O})}^2 := \sum_{|\alpha| \leq k} \int_{\mathcal{O}} (D^\alpha f)^2 dx.$$

It consists of functions that have  $|\alpha|$  weak derivatives, all of which are elements of  $L^2(\mathcal{O})$ . The closure of  $C_c^\infty(\mathcal{O})$  in  $H^k$  with respect to the  $\|\cdot\|_{H^k(\mathcal{O})}$ -norm is denoted by  $H_0^k(\mathcal{O})$ . One can see this as the elements of  $H^k(\mathcal{O})$  for which all derivatives up to  $k - 1$  vanish on the boundary of  $\mathcal{O}$ . One can construct spaces of non-integer smoothness  $\beta \notin \mathbb{N}$  in various (equivalent) ways, c.f. (Evans 2010; Triebel 2008; Di Nezza, Palatucci and Valdinoci 2012). For certain orthonormal bases, there exists a continuous embedding of the SVD spaces into Sobolev spaces (Taylor 2011). For a comprehensive discussion of the properties of Sobolev spaces, see (Maz'ya 2011).

In many cases, it is impossible to derive the existence of solutions of PDEs that are differentiable in the classical sense. This necessitates the need for a



weaker notion of solving a PDE. To ease the notation of this example, we take  $G$  to be elliptic. Then the function  $f \in H_0^1(\mathcal{O})$  is a *weak solution* of the PDE

$$\begin{cases} -\sum_{i,j=1}^d \frac{\partial}{\partial x_j} \left( a_{ij}(x) \frac{\partial}{\partial x_i} f \right) + \sum_{j=1}^d b_j(x) \frac{\partial}{\partial x_j} f + c(x)f = h, & x \in \mathcal{O}, \\ f = 0, & x \in \partial\mathcal{O}, \end{cases}$$

if for all  $g \in H_0^1(\mathcal{O})$ , the function  $f$  satisfies

$$\int_{\mathcal{O}} \sum_{i,j=1}^d a_{ij} \frac{\partial}{\partial x_i} f \frac{\partial}{\partial x_j} g + \sum_{j=1}^d b_j \frac{\partial}{\partial x_j} f g + c f g dx = \int_{\mathcal{O}} g h dx.$$

The idea behind this is that an appropriately differentiable solution  $f$  satisfies this relation for all compactly supported functions  $g \in C^\infty(\mathcal{O})$ . Proving the existence of weak solutions to PDEs turns out to be much more manageable compared to proving such results for classically differentiable functions.

For an introductory text to the theory of PDEs, we refer to (Evans 2010), and a more advanced treatment of the theory can be found in (Taylor 2011), (Lions and Magenes 1972a) and (Lions and Magenes 1972b)

## 1.7 Overview

In Chapter 2 of this thesis, we propose a method to use existing Bayesian methods for linear problems for certain non-linear inverse problems. Using established techniques from the linear inverse problem literature, we achieve minimax rates for a general class of non-linear problems. The results are illustrated by a simulation study.

In Chapter 3, we extend existing methods for distributed computation in linear regression problems to the linear inverse problem setting. It is combined with the technique from Chapter 2, to obtain minimax rates in non-linear inverse problems using a distributed computational method. A simulation study of this method is presented and compared to earlier results for a non-linear inverse problem.

In Chapter 4, we prove an additional version of the Bernstein-von Mises theorem in a misspecified setting. We give multiple examples of a common Gaussian misspecification for several hierarchical, non-linear operators and numerically determine the asymptotic behaviour of each model.

## 1.8 Notation

For two sequences  $a_N, b_N$  we write  $a_N \lesssim b_N$  if there exists a constant  $C > 0$  such that  $a_N/b_N \leq C$ , for every  $N$ . We write  $a_N \asymp b_N$  if  $a_N \lesssim b_N$  and  $b_N \lesssim a_N$  hold simultaneously. The letters  $c$  and  $C$  denote constants not depending on  $N$ ; their values might change from line to line. The Euclidean norm of the finite-dimensional vector  $v$  is denoted by  $\|v\|$ . The same notation  $\|V\|$  is used for the operator norm of a  $d \times d$  matrix  $V$ .



# Chapter 2

## Linear methods for non-linear problems

### 2.1 Introduction

Consider the recovery of a function  $f : \mathcal{O} \rightarrow \mathbb{R}$  on a known domain  $\mathcal{O} \subset \mathbb{R}^d$  from a noisy observation of the function  $u_f$  satisfying a partial differential equation of the form

$$\begin{cases} \mathcal{L}u_f = c(f, u_f), & \text{on } \mathcal{O}, \\ u_f = g, & \text{on } \Gamma \subseteq \partial\mathcal{O}. \end{cases} \quad (2.1)$$

Here  $\mathcal{L}$  is a differential operator and  $g : \Gamma \rightarrow \mathbb{R}$  is a known function on a subset  $\Gamma$  of the boundary  $\partial\mathcal{O}$  of  $\mathcal{O}$ . The map  $c$  transforms the pair of functions  $(f, u_f)$  into a function on  $\mathcal{O}$ . In most of our examples, this map takes the form of a pointwise transformation, such as  $c(f, u_f)(x) = \bar{c}(f(x), u_f(x))$ , for a map  $\bar{c} : \mathbb{R}^2 \rightarrow \mathbb{R}$ , but it may also act on  $(f, u_f)$  as functions belonging to given function spaces (and hence also involve derivatives of  $f$  or  $u_f$ ). The idea of Equation (2.1) is to isolate possible non-linearities in the PDE in the term  $c(f, u_f)$ , and to choose the linear transformation  $\mathcal{L}u_f$  rich enough to recover the function  $f$ . We make the latter precise by assuming that there is a (sufficiently regular) solution map  $e$  such that

$$f = e(\mathcal{L}u_f). \quad (2.2)$$

In several examples, the map  $e$  takes a simple, pointwise form, which makes our general method straightforward to implement. More generally, the map  $e$  may act on  $\mathcal{L}u_f$  as an element of a function space, and for practical implementation

a numerical implementation of the recovery of  $f$  from  $\mathcal{L}u_f$  suffices. The general idea of the chapter is to solve the linear inverse problem of recovering  $\mathcal{L}u_f$  from the observations by the Bayesian method, and next recover  $f$  by an implementation of Equation (2.2). We prove that the method can attain optimal results, and confirm the feasibility of the approach also by numerical experiments.

The basic observational model is the white noise model, which consists of observing the Gaussian process  $(Y_N(h) : h \in L^2)$  indexed by  $L^2 = L^2(\mathcal{O})$ . This process is given by

$$Y_N(h) = \langle h, u_f \rangle_{L^2} + \frac{1}{\sqrt{N}} \dot{\mathbb{W}}(h),$$

for  $\dot{\mathbb{W}} = (\dot{\mathbb{W}}(h) : h \in L^2)$  the mean zero Gaussian process with covariance function  $\text{Cov}(\dot{\mathbb{W}}(h_1), \dot{\mathbb{W}}(h_2)) = \langle h_1, h_2 \rangle_{L^2}$ . Informally this means observing the function  $u_f$  plus white noise, which we can write in the form

$$Y_N = u_f + \frac{1}{\sqrt{N}} \dot{\mathbb{W}}. \quad (2.3)$$

Next to the preceding continuous interpretation of Equation (2.3), we also consider the practically relevant discrete observational setting, where  $Y_N$  is observed at  $N$  given design points  $x_{N,1}, \dots, x_{N,N} \in \mathcal{O}$ . In this case the observation  $Y_N$  is a vector in  $\mathbb{R}^N$  with coordinates  $Y_N(i) = u_f(x_{N,i}) + W_i$ , where  $W_1, \dots, W_N$  are i.i.d. standard normal variables. In our general discussion, we shall use Equation (2.3) to refer to both models.

Examples of the structure in Equation (2.1) include the time-independent Schrödinger equation, the heat equation with absorption term, considered in (Kekkonen 2022), with  $\mathcal{L} = -\Delta_x/2 + \partial_t$ , the one-dimensional Darcy problem with  $\mathcal{L}$  the inverse Volterra operator, and the one-dimensional Darcy problem with  $\mathcal{L}$  equal to  $\frac{d^2}{dx^2}$ . Details for these examples and additional examples are given in Sections 2.4–2.7. It is not clear that the approach works in general, but the list of examples is encouraging.

A non-homogeneous boundary condition ( $u_f = g$  on  $\Gamma$  with  $g$  nonzero) still renders the problem of recovering  $\mathcal{L}u_f$  non-linear. To remove also this non-linearity, define  $K$  as the solution operator  $u \mapsto Ku$  of the homogeneous partial differential equation

$$\begin{cases} \mathcal{L}Ku = u & \text{on } \mathcal{O}, \\ Ku = 0 & \text{on } \Gamma \subseteq \partial\mathcal{O}. \end{cases} \quad (2.4)$$

Thus the linear operator  $K$  is the inverse of  $\mathcal{L}$  under the Dirichlet boundary condition. For our purpose, it suffices that Equation (2.4) is valid and  $K$  is defined on the set of functions  $\{\mathcal{L}u : u = 0 \text{ on } \Gamma\}$ . Furthermore, for the given boundary function  $g$  in (2.1), let  $\tilde{g} : \mathcal{O} \rightarrow \mathbb{R}$  be the function that satisfies

$$\begin{cases} \mathcal{L}\tilde{g} = 0 & \text{on } \mathcal{O}, \\ \tilde{g} = g & \text{on } \Gamma \subseteq \partial\mathcal{O}. \end{cases} \quad (2.5)$$

Solutions to the Equation (2.4)–(2.5) exist and are unique under mild regularity conditions on the operator  $\mathcal{L}$ , the domain  $\mathcal{O}$ , the boundary set  $\Gamma$  and the function  $g$ . In particular, suppose that the solution to the problem to Equation (2.5) with a homogeneous Dirichlet boundary condition is unique, i.e. the function  $v = 0$  is the only solution to the problem  $\mathcal{L}v = 0$  on  $\mathcal{O}$  and  $v = 0$  on  $\Gamma$ . Because the definitions (2.4) and (2.5) imply that  $\mathcal{L}(K\mathcal{L}u_f + \tilde{g}) = \mathcal{L}K(\mathcal{L}u_f) + 0 = \mathcal{L}u_f$  on  $\mathcal{O}$ , and  $K\mathcal{L}u_f + \tilde{g} = 0 + g$  on  $\Gamma$ , it then follows that

$$u_f = K\mathcal{L}u_f + \tilde{g}. \quad (2.6)$$

Since the function  $\tilde{g}$  is known, this and Equation (2.3) lead to the adapted observational model

$$\tilde{Y}_N := Y_N - \tilde{g} = K(\mathcal{L}u_f) + \frac{1}{\sqrt{N}}\dot{\mathbb{W}}. \quad (2.7)$$

Thus the problem of estimating  $\mathcal{L}u_f$  has been reduced to solving the linear inverse problem  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\dot{\mathbb{W}}$ , defined by the operator  $K$ . The solution  $v$  to this problem is next substituted as  $v = \mathcal{L}u_f$  in the inversion map from Equation (2.2) to find  $f$ . The induced distribution of  $f$  when this map is applied to a posterior distribution of  $v$  is the posterior distribution of  $f$  for the same prior (see Section 2.2 for a precise statement).

A Bayesian approach for solving non-linear inverse problems could start by putting a prior on the function  $f$ , as in (Nickl 2018; Giordano and Nickl 2020; Monard, Nickl and Paternain 2021a; Monard, Nickl and Paternain 2021b; Kerkkonen 2022; Nickl 2023), who put a Gaussian process prior on  $f$  or  $\log f$  or independent priors on the coefficients of an orthonormal wavelet basis of the domain. However, since our aim is to recover  $\mathcal{L}u_f$  in the observational model (2.7), an easier approach is to put a prior on  $\mathcal{L}u_f$  or (almost equivalently) on  $u_f$ . Methods for computing the posterior distribution of  $\mathcal{L}u_f$ , as well as theoretical results on this posterior distribution are then available (e.g. (Knapik, van der Vaart and van Zanten 2011; Knapik, van der Vaart and van Zanten 2013; Ray 2013; Agapiou, Larsson and Stuart 2013; Knapik, Szabó et al. 2016; Gugushvili,

van der Vaart and Yan 2020; Yan 2020; Yan and van der Vaart 2023)). The main purpose of the present chapter is to show that such methods can also lead to accurate recovery of the function  $f$ . Nothing appears to be lost at the level of contraction rates of the posterior distribution. Furthermore, our numerical experiments also show satisfactory performance.

As usual, the posterior distribution is the conditional distribution given the observed data. In the Bayesian setup, we assume that the function  $u_f$  is generated from a prior distribution, and given  $u_f$  the data  $Y_N$  is distributed according to Equation (2.3). Next, the posterior distributions of  $u_f$ ,  $\mathcal{L}u_f$ , or  $f$  are their conditional distributions given  $Y_N$ . Because  $f$  is determined uniquely by  $\mathcal{L}u_f$  in view of Equation (2.2), the posterior distribution of  $\mathcal{L}u_f$  induces the posterior distribution of  $f$ . We denote all posterior distributions by  $\Pi_N(\cdot \mid Y_N)$ , where the argument  $\cdot$  can be specialised to a set concerning  $u_f$ ,  $\mathcal{L}u_f$  or  $f$ .

We obtain posterior contraction rates, as usual (Ghosal, Ghosh and van der Vaart 2000), in the non-Bayesian setup where the observation is assumed to be generated according to the model (2.1) and Equation (2.3) with a fixed parameter  $f_0$ . For instance, a contraction rate for  $f$  relative to a given norm  $\|\cdot\|$  is a sequence  $\epsilon_N \rightarrow 0$  such that  $\Pi_N(f : \|f - f_0\| \geq M_N \epsilon_N \mid Y_N) \rightarrow 0$  in probability, for every  $M_N \rightarrow \infty$ . This rate can be compared to an optimal contraction rate, which is determined, in a minimax sense, by the smoothness of the function  $f_0$ , and is typically obtainable by using a prior with similar smoothness. By mixing over priors of different smoothness levels (hierarchical Bayes) or using a prior of data-determined smoothness level (empirical Bayes), optimal contraction rates for a range of different smoothness levels are obtainable by a single prior. An advantage of our approach is that such “adaptive” contraction rates for the linear inverse problem of estimating  $\mathcal{L}u_f$  carry over into adaptive rates for the recovery of  $f$ . Adaptive priors for the linear inverse problem have been constructed by using a scale of priors of fixed smoothness simultaneously through empirical or hierarchical Bayes methods (Knapik, Szabó et al. 2016; Gugushvili, van der Vaart and Yan 2020; Szabó, van der Vaart and van Zanten 2013).

Besides using a posterior distribution for estimating the parameter at a (nearly) minimax rate, it also can be used for uncertainty quantification. This can be justified, or not, from a non-Bayesian point of view by the coverage of credible sets, i.e. the probability that a (central) set of given posterior probability covers the true parameter  $f_0$ . The coverage of credible intervals for smooth functionals of the parameter, or of credible balls in a very weak norm, follows from a suitable version of the Bernstein-von Mises theorem (Castillo 2012; Castillo and Nickl 2014; Ray 2017; Nickl 2018; Monard, Nickl and Paternain 2021a). For infinite-dimensional credible sets that involve a bias-variance trade-off, the situation is

more involved. This problem has been studied for linear inverse problems in (Szabó, van der Vaart and van Zanten 2015; Sniekers and van der Vaart 2020). Since our method is based on transforming the non-linear inverse problem at hand into a linear one, the positive results showing the existence of reasonable credible sets in the linear inverse problem translate into existence in the problem of interest.

The approach of this chapter is partly motivated by computational efficiency. The recent paper (Nickl and Wang 2020) derived a Langevin-type algorithm to compute the posterior mean, which provenly takes only polynomial time. The algorithm is based on only the forward map  $f \mapsto u_f$ , which is an advantage for many problems. However, repeatedly evaluating the forward map in MCMC algorithms may still be heavy for some practical applications with medium or large data sets. In order to compute a draw from the posterior of  $f$ , the inverse map from Equation (2.2) only needs to be computed once for a draw from the posterior of  $\mathcal{L}u_f$ . This provides a significant advantage compared to MCMC methods, where the forward map often needs to be computed multiple times if one wants to generate a single draw from the posterior. A further gain is that our approach may combine well with distributed computation based on partitioning the data in space. Since the observation is a noisy version of the function  $u_f : \mathcal{O} \rightarrow \mathbb{R}$ , a (possibly overlapping) partition of the domain  $\mathcal{O}$ , followed by separate reconstruction of  $\mathcal{L}u_f$  on the partitioning sets, is feasible. Results from Chapter 3 show that such a distributional approach may give theoretically optimal and practically competitive results compared to other computational shortcuts. In particular, spatial partitioning allows priors and posteriors to adapt to the smoothness of the true function (for the regression problem without inversion, see Chapter 4 of (Hadji 2023)). In comparison, an approach based on partitioning the domain of the function of  $f$  seems impossible, due to the structure of the partial differential equation.

The chapter is organised as follows. In Section 2.2 we formalise the relation between the non-linear inverse problem of interest and the auxiliary linear inverse problem. Next in Section 2.3 we review and extend Bayesian methods for the linear inverse problem, describe prior choices, with both fixed and data-driven choices of the tuning parameter, and state the resulting results on contraction rates and frequentist coverage of credible sets. In Sections 2.4–2.7 we apply our approach to four examples of PDEs. In Section 2.8 we present numerical simulations that illustrate our approach in various settings. We conclude with a discussion of the approach and open questions in Section 2.9. Sections 2.A–2.C provide technical complements. Sections 2.D and 2.E present posterior contraction rates in smoothness norms for the linear inverse problem



in the Gaussian white noise, and the discrete observation settings, respectively.

## 2.2 General method

In this section, we present the skeleton of our approach, specialised to rates of contraction and uncertainty quantification. In the next sections, we verify the generic assumptions of the present section for combinations of priors and PDE models.

Given some prior distribution on the function  $v = \mathcal{L}u_f$ , we consider the two posterior distributions corresponding to the models (2.7) and Equation (2.3):

- L. the posterior  $\tilde{\Pi}_N(v \in \cdot \mid \tilde{Y}_N)$  of  $v$  in the observational model  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}} \dot{\mathbb{W}}$ ;
- N. the posterior  $\Pi_N(f \in \cdot \mid Y_N)$  of  $f$  in the observational model  $Y_N = u_f + N^{-\frac{1}{2}} \dot{\mathbb{W}}$ .

The linear model L and the non-linear model N may be interpreted in terms of the white noise model, or be understood as vectors in  $\mathbb{R}^N$ , for given  $x_{N,1}, \dots, x_{N,N} \in \mathcal{O}$ , equal to the sum of the vectors  $(Kv(x_{N,i}) : i = 1, \dots, N)$  and  $(u_f(x_{N,i}) : i = 1, \dots, N)$  and a standard normal vector, in the linear and non-linear model, respectively.

The two models are linked through the definition of  $u_f$  through (2.1) and the identity in Equation (2.6). Together with Equation (2.2) these relations allow us to deduce properties of the second posterior distribution from the first. Thus we adopt (2.1), (2.6), and (2.2) in the remainder of the chapter.

### 2.2.1 Posterior contraction rate

We assume that the posterior distribution of  $v$  is a Borel law on a normed space  $(V, \|\cdot\|)$  and possesses a rate of contraction  $\epsilon_N$  to  $v_0$  in the sense that  $\tilde{\Pi}_N(v : \|v - v_0\| \geq \epsilon_N M_N \mid \tilde{Y}_N) \xrightarrow{P} 0$ , for any  $M_N \rightarrow \infty$ , under the assumption that  $v_0$  gives the true distribution of  $\tilde{Y}_N = Kv_0 + N^{-\frac{1}{2}} \dot{\mathbb{W}}$ .

**Proposition 2.2.1.** *Suppose that the posterior distribution  $\tilde{\Pi}_N(v \in \cdot \mid \tilde{Y}_N)$  of  $v$  contracts under  $v_0$  to  $v_0 = \mathcal{L}u_{f_0}$  at rate  $\epsilon_N$  in  $(V, \|\cdot\|)$  and satisfies  $\Pi_N(v \in V_N \mid \tilde{Y}_N) \xrightarrow{P} 1$  for given sets  $V_N \subset V$ . If Equation (2.2) holds for a map  $e$  such that  $e : V_N \rightarrow L^2$  is Lipschitz at  $v_0$ , then the posterior distribution of  $f$  attains a rate of contraction  $\epsilon_N$  under  $f_0$  relative to the  $L^2$ -norm.*

*Proof.* By assumption the prior distributions of  $\mathcal{L}u_f$  and  $v$  are the same, and in view of Equations (2.6) and (2.7), the conditional distributions of  $Y_N - \tilde{g} = K\mathcal{L}u_f + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$  given  $\mathcal{L}u_f$  and  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$  given  $v$  are the same as well, under the identification  $v = \mathcal{L}u_f$ . It follows that the conditional distributions of  $\mathcal{L}u_f$  given  $Y_N - \tilde{g}$  and of  $v$  given  $\tilde{Y}_N$  are the same. More precisely, if  $(y, B) \mapsto L(\tilde{Y}_N, \cdot)$  is a Markov kernel such that  $v \mid \tilde{Y}_N \sim L(\tilde{Y}_N, \cdot)$ , then  $\mathcal{L}u_f \mid Y_N \sim L(Y_N - \tilde{g}, \cdot)$ .

By Equation (2.2), we have that  $f - f_0 = e(\mathcal{L}u_f) - e(\mathcal{L}u_{f_0})$  and hence  $\Pi_N(f : \|f - f_0\|_{L^2} < \epsilon, \mathcal{L}u_f \in V_N \mid Y_N) = \Pi_N(f : \|e(\mathcal{L}u_f) - e(\mathcal{L}u_{f_0})\|_{L^2} < \epsilon, \mathcal{L}u_f \in V_N \mid Y_N)$ . By the preceding paragraph this is the same as  $L(Y_N - \tilde{g}, B)$ , for  $B = \{v \in V_N : \|e(v) - e(v_0)\|_{L^2} < \epsilon\}$ . Since  $Y_N - \tilde{g}$  is distributed as  $\tilde{Y}_N$  under  $v_0$ , we obtain that  $L(Y_N - \tilde{g}, B) \sim \tilde{\Pi}_N(B \mid \tilde{Y}_N)$ . The assumptions that  $e$  is Lipschitz on  $V_N$  and that  $\tilde{\Pi}_N(V_N \mid \tilde{Y}_N) \xrightarrow{P} 1$  show that there exists a constant  $C$  such that  $\tilde{\Pi}_N(B \mid \tilde{Y}_N) \leq \tilde{\Pi}_N(v : \|v - v_0\| \leq C\epsilon \mid \tilde{Y}_N)$  on a set of probability tending to 1. The last probability tends to zero in probability for  $\epsilon = M_N \epsilon_N$  and every  $M_N \rightarrow \infty$ .  $\square$

The proposition uses the  $L^2$ -norm on the functions  $f$  for definiteness, and allows a general norm  $\|\cdot\|$  on the functions  $\mathcal{L}u_f$  for flexibility. It connects the posteriors for the two problems (2.1) and (2.7) in an abstract way, and needs to be made precise for particular problems. The input is a  $\|\cdot\|$ -contraction rate in the problem  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$  at  $v_0 = \mathcal{L}u_{f_0}$  and the guarantee that the posterior in this problem concentrates on sets  $V_N$  on which the solution map from Equation (2.2) is Lipschitz.

In our examples, the Lipschitz assumption relative to a natural norm, is usually mild. For instance, it can be verified by ascertaining that the posterior distribution  $\tilde{\Pi}_N(Kv \in \cdot \mid \tilde{Y}_N)$  of  $Kv$  concentrates on functions that are bounded away from zero, or functions whose derivative is well-behaved. We show this by establishing consistency of this posterior distribution for a relevant norm combined with the relevant assumption on the true function  $Kv_0$ . Because  $K$  is smoothing, uniform consistency for  $Kv$  is typically automatic for the usual priors and often sufficient.

## 2.2.2 Uncertainty quantification

Credible sets for the linear inverse problem map naturally into credible sets for the non-linear problem of interest. Assume we are given a credible set  $\tilde{C}_N(\tilde{Y}_N)$  for  $v$  based on the observation  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$  with a prescribed probability under the posterior distribution  $\tilde{\Pi}_N(\cdot \mid \tilde{Y}_N)$ . We transform this in the set of

functions  $f$  given by

$$C_N(Y_N) := \{f : \mathcal{L}u_f \in \tilde{C}_N(Y_N - \tilde{g})\} = \{e(v) : v \in \tilde{C}_N(Y_N - \tilde{g})\}, \quad (2.8)$$

where  $e$  is the solution map given in Equation (2.2).

The *credible level* of the credible set  $\tilde{C}_N(\tilde{Y}_N)$  for  $v$  in the model  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{W}$  is by definition the posterior probability  $\tilde{\Pi}_N(\tilde{C}_N(\tilde{Y}_N) \mid \tilde{Y}_N)$  and its *coverage* at  $v_0$  is by definition the probability  $\Pr_{v_0}(\tilde{C}_N(\tilde{Y}_N) \ni v_0)$  if  $\tilde{Y}_N = Kv_0 + N^{-\frac{1}{2}}\tilde{W}$ . For the model  $Y_N = u_f + N^{-\frac{1}{2}}\tilde{W}$  the same quantities are defined analogously relative to the parameter  $f$  with true value  $f_0$ .

In the following proposition, we identify  $\tilde{Y}_N$  and  $Y_N - \tilde{g}$ , so that the assertions can be understood in an almost sure sense.

**Proposition 2.2.2.** *The credible levels of the credible sets  $\tilde{C}_N(\tilde{Y}_N)$  for  $v$  and  $C_N(Y_N)$  for  $f$  are equal, and so are the coverage levels of these sets at  $v_0 = \mathcal{L}u_{f_0}$  and  $f_0$ , respectively. Furthermore, if the map  $e : V_N \rightarrow L^2$  in Equation (2.2) is uniformly Lipschitz at points  $\bar{v}_N \in V_N$  and  $\tilde{C}_N(\tilde{Y}_N) \subset V_N$ , then on the event  $\bar{v}_N \in \tilde{C}_N(\tilde{Y}_N)$  the  $L^2$ -diameters of the sets  $C_N(Y_N)$  are of the same order in probability under  $f_0$  as the  $L^2$ -diameters of the sets  $\tilde{C}_N(\tilde{Y}_N)$  under  $v_0$ .*

*Proof.* As noted in the proof of Proposition 2.2.1, if the Markov kernel  $(y, B) \mapsto L(\tilde{Y}_N, \cdot)$  gives the posterior distribution of  $v$  given  $\tilde{Y}_N$ , then  $L(Y_N - \tilde{g}, \cdot)$  gives the posterior distribution of  $\mathcal{L}u_f$  given  $Y_N$ . By the definition of  $C_N(Y_N)$  it follows that  $L(\tilde{Y}_N, \tilde{C}_N(\tilde{Y}_N)) = \Pi_N(f \in C_N(Y_N) \mid Y_N) = \Pi_N(\mathcal{L}u_f \in \tilde{C}_N(Y_N - \tilde{g}) \mid Y_N) = L(Y_N - \tilde{g}, \tilde{C}_N(Y_N - \tilde{g}))$ , which is the same as  $\tilde{\Pi}_N(\tilde{C}_N(\tilde{Y}_N) \mid \tilde{Y}_N)$ , since  $\tilde{Y}_N = Y_N - \tilde{g}$ . Similarly  $\Pr_{f_0}(f_0 \in C_N(Y_N)) = \Pr_{f_0}(\mathcal{L}u_{f_0} \in \tilde{C}_N(Y_N - \tilde{g}))$  is the same as  $\Pr_{v_0}(v_0 \in \tilde{C}_N(\tilde{Y}_N))$ .

The  $L^2$ -diameter  $\sup\{\|f - g\|_{L^2} : f, g \in C_N(Y_N)\}$  of  $C_N(Y_N)$  is equal to  $\sup\{\|e(v) - e(w)\|_{L^2} : v, w \in \tilde{C}_N(\tilde{Y}_N)\} \leq 2 \sup\{\|e(v) - e(\bar{v}_N)\|_{L^2} : v \in \tilde{C}_N(\tilde{Y}_N)\}$ , which is bounded by a constant times  $\sup\{\|v - \bar{v}_N\| : v \in \tilde{C}_N(\tilde{Y}_N)\}$ , because  $\tilde{C}_N(\tilde{Y}_N) \subset V_N$  and  $e$  is uniformly Lipschitz at  $\bar{v}_N \in V_N$ . On the event  $\bar{v}_N \in \tilde{C}_N(\tilde{Y}_N)$  the right side is bounded above by the diameter of  $\tilde{C}_N(\tilde{Y}_N)$ .  $\square$

Thus credible and coverage levels translate from the linear to the non-linear inverse problem for arbitrary credible sets. To retain also the size of the sets it may be necessary to restrict the starting sets  $\tilde{C}_N(\tilde{Y}_N)$  to sets  $V_N$  on which the solution map from Equation (2.2) is locally Lipschitz.

In our examples, the latter can be achieved without affecting the credible or coverage levels. Indeed, in general as soon as  $\mathbb{E}_{v_0}\tilde{\Pi}_N(V_N \mid \tilde{Y}_N) \rightarrow 1$  and  $\Pr_{v_0}(V_N \ni v_0) \rightarrow 1$ , the credible and coverage levels of the sets  $\tilde{C}_N(\tilde{Y}_N)$  and

$\tilde{C}_N(\tilde{Y}_N) \cap V_N$  are asymptotically the same. In our examples, we choose the sets  $V_N$  for instance a uniform ball  $V_N = \{v : \|Kv - K\hat{v}_N\|_\infty \leq c_0\}$  of small radius around uniformly consistent estimators  $K\hat{v}_N$  of  $Kv_0$ . Then the uniform consistency gives  $\Pr_{v_0}(V_N \ni v_0) \rightarrow 1$ , and consistency relative to the uniform norm of the posterior distribution of  $Kv$  gives  $\mathbb{E}_{v_0} \tilde{\Pi}_N(V_N \mid \tilde{Y}_N) \rightarrow 1$ .

## 2.3 Bayesian analysis in the linear problem

In this section, we review and extend results on the problem of recovering a function  $v$  from the noisy observation  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{W}$ , where  $K$  is the operator given in Equation (2.4). When applied to solving the non-linear problem from Equation (2.3), the function  $v$  will be taken equal to  $v = \mathcal{L}u_f$ . We restrict to mildly ill-posed problems in the sense of (Cavalier 2008). In Sections 2.3.1–2.3.3 we consider the white noise model, whereas in Section 2.3.4 we consider the discrete observational model.

We consider Gaussian priors on  $v$ , or mixtures thereof, possibly with tuning parameters set by the empirical Bayes method. (Non-Gaussian priors were considered in (Agapiou, Larsson and Stuart 2013; Ray 2013) and (Gugushvili, van der Vaart and Yan 2020; Yan 2020).) Without loss of generality such a prior can be constructed by equipping the coefficients of  $v$  in a basis expansion with (conditionally) independent univariate normal distributions. For a given orthonormal basis  $(h_i)$  of  $L^2$  and  $v = \sum_{i=1}^\infty v_i h_i$ , we set

$$(v_1, v_2, \dots) \mid \tau, \alpha \sim \bigotimes_{i=1}^\infty N(0, \tau^2 i^{-1-2\alpha/d}). \quad (2.9)$$

The hyperparameters  $\tau$  or  $\alpha$  determine the smoothness of the prior, and may be chosen fixed, given a hyper prior or set by the empirical Bayes method.

### 2.3.1 Smoothness scales

The basis  $(h_i)$  give rises to a scale of function classes  $S^s$ , for  $s \geq 0$ , consisting of all functions  $v = \sum_{i=1}^\infty v_i h_i \in L^2$  with  $\|v\|_{S^s} < \infty$ , for

$$\|v\|_{S^s}^2 = \sum_{i=1}^\infty v_i^2 i^{2s/d}. \quad (2.10)$$

For  $s < 0$  we define the same norm on the sequences  $(v_i)$  and define  $S^s$  as the corresponding normed space; this can also be identified with the dual space of

$S^{-s}$ . Assume that the operator  $K : S^0 \rightarrow L^2$  is *smoothing of order  $p$*  in the scale  $S^s$  in the sense that

$$\|Kv\|_{L^2} \asymp \|v\|_{S^{-p}}, \quad v \in S^0. \quad (2.11)$$

**Theorem 2.3.1.** *Let  $K : S^0 \rightarrow L^2$  be a linear operator satisfying Equation (2.11).*

- (i) *If  $v_0 \in S^\beta$  and  $-p \leq \delta < \alpha \wedge \beta$ , then the  $S^\delta$ -contraction rate to  $v_0$  of the posterior distribution of  $v$  in the white noise model is of the order  $N^{-(\alpha \wedge \beta - \delta)/(2\alpha + 2p + d)}$ .*
- (ii) *If either  $\sup_{i \in \mathbb{N}} i^{u/d} \|Kh_i\|_\infty < \infty$ , for some  $u$  with  $d/2 < \alpha \wedge \beta + u$ , or  $K : S^{\gamma-p} \rightarrow S^\gamma$  is continuous and  $\|\cdot\|_{H^\gamma} \lesssim \|\cdot\|_{S^\gamma}$  for some  $\gamma$  with  $d/2 < \gamma < \alpha \wedge \beta + p$ , then the posterior distribution of  $Kv$  is consistent for  $Kv_0$  relative to the uniform norm.*

*Proof.* The contraction rate (i) is an extension of Theorem 6.1 in (Gugushvili, van der Vaart and Yan 2020) to the  $\|\cdot\|_{S^\delta}$ -norm, proved in Theorem 2.D.2.

For the uniform consistency (ii) of the posterior distribution of  $Kv$  under the first condition, we argue that  $Kv = \sum_{i=1}^\infty v_i Kh_i$ , and hence  $\|Kv\|_\infty^2 \leq \|v\|_{S^\delta}^2 \sum_{i=1}^\infty \|Kh_i\|_\infty^2 i^{-2\delta/d}$ , by the Cauchy-Schwarz inequality. Under the condition on the uniform norms of the functions  $Kh_i$ , this is bounded above by a multiple of  $\|v\|_{S^\delta}^2$ , for  $(2\delta + 2u)/d > 1$ , equivalently  $\delta + u > d/2$ . Thus the uniform consistency follows from the contraction in  $S^\delta$ , which takes place for  $\delta < \alpha \wedge \beta$ . Under the condition  $\alpha \wedge \beta + u > d/2$ , a value of  $\delta$  that satisfies both requirements exists. Under the second condition, we use Sobolev embedding to see that  $\|Kv\|_\infty \lesssim \|Kv\|_{H^\gamma}$ , for  $\gamma > d/2$ , which by assumption is bounded above by  $\|Kv\|_{S^\gamma} \lesssim \|v\|_{S^{\gamma-p}}$ . Thus the uniform consistency follows from contraction in  $S^\delta$  for  $\delta = \gamma - p$ , which takes place if  $\delta < \alpha \wedge \beta$ .  $\square$

The preceding theorem assumes fixed hyperparameters. Theorem 7.2 of (Gugushvili, van der Vaart and Yan 2020) shows that the mixture prior obtained by choosing a parameter  $\tau$  from an inverse Gamma distribution and a fixed  $\alpha$  gives an  $L^2$ -contraction rate  $N^{-\beta/(d+2\beta+2p)}$ , for any  $\beta \in (0, \alpha]$ , but has not been extended to rates in  $S^\delta$ .

### 2.3.2 Eigenbasis

If  $K : L^2 \rightarrow L^2$  is compact, then we may employ the eigenbasis  $(h_i)$  of the self-adjoint operator  $K^\top K : L^2 \rightarrow L^2$ . If  $\kappa_i^2$  are the eigenvalues of  $K^\top K$ , then

$\|Kv\|_{L^2}^2 = \langle K^\top Kv, v \rangle_{L^2} = \sum_{i=1}^{\infty} v_i^2 \kappa_i^2$ . Hence  $K$  is smoothing of order  $p$  in the sense of Equation (2.11) if

$$\kappa_i \asymp i^{-p/d}. \quad (2.12)$$

If  $K$  is self-adjoint, then  $K^s v_i = \kappa_i^s v_i$ , and  $K$  is also smoothing in the extended sense that  $\|Kv\|_{S^s} \asymp \|v\|_{S^{s-p}}$ , for every  $s \in \mathbb{R}$ .

For the prior constructed relative to the eigenbasis, the white noise problem is equivalent to observing a sequence  $(\tilde{Y}_{N,1}, \tilde{Y}_{N,2}, \dots)$  of independent normal variables with  $\tilde{Y}_{N,i} \mid v \sim N(\kappa_i v_i, 1/N)$ , equipped with conditionally independent priors  $v_i \mid \alpha, \tau \sim N(0, \tau^2 i^{-1-2\alpha/d})$ . Conditionally on the hyperparameters, the problem is conjugate and the posterior can be obtained explicitly. The contraction rates for fixed hyperparameters follow from Theorem 2.3.1, but contraction rates for priors with hyperparameters set by the empirical and hierarchical Bayes methods were studied in (Szabó, van der Vaart and van Zanten 2013; Knapik, Szabó et al. 2016; Rousseau and Szabo 2017), and credible balls and adaptive credible balls in (Knapik, van der Vaart and van Zanten 2011; Szabó, van der Vaart and van Zanten 2015). (The parameter  $\alpha$  in (Knapik, van der Vaart and van Zanten 2011; Szabó, van der Vaart and van Zanten 2013; Szabó, van der Vaart and van Zanten 2015; Knapik, Szabó et al. 2016) is denoted by  $\alpha/d$  in the present chapter.)

The hierarchical Bayes method can be implemented with fixed  $\tau$  and a prior on  $\alpha$  with density  $\lambda$  on  $(0, \infty)$  chosen such that, for every  $c_1 > 0$  there exist  $c_2 \geq 0$ ,  $c_3 \in \mathbb{R}$  and  $c_4 > 1$ , with  $c_3 > 1$  if  $c_2 = 0$ , such that for  $\alpha \geq c_1$ ,

$$c_4^{-1} \alpha^{-c_3} e^{-c_2 \alpha} \leq \lambda(\alpha) \leq c_4 \alpha^{-c_3} e^{-c_2 \alpha}. \quad (2.13)$$

The empirical Bayes method chooses  $\alpha$  equal to the maximum likelihood estimator, restricted to  $[0, \log N]$ , based on the Bayesian marginal distribution of the observation  $\tilde{Y}_N$ . This can be seen to be

$$\hat{\alpha}_N = \arg \min_{\alpha} \sum_{i=1}^{\infty} \left( \log \left( 1 + \frac{N}{i^{1+2\alpha/d} \kappa_i^{-2}} \right) - \frac{n^2}{i^{1+2\alpha/d} \kappa_i^{-2} + N} \tilde{Y}_{N,i}^2 \right). \quad (2.14)$$

Alternatively, the parameter  $\alpha$  may be fixed and the parameter  $\tau$  chosen by empirical Bayes or equipped with a prior, see (Szabó, van der Vaart and van Zanten 2013).

The papers (Knapik, van der Vaart and van Zanten 2011; Szabó, van der Vaart and van Zanten 2015) studied credible balls of the type, for  $\hat{v}_N$  the posterior mean and  $r_N$  set to the smallest value such that  $\{v : \|v - \hat{v}_N\|_{L^2} \leq r_N\}$  possesses a prescribed posterior probability in  $(0, 1)$ ,

$$\tilde{C}_N(\tilde{Y}_N) = \{v : \|v - \hat{v}_N\|_{L^2} \leq cr_N\}. \quad (2.15)$$

In (Knapik, van der Vaart and van Zanten 2011) priors with fixed hyperparameters  $(\tau, \alpha)$  were shown to give frequentist coverage provided  $\alpha$  undershoots the true smoothness. As it is impossible to construct confidence sets of uniform coverage and rate-adaptive size, see e.g. (Cai and Low 2004; Robins and van der Vaart 2006), neither the empirical nor the hierarchical Bayes methods can provide credible sets with prescribed uniform frequentist coverage over all possible parameters. However, in (Szabó, van der Vaart and van Zanten 2015; Sniekers and van der Vaart 2015) the credible sets  $\tilde{C}_N(\tilde{Y}_N)$  are shown to be confidence sets for true parameters whose coefficients decrease not too irregularly, so-called *polished tail* sequences.

We record these results, together with extensions of contraction in  $S^\delta$  and uniform consistency, in the following theorem.

**Theorem 2.3.2.** *Consider the prior based on the eigenbasis with coefficients satisfying Equation (2.9) with fixed  $\tau$  and  $\alpha$  determined by the hierarchical or empirical Bayes method in Equation (2.13) or (2.14).*

- (i) *If  $v_0 \in S^\beta$  and  $-p \leq \delta < \beta$ , then the  $S^\delta$ -contraction rate of the posterior distribution of  $v$  in the white noise model is  $l_N N^{-(\beta-\delta)/(2\beta+2p+d)}$ , for  $l_N = (\log N)^{3/2}(\log \log N)^{\frac{1}{2}}$ .*
- (ii) *If either  $\sup_{i \in \mathbb{N}} i^{u/d} \|Kh_i\|_\infty < \infty$ , for some  $u$  with  $d/2 < \beta + u$ , or  $\|\cdot\|_{H^\gamma} \lesssim \|\cdot\|_{S^\gamma}$  for some  $\gamma$  with  $d/2 < \gamma < \beta + p$ , then the posterior distribution of  $Kv$  is consistent for the uniform norm.*
- (iii) *If the prior is constructed with fixed hyperparameters  $\tau$  and  $\alpha$ , then the coverage tends to 1 provided  $\alpha < \beta$  and the  $L^2$ -diameter is of the order  $O_P(N^{-(\alpha \wedge \beta)/(2\alpha+2p+d)})$ .*
- (iv) *If  $v_0$  satisfies the polished tail condition, then for sufficiently large  $c$  the coverage of the sets in Equation (2.15) tends to 1 and the  $L^2$ -diameter is of the order  $O_P(l_N N^{-\beta/(2\beta+2p+d)})$ .*

*Proof.* Assertions (iii) and (iv) follow immediately from the results in (Knapik, van der Vaart and van Zanten 2011) and (Szabó, van der Vaart and van Zanten 2015). Assertion (ii) follows from assertion (i) by the arguments in the proof of Theorem 2.3.1.

For the proof of (i), we extend the results of (Knapik, Szabó et al. 2016), which are based on the explicit representation of the posterior distribution: for  $\lambda_i(\alpha) = i^{-1-2\alpha/d}$  and  $\kappa_i^2 \asymp i^{-2p/d}$ , the posterior distribution satisfies  $v_i - v_{i,0} \mid$

$(\tilde{Y}_N, \alpha) \sim N(m_i(\alpha), \sigma_i^2(\alpha))$ , for

$$m_i(\alpha) := \frac{n\lambda_i(\alpha)\kappa_i\tilde{Y}_{N,i}}{1 + N\lambda_i(\alpha)\kappa_i^2} - v_{0,i}, \quad \sigma_i^2(\alpha) := \frac{\lambda_i(\alpha)}{1 + N\lambda_i(\alpha)\kappa_i^2}.$$

Theorem 1 in (Knapik, Szabó et al. 2016) shows that the empirical Bayes estimator from Equation (2.14) is, with probability tending to 1 under  $v_0$ , bounded below and above by deterministic sequences  $\underline{\alpha}_N$  and  $\bar{\alpha}_N$ , characterised as crossing points of a certain deterministic criterion function  $h_N$ , derived from the likelihood. In the proof of Theorem 3 in the same reference, it is shown that in the hierarchical Bayes case, the posterior probability that  $\alpha$  falls in the interval  $[\underline{\alpha}_N, \bar{\alpha}_N]$  tends in probability to one. The lower bound satisfies  $\underline{\alpha}_N \geq \beta - c_0/\log N$ , for some constant  $c_0$ , which ensures that in both cases the rate of contraction does not fall essentially below the optimal rate for a  $\beta$ -smooth parameter.

It follows that, for the proof of (i) for both methods it suffices to show, for any  $M_N \rightarrow \infty$ ,

$$\sup_{\underline{\alpha}_N \leq \alpha \leq \bar{\alpha}_N} \Pi_N \left( v : \|v - v_0\|_{S^\delta} > M_N l_N N^{-\frac{\beta-\delta}{d+2\beta+2p}} \mid \tilde{Y}_N, \alpha \right) \xrightarrow{P_{v_0}} 0. \quad (2.16)$$

By Markov's inequality, the posterior probability can be bounded above by  $(M_N \epsilon_N)^{-2} \mathbb{E}(\|v - v_0\|_{S^\delta}^2 \mid \tilde{Y}_N, \alpha)$ , for  $\epsilon_N$  the rate of contraction. The second moment can be computed using the explicit form of the posterior distribution. It follows that the preceding display is valid if the following two relations hold:

$$\begin{aligned} \sup_{\underline{\alpha}_N \leq \alpha \leq \bar{\alpha}_N} \sum_{i=1}^{\infty} m_i^2(\alpha) i^{2\delta/d} &= O_{P_{v_0}} \left( l_N^2 N^{-\frac{2(\beta-\delta)}{d+2\beta+2p}} \right), \\ \sup_{\underline{\alpha}_N \leq \alpha \leq \bar{\alpha}_N} \sum_{i=1}^{\infty} \sigma_i^2(\alpha) i^{2\delta/d} &= O \left( l_N^2 N^{-\frac{2(\beta-\delta)}{d+2\beta+2p}} \right). \end{aligned}$$

The left side of the second equation is deterministic and monotonously decreasing in  $\alpha$ , and hence it is sufficient to consider the series at  $\alpha = \underline{\alpha}_N$ . By elementary calculus, this can be seen to be of order  $N^{-2(\underline{\alpha}_N - \delta)/(d+2\underline{\alpha}_N+2p)} \lesssim N^{-2(\beta-\delta)/(d+2\beta+2p)}$ .

The analysis of the first relation is considerably harder, as it is random and non-monotone. However, the relation can be verified by the same arguments as in (Knapik, Szabó et al. 2016).  $\square$

In the preceding results, the proof of uniform consistency is essentially based on Sobolev embedding. The following lemma shows that this approach is generally possible for parabolic or elliptic differential operators.



**Lemma 2.3.3.** *Let  $(S^s)$  be the eigenscale of the inverse operator  $K$  (as in Equation (2.4)) of a parabolic or elliptic differential operator  $\mathcal{L}$  under the Dirichlet boundary condition. Then  $K : S^\gamma \rightarrow L_\infty$  is continuous for every  $\gamma > d$ .*

*Proof.* See Theorem B.2 in (Kinderlehrer and Stampacchia 2000) for elliptic  $\mathcal{L}$ , and Theorem 4.5 in (Tröltzsch 2010) (adapted to Dirichlet boundary condition) for parabolic  $\mathcal{L}$ .  $\square$

### 2.3.3 Sobolev spaces

The scale of Sobolev spaces  $H^s$  is of special interest in connection to differential operators. There are several bases that generate this scale, popular ones being the wavelet bases. Although a wavelet basis is most conveniently indexed by a double index in the form  $(h_{j,k})$ , for a resolution level  $j \in \mathbb{N} \cup \{0\}$  and location index  $k$  (taking of the order  $2^{jd}$  different values at level  $j$ ), the base functions can be ordered in a sequence and the resulting scale takes the form as in Section 2.3.1 (see Example 2.D.1). The Gaussian prior on this sequence constructed as in Section 2.3.1 takes the form  $\sum_{j=1}^{\infty} \sum_k v_{j,k} h_{j,k}$  in the wavelet domain, for independent variables  $v_{j,k} \sim N(0, \tau^2 2^{-2j(\alpha+d/2)})$ .

Thus we obtain the results of Theorem 2.3.1 for a wavelet prior, provided the operator  $K$  is smoothing in the Sobolev scale, as in Section 2.3.1. Unfortunately, this appears to be not often true for the full Sobolev scale, due to the boundary conditions connected to the differential operator. For instance, the inverse of the Laplacian is smoothing (of order 2) relative to the scale  $H_0^1 \cap H^s$ , where  $H_0^1$  is the space of weak solutions of Equation (2.4), the subscript 0 arising from the Dirichlet boundary condition (see Proposition 2.A.1). In such a case the contraction rates provided by Theorem 2.3.1 are still valid for priors constructed from wavelet bases that observe the boundary conditions. We conjecture that an unrestricted wavelet basis would give the same results.

More abstractly, the operator  $K$  will be smoothing relative to the scale generated by the differential operator  $\mathcal{L}$ , and we may construct a Gaussian prior with covariance operator equal to the inverse of  $\mathcal{L}$  (see Proposition 8.19 in (Engl, Hanke and Neubauer 1996)). The equivalence of Hilbert scales resulting from linear differential operators  $\mathcal{L}$  and Sobolev spaces has been investigated extensively, see for instance Theorem 10.1 in (Seeley 1965) for elliptic differential operators  $\mathcal{L}$ .

### 2.3.4 Discrete observations

The preceding results for the white noise model extend to the discrete observational scheme, in which for given “design points”  $x_{N,1}, \dots, x_{N,N} \in \mathcal{O}$ , we observe  $\tilde{Y}_N(1), \dots, \tilde{Y}_N(N)$  given by, for i.i.d. standard normal variables  $Z_i$ ,

$$\tilde{Y}_N(i) = Kv(x_{N,i}) + Z_i, \quad i = 1, \dots, N.$$

The key is an interpolation technique, mapping discrete signals to the continuous domain, as employed in Chapter 8 of (Yan 2020) and (Yan and van der Vaart 2023).

Denote the positive semi-definite Hermitian form by

$$\langle v, w \rangle_N := \frac{1}{N} \sum_{i=1}^N v(x_{N,i})w(x_{N,i}),$$

which can be seen as a semi-definite inner product. Let  $\|v\|_N := \sqrt{\langle v, v \rangle_N}$  be the induced semi-norm. Assume that for every  $N \in \mathbb{N}$ , there exists an  $N$ -dimensional subspace  $\mathbb{L}_N \subset L^2(\mathcal{O}) \cap C(\mathcal{O})$  such that, for constants  $0 < C_1 < C_2 < \infty$ , independent of  $N$ ,

$$C_1\|v\|_{L^2} \leq \|v\|_N \leq C_2\|v\|_{L^2}, \quad v \in \mathbb{L}_N. \quad (2.17)$$

Let  $\mathcal{I}_N : L^2(\mathcal{O}) \cap C(\mathcal{O}) \rightarrow \mathbb{L}_N$  be the map such that  $\mathcal{I}_N v$  is the unique element of  $\mathbb{L}_N$  that interpolates  $v$  at the design points, i.e.  $\mathcal{I}_N v(x_{N,i}) = v(x_{N,i})$ , for every  $i = 1, \dots, N$ . Then assume that for every  $s$  in some interval  $(s_{\min}, s_{\max})$ , and some sequence  $\delta_{N,s} \rightarrow 0$ ,

$$\|\mathcal{I}_N v - v\|_{L^2} \lesssim \delta_{N,s} \|f\|_{S^s}. \quad (2.18)$$

**Theorem 2.3.4.** *If Equations (2.17)–(2.18) hold with  $\delta_{N,\beta} \lesssim cN^{-\beta/(2\alpha+2p+d)}$ , then under the conditions of Theorem 2.3.1 the assertions (i)–(ii) of the theorem are also valid in the discrete observational model.*

*Proof.* The assertion on the contraction rate is an extension of results of (Yan 2020; Yan and van der Vaart 2023) to smoothness norms, which is proved in Section 2.E.

The remainder of the proof is the same as the proof of Theorem 2.3.1.  $\square$

## 2.4 Schrödinger equation

For a bounded domain  $\mathcal{O} \subset \mathbb{R}^d$ , a given nonnegative function  $f \in L^2$  and a bounded, measurable function  $g : \partial\mathcal{O} \rightarrow \mathbb{R}$ , let  $u_f$  be the solution to the time-independent *Schrödinger equation*

$$\begin{cases} \frac{1}{2}\Delta u_f = f u_f & \text{on } \mathcal{O}, \\ u_f = g & \text{on } \partial\mathcal{O}. \end{cases} \quad (2.19)$$

The goal is to recover the unknown “potential”  $f$  using a noisy observation of  $u_f$ . This model serves as a prototypical, benchmark example for elliptic PDEs, with a substantial interest in its own right. It was investigated in the literature using Bayesian methods with a prior distribution on the (logarithm) of the functional parameter  $f$  and the properties of the corresponding posterior were derived using multiscale analysis. In (Nickl, van de Geer and Wang 2020) minimax convergence rates were derived for the MAP estimator corresponding to an optimally rescaled Gaussian process prior in a non-adaptive setting. Posterior contraction rates were derived for a uniform prior on the coefficients in a series expansion (Nickl 2018) and a rescaled Gaussian process prior (Monard, Nickl and Paternain 2021a). Frequentist coverage guarantees for credible balls in weak Sobolev spaces were deduced from uniform non-parametric Bernstein-von Mises results (Nickl 2018; Monard, Nickl and Paternain 2021a).

The solution  $u_f$  to the Schrödinger equation has a probabilistic expression through the Feynman-Kac formula (see Theorem 4.7 in (Zhao and Chung 2012)):

$$u_f(x) = \mathbb{E}^x \left( g(X_\tau) e^{-\int_0^\tau f(X_s) ds} \right), \quad x \in \mathcal{O}. \quad (2.20)$$

Here  $(X_s)_{s \geq 0}$  denotes a  $d$ -dimensional Brownian motion with starting value  $X_0 = x \in \mathcal{O}$ , and  $\tau$  is its exit time from  $\mathcal{O}$ . It is known that  $\sup_x \mathbb{E}^x \tau$  is finite (c.f. Proposition 1.16 in (Zhao and Chung 2012)). If  $f$  is bounded and  $g$  is bounded away from zero on  $\partial\mathcal{O}$ , then Jensen’s inequality applied to Equation (2.20) shows that

$$\inf_{x \in \mathcal{O}} u_f(x) \geq \inf_{y \in \partial\mathcal{O}} g(y) \inf_{z \in \mathcal{O}} e^{-\|f\|_\infty \mathbb{E}^z \tau} > 0.$$

This justifies an inversion formula of the form (2.2):

$$f = \frac{\Delta u_f}{2u_f}. \quad (2.21)$$

We conclude that to estimate  $f$ , it suffices to estimate  $\Delta u_f$  and  $u_f$ . In fact, an estimator for  $\Delta u_f$  suffices, since  $u_f$  can be recovered without error from  $\Delta u_f$  and the known boundary condition  $u_f = g$  on  $\partial\mathcal{O}$ .

This motivates us to take  $\mathcal{L}$  as in (2.1) equal to the Laplacian  $\mathcal{L} = \Delta$ . Its inverse  $K$  as in Equation (2.4) is then the inverse Laplacian under the Dirichlet boundary condition, and the function  $\tilde{g}$  in Equation (2.4) is the solution to the *Poisson equation* with boundary function  $g$ . We record the existence of this operator and function in the following lemma. We assume that the boundary points of  $\mathcal{O}$  are *regular* in the sense of (Gilbarg and Trudinger 2015), page 25; a sufficient condition is that  $\mathcal{O}$  satisfies the exterior sphere condition: for every  $\xi \in \partial\mathcal{O}$  there exists a closed ball  $B \subset \mathbb{R}^d$  such that  $B \cap \bar{\mathcal{O}} = \{\xi\}$ . This is true in particular for the domain  $\mathcal{O} = (0, 1)^d$ .

**Lemma 2.4.1.** *Let  $\mathcal{L} = \Delta$  be the Laplacian. There exists a compact, self-adjoint linear operator  $K : L^2 \rightarrow L^2$  such that Equation (2.4) holds. If  $\mathcal{O}$  is regular, and  $g : \partial\mathcal{O} \rightarrow \mathbb{R}$  is continuous, then there exists a function  $\tilde{g} : \mathcal{O} \rightarrow \mathbb{R}$  satisfying (2.5). If the boundary  $\partial\mathcal{O}$  is  $C^{m+2}$ , for  $m \in \mathbb{N} \cup \{0\}$ , then  $K : H^m \rightarrow H^{m+2}$  is continuous.*

*Proof.* Theorem 3 in Section 6.2 of (Evans 2010) with  $L$  minus the Laplacian together with the example following the theorem show that there exists a unique (weak) solution  $Ku \in H_0^1$  to Equation (2.4) with  $\mathcal{L}$  the Laplacian, for every  $u \in L^2$ . The linearity of the map  $u \mapsto Ku$  is clear, while compactness and self-adjointness is noted in the proof of Theorem 1 on page 335 of (Evans 2010). The continuity  $K : H^m \rightarrow H^{m+2}$  follows from Theorem 5 on page 323 of (Evans 2010).

By Theorem 2.14 in (Gilbarg and Trudinger 2015), the Dirichlet problem on a bounded domain has a solution for every continuous boundary function  $g$  if and only if the boundary points of the domain are regular. By the remark made below the stated theorem, a sufficient condition for this is that the domain satisfies the exterior sphere condition.  $\square$

Thus we can proceed with the general approach outlined in the introduction with  $\mathcal{L} = \Delta$  the Laplacian. Because  $u_f = K\Delta u_f + \tilde{g}$ , Equation (2.21) shows that the solution map in Equation (2.2) can be written as  $f = e(\Delta u_f)$  for

$$e(v) = \begin{cases} \frac{v}{2(Kv + \tilde{g})} & \text{if } \text{ess inf } Kv + \tilde{g} > 0, \\ 0 & \text{otherwise.} \end{cases}$$

**Lemma 2.4.2.** *If  $u_{f_0} = Kv_0 + \tilde{g}$  satisfies  $\inf_{x \in \mathcal{O}} u_{f_0}(x) \geq c_0$ , for some  $c_0 > 0$ , then the map  $e : V \subset L^2 \rightarrow L^2$ , for  $V = \{v : \|Kv - Kv_0\|_\infty \leq c_0/2\}$ , is Lipschitz at every uniformly bounded  $\bar{v} \in V$  with Lipschitz constant bounded by a multiple of  $1 \vee \|\bar{v}\|_\infty$ .*

*Proof.* If  $Kv_0 + \tilde{g} \geq c_0$  and  $\|Kv - Kv_0\|_\infty \leq c_0/2$ , then  $Kv + \tilde{g} \geq c_0/2$ . By the triangle inequality, uniformly in  $v \in V$ ,

$$|e(v) - e(\bar{v})| = \left| \frac{v/2}{Kv + \tilde{g}} - \frac{\bar{v}/2}{K\bar{v} + \tilde{g}} \right| \leq \frac{|v - \bar{v}|}{c_0} + \frac{|Kv - K\bar{v}| \|\bar{v}\|_\infty}{c_0^2}.$$

Thus the  $L^2$ -norm of left-side is bounded above by a multiple of  $\|v - \bar{v}\|_{L^2} + \|Kv - K\bar{v}\|_{L^2} \|\bar{v}\|_\infty \lesssim \|v - \bar{v}\|_{L^2} (1 \vee \|\bar{v}\|_\infty)$ , by the continuity of  $K$ .  $\square$

**Theorem 2.4.3.** *Suppose that  $u_{f_0} = Kv_0 + \tilde{g}$  is bounded away from zero and  $\|v_0\|_\infty < \infty$ . If the posterior distribution of  $Kv$  based on the observation  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$  is consistent for the uniform norm, then the posterior distribution of  $f$  given  $Y_N$  contracts under  $f_0$  at the same rate as the posterior distribution of  $v$  given  $\tilde{Y}_N$  under  $v_0$ . Furthermore, the credible sets  $C_N(Y_N)$  given in Equation (2.8) have diameter of the same order in probability under  $f_0$  as the credible sets  $\tilde{C}_N(\tilde{Y}_N)$  under  $v_0$  provided the latter sets are contained in  $\{v : \|Kv - Kv_0\|_\infty < c_0/2\}$ , for  $c_0 = \inf_{x \in \mathcal{O}} u_{f_0}(x)$ , and contain  $\bar{v}_N$  with  $\|\bar{v}_N\|_\infty = O_P(1)$ .*

*Proof.* We combine Lemma 2.4.2 with Proposition 2.2.1 to obtain the contraction rate, and with Proposition 2.2.2 to obtain the coverage of credible sets.  $\square$

Thus we focus on priors for  $v$  in the problem  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\tilde{\mathbb{W}}$ , where  $K$  is the inverse Laplacian under the Dirichlet boundary condition, that attain fast  $L^2$ -contraction rates and at the same time give uniformly consistent posteriors of  $Kv$ . Theorems 2.3.1, 2.3.2 and 2.3.4 give many possibilities. For a concrete example consider the domain  $\mathcal{O} = (0, 1)^d$ .

It can be verified that the eigenfunctions of the inverse Laplacian on  $\mathcal{O} = (0, 1)^d$  with Dirichlet boundary condition are given by

$$h_{i_1, \dots, i_d}(x_1, \dots, x_d) = 2^{d/2} \prod_{j=1}^d \sin(i_j \pi x_j), \quad (i_1, \dots, i_d) \in \mathbb{N}^d, \quad (2.22)$$

with corresponding eigenvalues

$$\kappa_{i_1, \dots, i_d} = \frac{1}{(\sum_{j=1}^d i_j^2) \pi^2}.$$

The indexing by the  $d$ -tuples  $(i_1, \dots, i_d)$  precludes an immediate application of contraction results for sequence priors. However the sequence  $k_1 \geq k_2 \geq \dots$  resulting from ordering the array of eigenvalues  $(\kappa_i : i \in \mathbb{N}^d)$  in decreasing magnitude has polynomial decrease  $k_\ell \asymp \ell^{-2/d}$ , as  $\ell \rightarrow \infty$  (see Lemma 2.B.2 in Section 2.B). Furthermore, if  $v_1, v_2, \dots$  are the array of coefficients  $(\nu_i : i \in \mathbb{N}^d)$  in corresponding order of a function  $v = \sum_{i \in \mathbb{N}^d} \nu_i h_i$ , then the square norm in Equation (2.10) of the smoothness scale  $(S^s)_{s \in \mathbb{R}}$  generated by the basis  $(h_i)$  ordered in a sequence satisfies, for  $s \geq 0$ ,

$$\|v\|_{S^s}^2 = \sum_{\ell=1}^{\infty} v_\ell^2 \ell^{2s/d} \asymp \sum_{i \in \mathbb{N}^d} \nu_i^2 \left( \sum_{j=1}^d i_j^2 \right)^s.$$

Although the eigenfunctions are not always explicit, similar results are valid for the Laplacian on more general domains; see Equation (3.1) in (Ivrii 2016).

The next lemma verifies the interpolation assumptions of Theorem 2.3.4 (see Section 2.B for a proof).

**Lemma 2.4.4.** *Let  $(S^s)$  be the scale generated by the eigenfunctions from Equation (2.22) and let  $\mathbb{L}_N := \{\sum_{i \in \mathbb{N}^d, \|i\|_\infty \leq N^{1/d}} c_i h_i : c_i \in \mathbb{R}\}$ . If  $\beta > d/2$ , then for every  $v \in S^\beta$ , there exists an element  $\mathcal{I}_N v \in \mathbb{L}_N$  such that  $\mathcal{I}_N v(x_i) = v(x_i)$  for every  $i = 1, \dots, N$ , and*

$$\|\mathcal{I}_N v - v\|_{L^2} \lesssim N^{-\beta/d} \|v\|_{S^\beta}. \quad (2.23)$$

Furthermore, there exist constants  $0 < C_1 < C_2 < \infty$  such that

$$C_1 \|v\|_{L^2} \leq \|v\|_N \leq C_2 \|v\|_{L^2}, \quad \forall v \in \mathbb{L}_N.$$

Given the eigenfunctions  $(h_i)$ , for  $i \in \mathbb{N}^d$ , equip  $v = \Delta u_f$  with the prior  $\sum_{i=1}^{\infty} \nu_i h_i$ , for  $\nu_i \stackrel{\text{ind}}{\sim} N(0, \kappa_i^{d/2+\alpha})$ , where the value of  $\alpha$  is chosen either fixed or determined by the empirical Bayes or hierarchical Bayes approaches. For the array  $(h_i)$  ordered in by decreasing eigenvalues, this corresponds to a prior of the form  $v = \sum_{\ell=1}^{\infty} v_\ell h_\ell$ , with  $v_\ell \stackrel{\text{ind}}{\sim} N(0, \ell^{-1-2\alpha/d})$ , as in Section 2.3. To quantify the uncertainty of the procedure, consider credible sets of the form

$$C_N(Y_N) = \left\{ \frac{v}{2(Kv + \tilde{g})} : \|v - \hat{v}_N\|_{L^2} \leq r_N, \|Kv - K\hat{v}_N\|_\infty \leq \delta_0 \right\}, \quad (2.24)$$

for  $\hat{v}_N$  the posterior mean of  $v$ , the constant  $r_N$  determined as the  $1 - \gamma$ -quantile of the posterior distribution of  $\|v - \hat{v}_N\|_{L^2}$  and a fixed constant  $\delta_0 > 0$ .

**Theorem 2.4.5.** *Assume that  $\Delta u_{f_0} \in S^\beta$ , for  $\beta > d/2$  and  $\inf_x u_{f_0}(x) > 2\delta_0 > 0$ .*

- (i) *The  $L^2$ -contraction rate of the posterior distribution of  $f$  to  $f_0$  in the white noise model is  $N^{-(\alpha \wedge \beta)/(d+2\alpha+4)}$  if the regularity of the prior is chosen deterministic and equal to  $\alpha$  with  $\alpha \wedge \beta + 2 > d/2$ .*
- (ii) *The  $L^2$ -posterior contraction rate is  $l_N N^{-\beta/(d+2\beta+4)}$  for a slowly varying sequence  $l_N$ , if  $\alpha$  is chosen by the empirical Bayes or hierarchical Bayes methods and  $\beta + 2 > d/2$ .*
- (iii) *For the prior with deterministic  $\alpha$  set to  $\alpha = \beta - c/\log N$  and sufficiently large  $c$ , the credible set  $C_N(Y_N)$  has frequentist coverage tending to one and  $L^2$ -diameter of the order  $O_P(N^{-\beta/(d+2\beta+4)})$ .*
- (iv) *For the prior with fixed  $\alpha > d/2$ , the same contraction rate is attained in the discrete observational scheme.*

*Proof.* The contraction rates (i) and (ii) in the Gaussian white noise model follow from combining Theorem 2.4.3 with Theorem 2.3.2. The frequentist coverage guarantee (iii) follows in the same way from combining Theorem 2.3.2 with Proposition 2.2.2. The contraction rate (iv) in the discrete observational model follows from combining Lemma 2.4.4 and Theorem 2.3.4.  $\square$

As discussed in Section 2.3.3, similar results are valid for other priors than those resulting from the eigenexpansion, for instance, priors based on certain wavelet expansions.

A sufficient condition for the assumption that the function  $\Delta u_{f_0}$  belongs to the Hilbert scale  $S^\beta$  is that it is compactly supported in  $\mathcal{O}$  and belongs to the Sobolev space  $H^\beta$ . See Proposition 2.A.1 and Chapter 5, Appendix A (page 471), in particular Equation (A.18) in (Taylor 2011).

## 2.5 Heat equation with absorption

For given (smooth) functions  $u_0 \in L^2([0, 1]^d)$  and  $g \in L^2(\partial[0, 1]^d \times [0, 1])$ , consider the solution  $u_f : [0, 1]^d \times [0, 1] \rightarrow \mathbb{R}$  to the parabolic partial differential equation

$$\begin{cases} \frac{\partial}{\partial t} u - \frac{1}{2} \Delta u = f u & \text{on } (0, 1)^d \times (0, 1], \\ u = g & \text{on } \partial[0, 1]^d \times [0, 1], \\ u = u_0 & \text{on } (0, 1)^d \times \{0\}. \end{cases} \quad (2.25)$$

It is assumed that  $u_0$  and  $g$  satisfy the compatibility condition  $g(x, 0) = u_0(x)$  for all  $x \in \partial[0, 1]^d$ . The white noise version of this problem was investigated in (Kekkonen 2022), where minimax contraction rates and frequentist guarantees for credible sets in a weak metric space were derived under the assumption that the regularity  $\beta$  of the function  $f \in H^\beta$  is known.

The equation is of the form (2.1) for the parabolic differential operator  $\mathcal{L} = \frac{\partial}{\partial t} - \frac{1}{2}\Delta$ , with  $\Gamma = (\partial[0, 1]^d \times [0, 1]) \cup ((0, 1)^d \times \{0\})$ . The function  $f$  can be recovered from the solution  $u_f$  to the equation as

$$f = \frac{\mathcal{L}u_f}{u_f}.$$

The existence and uniqueness of the inverse operator  $K \in L^2([0, 1]^d \times [0, 1])$  satisfying Equation (2.4) is proved in Theorem 3 and Theorem 4 of Section 7.1 in (Evans 2010) and the existence of the function  $\tilde{g}$  satisfying Equation (2.5) follows from Section 3 in (Doob 1955). This gives the solution map  $f = e(\mathcal{L}u_f)$  as in Equation (2.2) given by

$$e(v) = \begin{cases} \frac{v}{Kv + \tilde{g}} & \text{if } \text{ess inf } Kv + \tilde{g} > 0, \\ 0 & \text{otherwise.} \end{cases}$$

This is similar to the Schrödinger equation in Section 2.4, except for the different definition of  $K$ . Lemma 2.4 is true in exactly the same form for the present operator  $K$  and solution map.

In view of Proposition 2.2.1, an  $L^2$ -contraction rate of the posterior distribution of  $v$  to  $v_0$  based on the observation  $\tilde{Y}_N = Kv + N^{-\frac{1}{2}}\mathbb{W}$  implies the same  $L^2$ -contraction rate for the posterior distribution of  $f$ , provided the posterior distribution of  $Kv$  is consistent for  $Kv_0$  relative to the uniform norm and  $u_{f_0} = Kv_0 + \tilde{g}$  is bounded away from zero. Proposition 2.2.2 shows that the coverage and diameter of credible sets translate similarly. This is exactly as in Theorem 2.4.3, and applicable to many priors.

For a concrete example, consider the prior on the eigenfunctions of the operator  $K^\top K$ , which can be derived in closed form (see Lemma 2.C.1). We may endow  $v = \mathcal{L}u_f$  with a Gaussian prior  $\sum_{i \in \mathbb{N}} v_i h_i$ , for  $h_1, h_2, \dots$  the eigenfunctions ordered by decreasing eigenvalues and  $v_i \stackrel{\text{ind}}{\sim} N(0, i^{-1-\alpha/d})$ , with the hyperparameter  $\alpha$  either fixed to some value or selected with the empirical or hierarchical Bayes methods, as described in Section 2.3. For uncertainty quantification, we consider the credible set described in Equation (2.24).

Let  $(S^s)$  be the smoothness scale corresponding to the sequence of eigenfunctions, ordered by corresponding decreasing eigenvalues.



**Theorem 2.5.1.** *Assume that  $\mathcal{L}u_{f_0} \in S^\beta$ , for  $\beta > d + 1$ , and  $\inf_x u_{f_0}(x) > 2\delta_0 > 0$ .*

- (i) *In the white noise model the  $L^2$ -contraction rate of the posterior distribution for  $f$  to  $f_0$  is equal to  $N^{-(\alpha \wedge \beta)/((d+1)+2\alpha+4)}$  if the regularity of the prior is chosen deterministic and equal to  $\alpha > d+1$ , and equal to  $l_N N^{-\beta/((d+1)+2\beta+4)}$  for a slowly varying sequence  $l_N$  if  $\alpha$  is chosen by the hierarchical or empirical Bayes methods.*
- (ii) *For the prior with deterministic  $\alpha$  set to  $\alpha = \beta - c/\log N$  and sufficiently large  $c$ , the credible set in Equation (2.24) has frequentist coverage tending to one and  $L^2$ -diameter of the order  $N^{-\beta/((d+1)+2\beta+4)}$ .*

*Proof.* In view of Lemma 2.B.2 and Lemma 2.C.1, the eigenvalues  $k_\ell^2$  of  $K^\top K$  ordered by decreasing magnitude satisfy  $k_\ell \asymp \ell^{-2/d}$ . Thus the operator  $K$  is smoothing as in Equation (2.11) of order  $p = 2$  relative to the eigenscale  $(S^s)$ .

Consequently, Theorem 2.3.2 gives  $L^2$ -contraction at the rates as given in (i), and  $\|\cdot\|_{S^\gamma}$ -contraction, for every  $\gamma < \beta$ .

By Proposition 2.2.1 the  $L^2$ -contraction rate of the posterior of  $v$  is transferred to the same rate for  $f$  provided that the posterior distribution of  $Kv$  is consistent for the uniform norm. By Lemma 2.3.3, the latter follows from consistency in the  $\|\cdot\|_{S^\gamma}$ -norm for some  $\gamma > d + 1$ .

The frequentist coverage (ii) of the credible sets follows similarly with the help of Proposition 2.2.2.  $\square$

## 2.6 One-dimensional Darcy equation

For given functions  $h : (0, 1) \rightarrow \mathbb{R}$  and  $g : \{0, 1\} \rightarrow \mathbb{R}$ , consider the one-dimensional differential equation

$$\begin{cases} f'u_f + fu_f'' = h & \text{on } (0, 1), \\ u_f = g & \text{on } \{0, 1\}. \end{cases} \quad (2.26)$$

Because the solutions  $f$  to the first order linear ordinary differential equation  $f'u_f + fu_f'' = h$ , for given  $u$  (and hence  $u'$  and  $u''$ ), form a one-dimensional affine space, and we can at best retrieve  $u_f$  from the data, the function  $f$  can be recovered from the data at best up to an additive constant. We shall assume that the value  $f(0)$  is known.

The equation can be written in the form  $(fu'_f)' = h$ . Integration of this equation gives that  $fu'_f - f(0)u'_f(0) = H(x)$ , for  $H(x) = \int_0^x h(s) ds$ . Under the assumption that  $\inf_{x \in [0,1]} u'_f(x) > 0$ , we can solve  $f$  as

$$f(x) = \frac{H(x) + f(0)u'_f(0)}{u'_f(x)}. \quad (2.27)$$

We conclude that we can write the equation in the form as in Equation (2.1) with  $\mathcal{L}u = u'$  (and  $c(f, u_f) = (f(0)u'_f(0) + H)/f$ ) and a solution map in Equation (2.2) as a map from  $\mathcal{L}u_f = u'_f$  to  $f$ .

Let  $\tilde{g}$  be the function  $\tilde{g}(x) = g(0) + (g(1) - g(0))x$ , and define the operator  $K$  by

$$Kv(x) = \int_0^x v(s) ds - x \int_0^1 v(s) ds.$$

Then  $Kv(0) = Kv(1) = 0$ , for every  $v$ , and hence  $Ku' + \tilde{g} = g$  on  $\partial O$  and  $(Ku' + \tilde{g})'' = u''$  on  $O = (0, 1)$ . This implies that  $u_f = K\mathcal{L}u_f + \tilde{g}$ , verifying the identity in Equation (2.6). (The two boundary conditions of the non-homogeneous problem (2.26) cannot be removed by the general scheme (2.4)–(2.5) explained in the introduction when using the first order differential operator  $\mathcal{L}u = u'$ . Posing the problem instead in terms of the second order operator  $\mathcal{L}_0 u = -u''$ , we would find the same function  $\tilde{g}$  and the inverse operator  $-KK^\top$ . This would give  $u_f = -KK^\top u''_f + \tilde{g} = Ku'_f + \tilde{g}$ , and hence lead to the inverse equation as introduced.)

The adjoint of  $K$  is given by  $K^\top u = -\int_0^1 u(s) ds + \int_0^1 \int_0^t u(s) ds dt$ , and the eigenfunctions and corresponding eigenvalues of  $K^\top K$  can be seen to be

$$h_0 = 1, \quad \kappa_0^2 = 0,$$

and

$$h_i(x) = \sqrt{2} \cos(\pi i x), \quad i \in \mathbb{N}, \quad \kappa_i^2 = \frac{1}{(\pi i)^2}.$$

The orthogonality  $\int_0^1 K^\top u(s) ds = 0$  of the functions  $K^\top u$  to  $h_0$  motivates to consider the prior on  $v = u'_f$  equal to the distribution of  $\sum_{j \in \mathbb{N}} v_j h_j$ , for independent  $v_j \sim N(0, j^{-1-2\alpha})$ . Define  $(S^s)_{s \in \mathbb{R}}$  as the scale generated by the eigenexpansion of  $K^\top K$ .

**Theorem 2.6.1.** *Suppose  $\|H\|_\infty < \infty$ ,  $u'_{f_0} \in S^\beta$ , for  $\beta > \frac{1}{2}$ , and  $u_{f_0}$  satisfies  $\inf_{0 < x < 1} u'_{f_0}(x) > 0$ .*

- (i) For any  $\alpha \wedge \beta > \delta > \frac{1}{2}$  the contraction rate of the posterior distribution of  $f$  to  $f_0$  relative to the  $\|\cdot\|_\infty$ -norm is  $N^{-(\alpha \wedge \beta - \delta)/(1+2\alpha+2)}$ .
- (ii) If  $f(0) = 0$  and  $\alpha \wedge \beta > \frac{1}{2}$ , then the contraction rate of the posterior distribution of  $f$  to  $f_0$  relative to the  $L^2$ -norm is  $N^{-(\alpha \wedge \beta)/(1+2\alpha+2)}$ .

*Proof.* The solution map (2.2) is given by  $e(v) = (H + f(0)v(0))/v$ . On the set of functions  $V_N = \{v : v \geq c\}$ , for fixed  $c > 0$ , it satisfies

$$\begin{aligned} |e(v) - e(v_0)| &\leq \frac{f(0)|v(0) - v_0(0)|}{c} + (\|H\|_\infty + f(0)v_0(0)) \frac{|v_0 - v|}{c} \\ &\lesssim \|v - v_0\|_\infty. \end{aligned}$$

Since  $v_0 = u'_{f_0}$  is bounded away from zero by assumption, the posterior distribution of  $v$  will concentrate on  $V_N$  as soon as it is consistent for the uniform norm. Thus in view of Proposition 2.2.1, for assertion (i) it suffices to obtain a contraction rate for  $v$  relative to the uniform norm.

By Theorem 2.3.2 the posterior contraction rate for  $v$  relative to the  $\|\cdot\|_{S^\delta}$ -norm is  $N^{-(\alpha \wedge \beta - \delta)/(1+2\alpha+2)}$ . By assumption  $u_{f_0}(0) = u_{f_0}(1) = 0$ , whence  $v_0 = u'_{f_0}$  satisfies  $\langle v_0, 1 \rangle_{L^2} = 0$ . By construction of the prior  $\langle v, 1 \rangle_{L^2} = 0$  for any  $v$  in a set of posterior probability one. It follows that on this set

$$\|v - v_0\|_\infty = \sup_{0 < x < 1} \left| \sum_{j \in \mathbb{N}} (v_j - v_{0,j}) h_j(x) \right| \leq \|v - v_0\|_{S^\delta} \sup_{0 < x < 1} \sqrt{\sum_{j \in \mathbb{N}} \frac{h_j(x)^2}{j^{2\delta}}}.$$

Since the functions  $h_j$  are uniformly bounded, the right side is bounded by a universal multiple of  $\|v - v_0\|_{S^\delta}$ , for any fixed  $\delta > \frac{1}{2}$ . Thus the posterior contraction relative to the  $\|v - v_0\|_{S^\delta}$ -norm implies a posterior contraction relative to the uniform norm at the same rate.

To prove assertion (ii), we first note that if  $f(0) = 0$ , then the map  $e$  is Lipschitz relative to the  $L^2$ -norm on  $V_N$ . By Theorem 2.3.2 the posterior contraction rate for  $v$  relative to the  $L^2$ -norm is  $N^{-(\alpha \wedge \beta)/(1+2\alpha+2)}$ . By Proposition 2.2.1 this carries over to a posterior contraction rate for  $f$  provided that the posterior probability of the sets  $V_N$  tends to one. Under the assumption that  $v_0 = u'_{f_0}$  is bounded away from zero, the latter follows from posterior consistency of  $v$  for the uniform norm. By the argument of the preceding paragraph this follows from consistency relative to the  $\|\cdot\|_{S^\delta}$ -norm, for some  $\delta > \frac{1}{2}$ , and this is true for  $\alpha \wedge \beta > \frac{1}{2}$ .  $\square$

Other boundary conditions than the Dirichlet boundary condition in Equation (2.26) may motivate different priors. For instance, for the mixed boundary condition  $u_f''(0) = u_f'(1) = 0$ , we may use the same differential operator  $\mathcal{L}u = u'$ , but take the standard Volterra operator  $Kv = \int_0^{\cdot} v(s) ds$  as its inverse. The eigenfunctions of  $K^\top K$  satisfy  $\kappa^2 v'' = v$  with  $v'(0) = v(1) = 0$ , are the functions  $h_i(x) = \sqrt{2} \cos((i + \frac{1}{2})\pi x)$  with corresponding eigenvalues  $\kappa_i^2 = ((i + \frac{1}{2})\pi)^{-2}$ , for  $i \in \mathbb{N} \cup \{0\}$ . The functions in the eigenscale  $S^s$  generated by the operator  $K^\top K$  satisfy the boundary condition on  $u_f'$ , and the preceding corollary remains valid provided  $u_{f_0}' \in S^\beta$ .

## 2.7 Exponentiated Volterra operator

For a known positive constant  $g$ , consider the solution  $u_f : [0, 1] \rightarrow \mathbb{R}$  of the ordinary differential equation

$$\begin{cases} u_f' = fu_f & \text{on } (0, 1), \\ u_f = g & \text{at } 0. \end{cases} \quad (2.28)$$

This equation takes the form as in Equation (2.1) with  $\mathcal{O} = (0, 1)$  and  $\Gamma = \{0\}$ , differential operator  $\mathcal{L}u = u'$  and  $c(u, f) = uf$ . The inverse operator  $K$  defined in Equation (2.4) is the standard Volterra operator  $K : L^2([0, 1]) \rightarrow L^2([0, 1])$  defined by

$$Kv(x) := \int_0^x v(s) ds.$$

The function  $\tilde{g}$  as in Equation (2.5) is the constant  $\tilde{g} = g$ . It is immediate from the differential Equation (2.28) that  $f = u_f'/u_f$ , where we note that the function  $u_f$  is nonzero, as follows from the explicit expression of the solution:  $u_f(x) = g \exp(\int_0^x f(s) ds)$ . Thus the solution map in Equation (2.2) is given by

$$e(v) = \frac{v}{Kv + g}. \quad (2.29)$$

Because  $\|Kv\|_{L^2} \leq \|Kv\|_\infty \lesssim \|v\|_{L^2}$ , the map  $e : L^2([0, 1]) \rightarrow L^2([0, 1])$  is Lipschitz with respect to the  $L^2$ -norm in a sufficiently small neighbourhood  $V = \{v : \|v - v_0\|_{L^2} < c_0\}$  of any function  $v_0$  such that  $Kv_0 + g$  is bounded away from zero. Therefore  $L^2$ -contraction rates for  $v = u_f'$  in the linear Volterra inverse problem translate immediately into  $L^2$ -contraction rates for  $f$ .

One possible prior is the one based on the eigenexpansion. The eigenfunctions of the operator  $K^\top K : L^2(0, 1) \rightarrow L^2(0, 1)$  take the form

$$h_i(x) = \sqrt{2} \cos((i - 1/2)\pi x), \quad i \in \mathbb{N},$$

with corresponding eigenvalues

$$\kappa_i^2 = \frac{1}{\pi^2(i - \frac{1}{2})^2}.$$

Hence the operator  $K$  is smoothing as in Equation (2.11) of order  $p = 1$  in its eigenscale  $(S^\beta)_{\beta \in \mathbb{R}}$ .

We endow  $v = u'_f$  with the Gaussian prior distribution  $v = \sum_{i=1}^{\infty} v_i h_i$ , for  $v_i \stackrel{\text{ind}}{\sim} N(0, i^{-1-2\alpha})$ , for fixed  $\alpha$  or with  $\alpha$  determined by the empirical or hierarchical Bayes method.

**Theorem 2.7.1.** *If  $u'_{f_0} \in S^\beta$ , then the posterior distribution of  $f$  contracts in the  $L^2$ -norm at rate  $\epsilon_N = N^{-(\alpha \wedge \beta)/(1+2\alpha+2)}$  if  $\alpha$  is chosen fixed and at the rate  $\epsilon_N = l_N N^{-\beta/(1+2\beta+2)}$  for a slowly varying sequence  $l_N$  if  $\alpha$  is chosen by the empirical or hierarchical Bayes methods. Furthermore, the credible sets from Equation (2.8) with  $\tilde{C}_N(\tilde{Y}_N)$  the credible ball in Equation (2.15) possess frequentist coverage tending to one and diameter of the order  $\epsilon_N$ .*

*Proof.* Because  $Ku'_{f_0} + g$  is bounded away from zero, the map  $e : L^2([0, 1]) \rightarrow L^2([0, 1])$  is Lipschitz with respect to the  $L^2$ -norm. Therefore, the results are an immediate consequence of Theorem 2.3.2 and Propositions 2.2.1 and 2.2.2.  $\square$

## 2.8 Numerical analysis

We demonstrate the practical performance of our approach by applying it to the time-independent Schrödinger equation (2.19), discussed in Section 2.4. We consider synthetic data sets generated from various choices of a true function  $f_0$  in the Gaussian white noise observational model (2.3) and investigate the behaviour of the posterior distributions corresponding to different choices of the prior.

First, we consider the one-dimensional case ( $d = 1$ ) with the domain equal to the unit interval. We set the boundary conditions for  $u_f$  as  $g(0) = 1$  and  $g(1) = 2$ , so that the function  $\tilde{g}$  in Equation (2.5) is given by  $\tilde{g}(x) = 1 + x$ . The results are shown in Figures 2.1 and 2.2. Each subpanel of the figures illustrates the posterior distribution corresponding to a single simulation of data. From top

to bottom, the four panels correspond to signal-to-noise ratios  $N = 10^4, 10^6, 10^8$ , and  $10^{10}$ . From left to right the panels in the two figures correspond to different functions  $f_0$  and/or different prior distributions.

For each plot, we took 500 draws from the posterior distribution of the Laplacian  $v = \Delta u_f$ , and transformed these into draws of  $f$  using the solution operator from Equation (2.21). We approximated the posterior mean of  $f$  as the average of these 500 draws (the red curve), as well as pointwise 2.5% and 97.5% quantiles to form credible bands (the green curves). In each panel, we also plot 20 draws selected randomly out of the 500 draws to visualize typical samples from the posterior (shown as dashed grey curves).

In Figure 2.1 the true function (drawn in black) was taken equal to  $f_0(x) = \sum_{i=1}^{\infty} i^{-3/2} \sin(i) h_i$  in the left four panels, and the true function was equal to  $f_0(x) = 1 - 4(x - \frac{1}{2})^2 - \frac{3}{4} \exp(-500(x - \frac{1}{2})^2)$  in the right four panels. The first function is in  $S^\beta$ , for all  $\beta < 1$ . The prior was based on the eigenexpansion of the Laplacian with Gaussian coefficients from Equation (2.9) and regularity hyperparameter  $\alpha$  set with the empirical Bayes approach by maximising the marginal likelihood function given in Equation (2.14). For the first function, the approach works well: the posterior mean provides an accurate estimate of the true function and the credible band contains the true function  $f_0$ , thus resulting in reliable uncertainty quantification from the frequentist perspective. For the second function, the performance is less satisfying: a larger sample size (smaller signal-to-noise ratio) is needed to detect the bump that is present in this function. However, this behaviour appears to be less connected to the method as to the problem: it is hard to detect local spikes, especially if the regularity of the function  $f_0$  is unknown. Already the direct estimation problem using Gaussian process priors of the form (2.9) needs large sample sizes.

The two true functions in Figure 2.1 vanish at the boundary, which agrees with the prior set on the singular basis of the inverse Laplace operator. In Figure 2.2 we investigate the function  $f_0(x) = 1 + \sin(3\pi x)$ , which takes the value 1 at both boundary points. Our theoretical results show that even for this function our posterior contracts around the truth with the optimal  $L^2$ -norm and the  $L^2$ -credible ball provides high frequentist coverage. However, the pointwise behaviour at the boundary using the singular value decomposition basis fails, as the posterior is supported on functions vanishing at the boundary. This is demonstrated in the four panels in the left column of Figure 2.2, which is based on the same prior as previously. The posterior distribution shows good performance except near the boundary. To overcome the boundary problem we also implemented a prior based on the B-spline basis, which does not restrict the boundary. The resulting posterior distribution is shown in the right column of

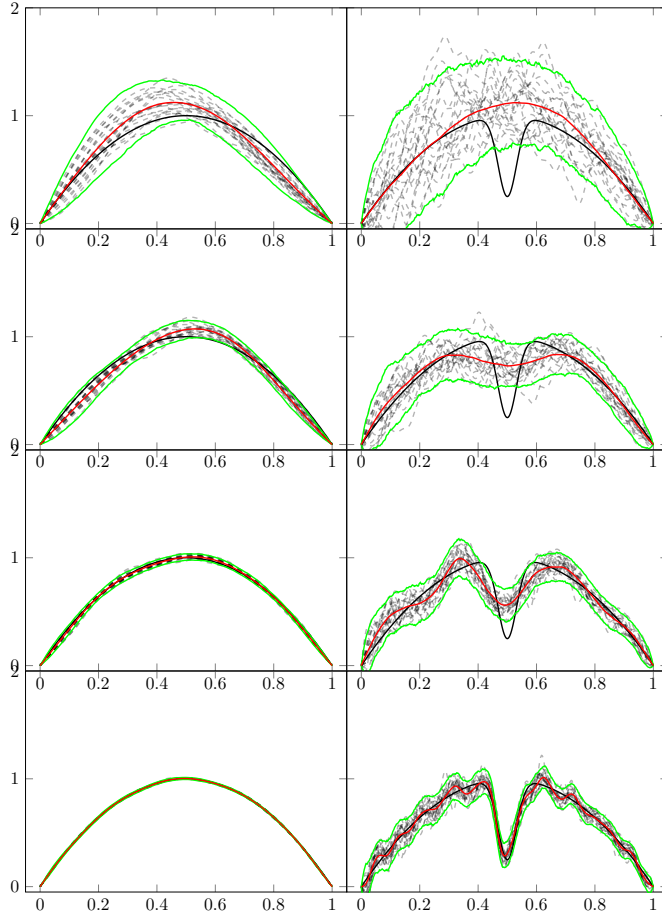


Figure 2.1: Posterior distributions of  $f$  based on data  $Y_N = u_f + N^{-\frac{1}{2}}\mathbb{W}$  for  $N = 10^4, 10^6, 10^8, 10^{10}$  (top to bottom) and prior based on eigenbasis. Black curve: true function  $f_0$ . Red curve: posterior mean. Green curves: pointwise 95% credible band. Dashed black curves: 20 draws from the posterior distribution. Left panels:  $f_0(x) = 1 + 4(x - \frac{1}{2})^2$ , right panels:  $f_0(x) = 1 - 4(x - \frac{1}{2})^2 - \frac{3}{4} \exp(-500(x - \frac{1}{2})^2)$ .

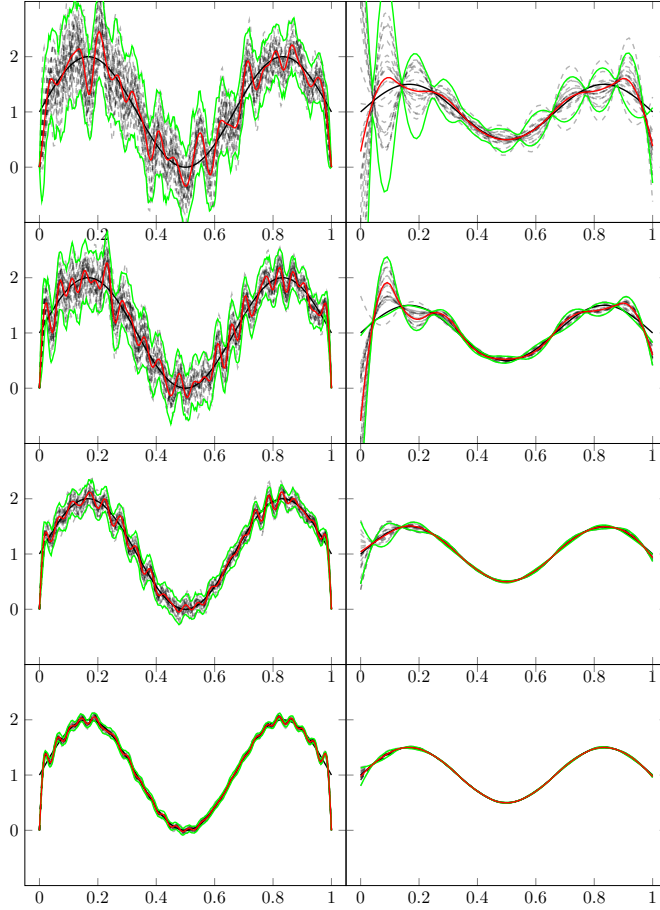


Figure 2.2: Posterior distributions of  $f$  based on data  $Y_N = u_f + N^{-\frac{1}{2}}\mathbb{W}$  for  $N = 10^4, 10^6, 10^8, 10^{10}$  (top to bottom). Black curve: true function  $f_0(x) = 1 + \sin(3\pi x)$ . Red curve: posterior mean. Green curves: pointwise 95% credible band. Dashed black curves: 20 draws from the posterior distribution. Left panels: prior based on eigenbasis. Right panels: prior based on spline basis.

Figure 2.2. Both the estimation and uncertainty quantification are more accurate at the boundary.

Next, we demonstrate the behaviour of our approach in a two-dimensional example, with true function and boundary condition on the domain  $(0, 1)^2$  given by

$$\begin{aligned} f_0(x, y) &= 2x(x-1)y(y-1)(2 + \sin(3\pi x) \sin(\pi y)), \\ g(x, y) &= 3 + xy^2 + 2y \sin(2\pi x) + x \cos(3\pi y), \end{aligned} \quad (2.30)$$



respectively. The four panels in the first column of Figure 2.3 all show a heat map of the function  $f_0$ . The other columns show heat maps of different aspects of the posterior distribution, each time based on single data obtained using increasing signal-to-noise ratio  $N = 10^4, 10^6, 10^8$  and  $10^{10}$ , from top to bottom: the posterior mean (second column), lower and upper 2.5% pointwise quantiles of the posterior (third and fourth column), and the absolute difference between the posterior mean and the true function (fifth column). The prior was based on the eigenbasis given in Equation (2.22) of the Laplacian with Gaussian coefficients from Equation (2.9) and regularity  $\alpha$  set by the empirical Bayes method. One can observe that similarly to Figure 2.1, we get good approximation and reliable uncertainty quantification, in this case even pointwise as  $f_0$  vanishes at the boundary.

## 2.9 Discussion

We have derived a novel, Bayesian linearisation method for non-linear inverse problems. This approach potentially decreases computation time, as routines for solving linear inverse problems are simpler and faster than for non-linear ones. Even if the linear inverse problem is non-conjugate (e.g. if the prior is not set on the singular value basis of  $K$ ) and iterative methods are needed, our method may be substantially faster as it does not require solving the forward problem at each iteration step.

We derived posterior contraction rates for Gaussian process priors, with both fixed and data-driven regularity hyperparameters (using either the empirical or hierarchical Bayes methods). These theorems on adaptation are the first results on PDE-constrained, non-linear inverse problems. We also showed that the frequentist coverage guarantees of credible sets in the corresponding linear inverse problems translate to the non-linear inverse problems. The results apply both to the benchmark Gaussian white noise model and to the practically more relevant discrete observational framework. These findings can be extended to priors going beyond the eigenfunctions of the differential operator  $\mathcal{L}$ . The numerical analysis with synthetic data sets shows that the proposed approach has good numerical properties, although proper handling of boundary conditions appears to require care.

A potential bottleneck of our approach is the solution operator in Equation (2.2). In this chapter, as a proof of concept, we cover a collection of benchmark PDE and ODE problems, where this operator takes a simple, explicit form. In general, a numerical implementation of the operator may be necessary. We

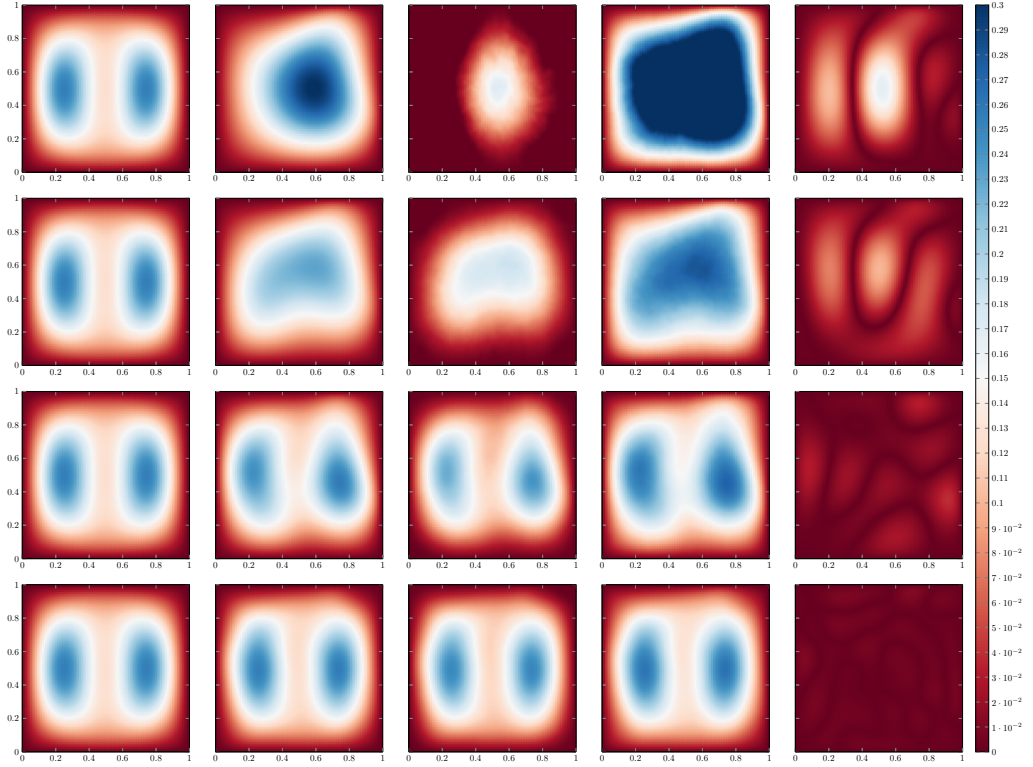


Figure 2.3: The true function  $f_0$  given in Equation (2.30), the posterior mean, the lower and upper 2.5% pointwise posterior quantiles, and the absolute approximation error of the posterior mean, resulting from the proposed Bayesian approach with signal-to-noise ratios  $N = 10^4, 10^6, 10^8, 10^{10}$  in the PDE constrained Gaussian white noise model (2.1)–(2.3).

can borrow here from the field of (numerical) analysis, where this is an active research area. Another interesting question for future research is scaling up our approach with various machine learning techniques, including distributed and variational inference. This linearised setup provides a convenient framework to derive such approximation algorithms compared to the original non-linear setting.



# Appendix

## 2.A Hilbert scales and Sobolev spaces

Let  $\mathcal{L}$  be a symmetric, second order, elliptic differential operator with  $C^\infty$  coefficient functions such that the function  $v = 0$  is the only solution to the problem  $\mathcal{L}v = 0$  on  $\mathcal{O}$  and  $v = 0$  on  $\partial\mathcal{O}$ . For  $u \in L^2$ , let  $Ku \in H_0^1$  be the weak solution to

$$\begin{cases} \mathcal{L}Ku = u & \text{on } \mathcal{O}, \\ Ku = 0 & \text{on } \partial\mathcal{O}. \end{cases}$$

Then  $K : L^2 \rightarrow L^2$  is self-adjoint and can be seen to be compact (it is bounded as a function  $K : L^2 \rightarrow H_0^1$  due to Theorem 6 in Section 6.2 in (Evans 2010), and since  $H_0^1$  is compactly embedded in  $L^2$  due to Theorem 7.22 in (Gilbarg and Trudinger 2015), it follows that  $K : L^2 \rightarrow L^2$  must be a compact operator).

Given its eigenfunctions  $(h_j)_{j \in \mathbb{N}}$  and eigenvalues  $\kappa_j$  and  $s \geq 0$ , let  $S^s$  be the set of functions  $v = \sum_{j=1}^\infty v_j h_j \in L^2$  for which  $\|v\|_{S^s}^2 := \sum_{j=1}^\infty v_j^2 \kappa_j^{-s}$  is finite. For  $s < 0$ , let  $S^s$  be the dual space of  $S^{-s}$ .

Recall that  $H^s$  is the Sobolev space of order  $s$ .

**Proposition 2.A.1.** *For  $s \geq 0$  such that  $\partial\mathcal{O} \in C^{\lceil s \rceil + 2}$ , we have  $S^s \subset H^s$ , with  $\|u\|_{H^s} \lesssim \|u\|_{S^s}$ , for every  $u \in S^s$ , with  $\|u\|_{H^s} \asymp \|u\|_{S^s}$  for every even integer  $s$ . For  $s = 2$  we have  $S^2 = H^2 \cap H_0^1$  if  $\partial\mathcal{O} \in C^2$ .*

*Proof.* The space  $S^s$  is equivalent to the set of functions  $K^s v$ , for  $v \in L^2$ , with the norm  $\|K^s v\|_{S^s} = \|v\|_{L^2}$ .

For every  $m \in \mathbb{N} \cup \{0\}$ , under the condition that  $\partial\mathcal{O} \in C^{m+2}$ , Theorem 5 in (Evans 2010) gives that  $Ku \in H^{m+2}$  if  $u \in H^m$ , with  $\|Ku\|_{H^{m+2}} \lesssim \|u\|_{H^m}$ .

For  $m = 0$ , this shows that  $S^2 = KL^2 \subset H^2$  and  $\|Ku\|_{H^2} \lesssim \|u\|_{L^2} \asymp \|Ku\|_{S^2}$ . By the definition of a weak solution with Dirichlet boundary condition, it is also true that  $KL^2 \subset H_0^1$  and hence  $S^2 \subset H^2 \cap H_0^1$  with  $\|u\|_{H^2} \lesssim \|u\|_{S^2}$ . Conversely if  $u \in H^2$ , then  $\mathcal{L}u \in L^2$  by the definitions of  $H^2$  and  $\mathcal{L}$  (or see (5.12))

in Chapter 4.5 of (Taylor 2011)). If also  $u \in H_0^1$ , then  $w = u$  solves  $\mathcal{L}w = \mathcal{L}u$  (trivially) on  $\mathcal{O}$  and  $w = 0$  on  $\partial\mathcal{O}$  and hence  $u = K\mathcal{L}u \in S^2$ , with norm  $\|u\|_{S^2} = \|\mathcal{L}u\|_{L^2} \lesssim \|u\|_{H^2}$ . This concludes the proof that  $S^2 = H^2 \cap H_0^1$  and that the norms are equivalent on this space.

Next, we prove by induction that  $S^{2m} \subset H^{2m}$  with equivalent norms, for every  $m \in \mathbb{N}$ . For  $m = 1$  the assertion is proved. From the definitions it follows that  $S^{2m+2} = KS^{2m}$  with norms satisfying  $\|K^{m+1}u\|_{S^{2m+2}} = \|u\|_{L^2} = \|K^m u\|_{S^{2m}}$ . By Theorem 5 in (Evans 2010) as quoted previously  $K^{m+1}u = K(K^m u)$  is contained in  $H^{2m+2}$  with norms satisfying the bound  $\|K^{m+1}u\|_{H^{2m+2}} \lesssim \|K^m u\|_{H^{2m}}$ . Therefore if the assertion is true for  $m$ , then it is true for  $m + 1$ .

The Sobolev spaces  $H^s$  for non-integer  $s \geq 0$  can be constructed as the interpolation spaces  $[L^2, H^m]_{s/m, 2}$ , for  $s \leq m \in \mathbb{N}$ . This means that the norm of  $u \in H^s$  is equivalent to

$$\left( \int_0^\infty \left( \frac{1}{t^{s/m}} \inf_{u=u_0+u_m} (\|u_0\|_{L^2} + t\|u_m\|_{H^m}) \right)^2 \frac{dt}{t} \right)^{\frac{1}{2}},$$

where the infimum is taken over all decompositions  $u = u_0 + u_m$  with  $u_0 \in L^2$  and  $u_m \in H^m$ , and the space  $H^s$  is exactly the set of functions  $u$  for which the expression is finite. By direct argument, it can be seen that the space  $S^s$  is the interpolation  $[L^2, S^m]_{s/m, 2}$ . Since  $S^m \subset H^m$ , for even  $m$ , with equivalent norms, it follows from the interpolation formula that  $S^s \subset H^s$  with  $\|u\|_{H^s} \lesssim \|u\|_{S^s}$ .  $\square$

**Lemma 2.A.2.** *For  $0 \leq s \leq m$ , the space  $S^s$  is the interpolation space  $S^s = [L^2, S^m]_{s/m, 2}$ .*

*Proof.* Let  $K(t, u) = \inf\{\|u_0\|_{L^2} + t\|u_m\|_{S^m} : u = u_0 + u_m\}$ , where the infimum is taken over  $u_0 \in L^2$  and  $u_m \in S^m$ . Because  $a^2 + b^2 t^2 \leq (a + bt)^2 \leq 2a^2 + 2b^2 t^2$ , for any  $a, b, t \geq 0$ ,

$$\begin{aligned} K^2(t, u) &\asymp \inf_{u=u_0+u_m} (\|u_0\|_{L^2}^2 + t^2\|u_m\|_{S^m}^2) \\ &= \inf_{u=u_0+u_m} \left( \sum_{i=1}^\infty (u_{i,0}^2 + u_{m,i}^2 t^2 \kappa_i^{-m}) \right) = \sum_{i=1}^\infty u_i^2 \wedge (u_i^2 t^2 \kappa_i^{-m}). \end{aligned}$$

It follows that the square interpolation norm is given by

$$\begin{aligned} \int_0^\infty \frac{K^2(t, u)}{t^{2s/m}} \frac{dt}{t} &= \sum_{i=1}^\infty u_i^2 \int_0^\infty \frac{1 \wedge (t^2 \kappa_i^{-2m})}{t^{2s/m}} \frac{dt}{t} \\ &= \sum_{i=1}^\infty u_i^2 \left( \int_0^{\kappa_i^{m/2}} t^{1-2s/m} dt \kappa_i^{-m} + \int_{\kappa_i^{m/2}}^\infty \frac{1}{t^{2s/m+1}} dt \right). \end{aligned}$$

The right side can be seen to be up to a multiple equivalent to  $\asymp \sum_{i=1}^\infty u_i^2 \kappa_i^{-s} = \|u\|_{S^s}^2$ .  $\square$

## 2.B Complements for Schrödinger equation

**Lemma 2.B.1.** *Let  $K : L^2 \rightarrow L^2$  be a self-adjoint linear operator with eigenvalues  $\kappa_j \asymp j^{-p/d}$  and such  $K^{-\beta/p}(H^\beta) \subset L^2$ . If  $f \in H^\beta$  and  $u_f \in H^\beta$  for some  $\beta > d/2$ , then  $K^{-1}(fu_f) \in S^\beta$ , for  $(S^s)$  the Hilbert scale generated by the eigenfunctions and eigenvalues of  $K$ .*

*Proof.* For every eigenfunction  $h_j$  and  $s \in \mathbb{R}$ , we have  $K^{-s}h_j = \kappa_j^{-s}h_j$ . Thus the self-adjointness of  $K$  gives

$$\langle fu_f, h_j \rangle_{L^2} = \langle fu_f, \kappa_j^{\beta/p} K^{-\beta/p} h_j \rangle_{L^2} = \kappa_j^{\beta/p} \langle K^{-\beta/p}(fu_f), h_j \rangle_{L^2}.$$

Since  $\kappa_j \asymp j^{-p/d}$ , it follows that

$$\sum_{j=1}^\infty \langle fu_f, h_j \rangle_{L^2}^2 j^{2\beta/d} \asymp \sum_{j=1}^\infty \langle K^{-\beta/p}(fu_f), h_j \rangle_{L^2}^2.$$

If  $f \in H^\beta$  and  $u_f \in H^\beta$ , then  $fu_f \in H^\beta$ , since  $\beta > d/2$ , and hence we have  $K^{-\beta/p}(fu_f) \in L^2$ , by assumption. Thus the right side of the display is finite, and hence so is the left side, which implies  $fu_f \in S^\beta$ .  $\square$

**Lemma 2.B.2.** *Assume that there exist constants  $C_1, C_2 > 0$  such that the eigenvalues  $\kappa_{i_1, \dots, i_d}^2$  satisfy  $C_2(\sum_{j=1}^d i_j^2)^{-p} \leq \kappa_{i_1, \dots, i_d}^2 \leq C_1(\sum_{j=1}^d i_j^2)^{-p}$  for some  $p > 0$ . Let  $k_1^2 \geq k_2^2 \geq \dots$  denote the singular values  $\kappa_{i_1, \dots, i_d}^2$  sorted in decreasing order. Then there exist constants  $\tilde{C}_1, \tilde{C}_2 > 0$  such that*

$$\tilde{C}_1 \ell^{-p/d} \leq k_\ell \leq \tilde{C}_2 \ell^{-p/d}, \quad \ell \in \mathbb{N}. \quad (2.31)$$

Furthermore, if  $v_1, v_2, \dots$  are the multi-indexed array  $\nu_{i_1, \dots, i_d}$ , for  $(i_1, \dots, i_d) \in \mathbb{N}^d$ , in corresponding order, then

$$\|v\|_{S^\beta}^2 \asymp \sum_{(i_1, \dots, i_d) \in \mathbb{N}^d} \nu_{i_1, \dots, i_d}^2 \left( \sum_{j=1}^d i_j^2 \right)^\beta. \quad (2.32)$$

*Proof.* Let  $z : \mathbb{N} \rightarrow \mathbb{N}^d$  be the bijection between the eigenvalues  $\{\kappa_i^2 : i \in \mathbb{N}^d\}$  and the sorted eigenvalues  $(k_\ell)_{\ell \in \mathbb{N}}$ .

We start with the first statement (2.31). Let us take the eigenvalue  $k_\ell$ , with  $\ell \in \mathbb{N}$ , and let us denote by  $i := z(\ell) \in \mathbb{N}^d$  its multi-index correspondence. Note that for  $\|i\|_\infty := \max(i_1, \dots, i_d)$ , we have  $\|i\|_\infty^2 \leq \sum_{j=1}^d i_j^2 \leq d\|i\|_\infty^2$ . Then, for any  $l \in \mathbb{N}^d$ ,  $\|l\|_\infty \leq d^{-\frac{1}{2}}(C_1/C_2)^{\frac{1}{2p}}\|i\|_\infty$ ,

$$\kappa_l^2 \geq C_1 \left( \sum_{j=1}^d l_j^2 \right)^{-p} \geq C_1 d^{-p} \|l\|_\infty^{-2p} \geq C_2 \|i\|_\infty^{-2p} \geq C_2 \left( \sum_{j=1}^d i_j^2 \right)^{-p} \geq \kappa_i^2.$$

Hence  $k_\ell^2 < k_{z^{-1}(l)}^2$ , and thus  $\ell > z^{-1}(l)$ . This implies that

$$\ell \geq (d^{-\frac{1}{2}}(C_1/C_2)^{1/(2p)}\|i\|_\infty)^d = d^{-d/2}(C_1/C_2)^{d/(2p)}\|i\|_\infty^d, \quad (2.33)$$

thus

$$k_\ell^2 \geq C_1 \left( \sum_{j=1}^d i_j^2 \right)^{-p} \geq C_1 (d\|i\|_\infty^2)^{-p} \geq \frac{C_1^2}{C_2 d^{2p}} \ell^{-2p/d}.$$

For the upper bound in Equation (2.31), we follow a similar strategy. First note that

$$k_\ell^2 = \kappa_i^2 \leq C_2 \left( \sum_{j=1}^d i_j^2 \right)^{-p} \leq C_2 \|i\|_\infty^{-2p}.$$

Then for any  $l \in \mathbb{N}^d$ ,  $\|l\|_\infty > (C_2/C_1)^{1/(2p)}\sqrt{d}\|i\|_\infty$ ,

$$\kappa_l^2 \leq C_2 \left( \sum_{j=1}^d l_j^2 \right)^{-p} \leq C_2 \|l\|_\infty^{-2p} < C_1 (d\|i\|_\infty^2)^{-p} \leq C_1 \left( \sum_{j=1}^d i_j^2 \right)^{-p} \leq \kappa_i^2,$$

implying that  $k_{z^{-1}(l)}^2 < k_\ell^2$ , and therefore  $z^{-1}(l) > \ell$ . Thus

$$\ell < (C_2/C_1)^{d/(2p)} d^{d/2} \|i\|_\infty^d, \quad (2.34)$$

which in turn implies  $k_\ell^2 \leq (C_2^2/C_1)d^p\ell^{-2p/d}$ .

Next, we deal with the second statement (2.32) of the lemma. In view of the Riemann series theorem and assertion (2.33), for  $i = z(\ell) \in \mathbb{N}^d$ ,

$$\begin{aligned} (\|f\|_{S^\beta}^*)^2 &= \sum_{i \in \mathbb{N}^d} \left( \sum_{j=1}^d i_j^2 \right)^\beta f_i^2 \leq \sum_{i \in \mathbb{N}^d} (d\|i\|_\infty^2)^\beta f_i^2 \\ &\leq d^\beta \sum_{\ell \in \mathbb{N}} ((C_2/C_1)^{1/p} d \ell^{2/d})^\beta f_{z(\ell)}^2 = (C_2^{1/p} d^2 / C_1^{1/p})^\beta \|f\|_{S^\beta}^2. \end{aligned}$$

For the other direction, by similar arguments and using Equation (2.34),

$$\begin{aligned} \|f\|_{S^\beta}^2 &= \sum_{i \in \mathbb{N}^d} z^{-1}(i)^{2\beta/d} f_{z^{-1}(i)}^2 \leq \sum_{i \in \mathbb{N}^d} ((C_2/C_1)^{d/(2p)} d^{d/2} \|i\|_\infty^d)^{2\beta/d} f_{z^{-1}(i)}^2 \\ &\leq (C_2^{1/p} d / C_1^{1/p})^\beta \sum_{i \in \mathbb{N}^d} \left( \sum_{j=1}^d i_j^2 \right)^\beta f_{z^{-1}(i)}^2 = (C_2^{1/p} d / C_1^{1/p})^\beta (\|f\|_{H_\beta}^*)^2. \end{aligned}$$

□

**Lemma 2.B.3.** *Let  $(S^s)$  be the scale generated by the eigenfunctions from Equation (2.22) and  $\beta > d/2$ . Then for every  $f \in S^\beta$ , there exists an element  $\mathcal{I}_N f \in \mathbb{L}_N$  such that  $\mathcal{I}_N f(x_i) = f(x_i)$  for every  $i = 1, \dots, N$ , and*

$$\|\mathcal{I}_N f - f\|_{L^2} \lesssim N^{-\beta/d} \|f\|_{S^\beta}. \quad (2.35)$$

Furthermore, there exist constants  $0 < C_1 < C_2 < \infty$ ,

$$C_1 \|w\|_{L^2} \leq \|w\|_N \leq C_2 \|w\|_{L^2}, \quad \forall w \in \mathbb{L}_N.$$

*Proof.* The proof of Equation (2.23) is based on Section 2.3 in (Reiß 2008). This part of the proof is not restricted to the singular value basis of the inverse Laplace operator, but holds more generally. Let  $P_N : L^2([0, 1]^d) \rightarrow \mathbb{L}_N$  be the orthogonal projection on  $\mathbb{L}_N$ . Then by the triangle inequality

$$\|\mathcal{I}_N g - g\|_{L^2} \leq \|\mathcal{I}_N f - P_N f\|_{L^2} + \|P_N f - f\|_{L^2}. \quad (2.36)$$

The second term on the right-hand side can be bounded as

$$\|P_N f - f\|_{L^2}^2 = \sum_{\|l\|_\infty > m} \langle f, h_l \rangle_{L^2}^2 \leq \sum_{\|l\|_\infty > m} \langle f, h_l \rangle_{L^2}^2 \frac{\|l\|_\infty^{2\beta}}{m^{2\beta}} \leq m^{-2\beta} \|f\|_{S^\beta}^2.$$



For the first term on the right-hand side of Equation (2.36), note that

$$\begin{aligned}
 \|\mathcal{I}_N f - P_N f\|_{L^2}^2 &= \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \langle f, h_{k(2m+1)+l} \rangle \right)^2 \\
 &\leq \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \left[ \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \|k(2m+1) + l\|_\infty^{2\beta} \langle f, h_{k(2m+1)+l} \rangle^2 \right) \right. \\
 &\quad \left. \times \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \|k(2m+1) + l\|_\infty^{-2\beta} \right) \right] \\
 &\lesssim \|f\|_{S^\beta}^2 \sup_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \|k(2m+1) + l\|_\infty^{-2\beta} \right),
 \end{aligned}$$

where the last inequality follows from  $\|v\|_\infty \asymp \sqrt{\sum_{j=1}^d v_j^2}$  and Lemma 2.B.2. Since for  $\|l\|_\infty \leq m$  by triangle inequality

$$m^{-1} \|k(2m+1) + l\|_\infty \geq \|2k\|_\infty - \|l/m\|_\infty \geq 2\|k\|_\infty - 1,$$

we find that

$$\|\mathcal{I}_N f - P_N f\|_{L^2}^2 \lesssim \|f\|_{S^\beta}^2 m^{-2\beta} \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} (2\|k\|_\infty - 1)^{-2\beta} \right).$$

As seen in the proof of Lemma 2.B.2, for  $\beta > d/2$  we have

$$\sum_{k \in \mathbb{N}_0^d \setminus \{0\}} (2\|k\|_\infty - 1)^{-2\beta} \asymp \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \|k\|_\infty^{-2\beta} \asymp \sum_{i=1}^{\infty} i^{-2\beta/d} = O(1).$$

We conclude the proof of the first statement by applying the upper bounds derived above for the two terms on the right-hand side of Equation (2.36) with  $m = N^{1/d}$ .

To show that  $\mathcal{I}_N f$  coincides with  $f$  on the design points, note that for  $x \in \mathbb{X} = \{(\frac{2i}{2m+1})_{i=0,\dots,m}\}^d$ ,

$$\begin{aligned}
 f(x) &= \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \langle f, h_l \rangle_{L^2} h_l(x) + \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \langle f, h_{k(2m+1)+l} \rangle_{L^2} h_{k(2m+1)+l}(x) \\
 &= \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \langle f, h_l \rangle_{L^2} h_l(x) + \sum_{\substack{l \in \mathbb{Z}^d \\ \|l\|_\infty \leq m}} \left( \sum_{k \in \mathbb{N}_0^d \setminus \{0\}} \langle f, h_{k(2m+1)+l} \rangle_{L^2} \right) h_l(x),
 \end{aligned}$$

where the last equality follows from

$$\begin{aligned} h_{k(2m+1)+l}(x) &= 2^{d/2} \prod_{j=1}^d \sin((k(2m+1) + l)_j \pi x_j) \\ &= 2^{d/2} \prod_{j=1}^d \sin(l_j \pi x_j + 2k_j \pi r_j) = 2^{d/2} \prod_{j=1}^d \sin(l_j \pi x_j) = h_l(x), \end{aligned}$$

where  $r_j = x_j(2m+1)/2 \in \mathbb{N}$ .

Finally, to prove the equivalence of the  $L^2$ -norm with the empirical Euclidean norm  $\mathbb{L}_N$ , note that in view of Lemma 2.B.4, for any  $f = \sum_{\|l\|_\infty \leq m} c_l h_l \in \mathbb{L}_N$ ,

$$\begin{aligned} \|f\|_N^2 &= \sum_{\|l\|_\infty \leq m} \sum_{\|k\|_\infty \leq m} c_l c_k \frac{1}{N} \sum_{x \in \mathbb{X}} h_l(x) h_k(x) \\ &= \left(1 + \frac{1}{2m}\right)^d \sum_{\|l\|_\infty \leq m} c_l^2 = \left(1 + \frac{1}{2m}\right)^d \|f\|_{L^2}^2. \end{aligned}$$

□

**Lemma 2.B.4.** *Let  $h_l(x) = 2^{d/2} \prod_{j=1}^d \sin(l_j \pi x_j)$  for  $l \in \mathbb{N}^d$  and  $x \in [0, 1]^d$ . Then for arbitrary  $l, k \in \mathbb{N}^d$  with  $\|l\|_\infty, \|k\|_\infty \leq m$ , we have*

$$\langle h_l, h_k \rangle_N = \begin{cases} \left(1 + \frac{1}{2m}\right)^d & \text{if } l = k, \\ 0 & \text{otherwise,} \end{cases}$$

where the empirical measure  $\mathbb{L}_N$  is defined on the regular grid given by  $\mathbb{X} = \left\{ \left( \frac{2i}{2m+1} \right)_{i=0, \dots, m} \right\}^d$ .

*Proof.* First we consider the  $d = 1$  case and then extend the result for general  $d \in \mathbb{N}$ . By standard trigonometric equivalence, for all  $l, k \in \mathbb{N}$

$$\frac{1}{m} \sum_{x \in \mathbb{X}} 2 \sin(l\pi x) \sin(k\pi x) = \frac{1}{m} \sum_{x \in \mathbb{X}} \cos((l-k)\pi x) - \cos((l+k)\pi x).$$

Then, in view of Lemma 2.B.5, for  $1 \leq l, k \leq m$ , the right-hand side of the preceding display is equal to

$$\begin{aligned} &\frac{1}{m} \left( m 1_{2m+1|(l-k)} - \frac{1}{2} 1_{2m+1 \nmid (l-k)} - \left( m 1_{2m+1|(l+k)} - \frac{1}{2} 1_{2m+1 \nmid (l+k)} \right) \right) \\ &= \frac{1}{m} \left( m 1_{l=k} - \frac{1}{2} 1_{l \neq k} + \frac{1}{2} \right) = \left( 1 + \frac{1}{2m} \right) 1_{l=k}. \end{aligned}$$

This finishes the proof for  $d = 1$ . To extend the results to  $d > 1$ , note that

$$\begin{aligned}\langle h_l, h_k \rangle_N &= \frac{1}{N} \sum_{x \in \mathbb{X}} 2^{d/2} \prod_{j=1}^d \sin(l_j \pi x_j) 2^{d/2} \prod_{j=1}^d \sin(k_j \pi x_j) \\ &= \prod_{j=1}^d \frac{1}{m} \sum_{x_j \in \mathbb{X}_j} 2 \sin(l_j \pi x_j) \sin(k_j \pi x_j) \\ &= \begin{cases} (1 + \frac{1}{2m})^d & l = k, \\ 0 & \text{otherwise.} \end{cases}\end{aligned}$$

This concludes the proof of the lemma. □

**Lemma 2.B.5.** *For any  $r, m \in \mathbb{N}$ ,*

$$\sum_{j=1}^m \cos\left(r\pi \frac{2j}{2m+1}\right) = \begin{cases} m & 2m+1 \mid r, \\ -\frac{1}{2} & 2m+1 \nmid r. \end{cases}$$

*Proof.* For  $r = v(2m+1)$  with  $v \in \mathbb{N}$ , we have

$$\sum_{j=1}^m \cos\left(r\pi \frac{2j}{2m+1}\right) = \sum_{j=1}^m \cos(2v\pi j) = m.$$

For  $r \in \mathbb{N}$  such that  $r \neq v(2m+1)$ , for any  $v \in \mathbb{N}$ , and  $z = \exp\left(ri\pi \frac{2}{2m+1}\right)$ , we have

$$\begin{aligned}\sum_{j=1}^m \cos\left(r\pi \frac{2j}{2m+1}\right) &= \sum_{j=1}^m \operatorname{Re}\left(\exp\left(ri\pi \frac{2j}{2m+1}\right)\right) \\ &= \operatorname{Re}\left(\sum_{j=1}^m z^j\right) \\ &= \operatorname{Re}\left(\frac{z(z^m - 1)}{z - 1}\right).\end{aligned}$$

Furthermore, since  $(z - 1)(\bar{z} - 1) = 2 - 2\operatorname{Re}(z) = 2 - 2\cos(r\pi\frac{2}{2m+1})$ , and

$$\begin{aligned} z(z^m - 1)(\bar{z} - 1) &= z^m - 1 - z^{m+1} + z \\ &= \exp\left(ri\pi\frac{2m+1}{2m+1}\right)\exp\left(-ri\pi\frac{1}{2m+1}\right) - 1 \\ &\quad - \exp\left(ri\pi\frac{2m+1}{2m+1}\right)\exp\left(ri\pi\frac{1}{2m+1}\right) + \exp\left(ri\pi\frac{2}{2m+1}\right) \\ &= (-1)^r\left(\exp\left(-ri\pi\frac{1}{2m+1}\right) - \exp\left(ri\pi\frac{1}{2m+1}\right)\right) - 1 \\ &\quad + \exp\left(ri\pi\frac{2}{2m+1}\right), \end{aligned}$$

so  $\operatorname{Re}(z(z^m - 1)(\bar{z} - 1)) = \cos(r\pi\frac{2}{2m+1}) - 1$ , we have

$$\operatorname{Re}\left(\frac{z(z^m - 1)}{z - 1}\right) = \frac{\cos(r\pi\frac{2}{2m+1}) - 1}{2 - 2\cos(r\pi\frac{2}{2m+1})} = -\frac{1}{2}.$$

finishing the proof of the lemma.  $\square$

## 2.C Complements for heat equation with absorption

**Lemma 2.C.1.** *The eigenvalues of the operator  $K^\top K$  for the operator  $K$  defined in Equation (2.4) for  $\mathcal{L} = \frac{\partial}{\partial t} - \frac{1}{2}\Delta$  satisfy*

$$\frac{\pi^{-2}}{(k+1)^2 + \sum_{j=1}^d i_j^2/4} \leq \lambda_{k,i_1,\dots,i_d} \leq \frac{1}{k^2\pi^2 + \sum_{j=1}^d i_j^2/4}, \quad i \in \mathbb{N}^d, k \in \mathbb{N}.$$

The corresponding eigenfunctions are given by

$$h_{i,k}(t, x) := \frac{-\frac{1}{\tan(\nu_{i,k})} \sin(\nu_{i,k}t) + \cos(\nu_{i,k}t)}{\sqrt{\frac{1}{\tan(\nu_{i,k})} + \frac{\nu_{i,k}}{\sin(\nu_{i,k})^2}}} 2^{d/2} \prod_{j=1}^d \sin(i_j\pi x_j), \quad (2.37)$$

for  $i \in \mathbb{N}^d$  and  $k \in \mathbb{N}$ , where  $\nu_{i,k} \in (k\pi, (k+1)\pi]$  is the solution to  $\nu_{i,k}/\tan(\nu_{i,k}) + \pi^2 \sum_{j=1}^d i_j^2/2 = 0$ . The set of eigenfunctions is uniformly bounded in  $L_\infty$ .

*Proof.* Let  $\mathcal{O} := (0, 1)^d$  and then the adjoint operator  $K^\top : L^2 \rightarrow L^2$  of  $K$  is defined as the solution operator  $v \mapsto K^\top v$  of the homogeneous differential equation, for some constant  $c \in \mathbb{R}$ ,

$$\begin{cases} -\frac{d}{dt}K^\top v - \frac{1}{2}\Delta K^\top v = v & (x, t) \in \mathcal{O} \times (0, 1], \\ K^\top v(t, x) = 0 & x \in \partial\mathcal{O}, t \in (0, 1], \\ K^\top v(1, x) = 0 & x \in \mathcal{O}. \end{cases} \quad (2.38)$$

Hence the operator  $K^\top K$  is given as the solution operator of  $v \mapsto K^\top K v$  of the homogeneous differential equation

$$\begin{cases} \left(-\frac{d^2}{dt^2} + \frac{1}{4}\Delta^2\right)K^\top K v = v & (x, t) \in \mathcal{O} \times (0, 1], \\ K^\top K v(t, x) = 0 & x \in \partial\mathcal{O}, t \in (0, 1), \\ K^\top K v(1, x) = 0 & x \in \mathcal{O}, \\ \left(-\frac{d}{dt} - \frac{1}{2}\Delta\right)K^\top K v(t, x) = 0 & x \in \partial\mathcal{O}, t \in (0, 1) \\ \left(-\frac{d}{dt} - \frac{1}{2}\Delta\right)K^\top K v(0, x) = 0 & x \in \mathcal{O}. \end{cases} \quad (2.39)$$

Thus for an eigenfunction  $h$  of  $K^\top K$  with eigenvalue  $\lambda > 0$  we have  $h = \left(-\frac{d^2}{dt^2} + \frac{1}{4}\Delta^2\right)\lambda h$ . Consider  $h(t, x) = g(t)f(x)$ , then

$$gf = -\lambda \frac{d^2}{dt^2}gf + \lambda \frac{1}{4}g\Delta^2 f. \quad (2.40)$$

Dividing both sides by  $gf$  results in

$$1 = -\lambda \frac{\frac{d^2}{dt^2}g}{g} + \frac{1}{4}\lambda \frac{\Delta^2 f}{f}. \quad (2.41)$$

Since the left-hand side is constant, and the two terms on the right do not share any variable, they have to be constants. Therefore, we get the following systems of equations

$$\begin{cases} \frac{d^2}{dt^2}g(t) = -\frac{c}{\lambda}g(t) & t \in (0, 1), \\ \left(-\frac{d}{dt} - \frac{1}{2}\Delta\right)g(t)f(x) = 0 & t = 0, x \in \mathcal{O}, \\ g(t) = 0 & t = 1. \end{cases} \quad (2.42)$$

and

$$\left\{ \begin{array}{ll} \Delta^2 f(x) = 4 \frac{1-c}{\lambda} f(x) & x \in \mathcal{O}, \\ f(x) = 0 & x \in \partial\mathcal{O}, \\ \left( -\frac{d}{dt} - \frac{1}{2}\Delta \right) g(t)f(x) = 0 & x \in \partial\mathcal{O}, t \in (0, 1). \end{array} \right. \quad (2.43)$$

We now assume that  $1 - c > 0$ . Then the second system can be solved by  $f(x) = 2^d \prod_{j=1}^d \sin(i_j \pi x_j)$  for  $i_j \in \mathbb{N}$ . It follows that  $\lambda_i = 4(1 - c)\mu_i^{-2}$ , with  $\mu_i = \pi^2 \sum_{j=1}^d i_j^2$ . Now assume that  $c > 0$ . Then the general solution of the second order ODE in Equation (2.42) is

$$g(t) = c_1 \sin(\nu t) + c_2 \cos(\nu t) \quad (2.44)$$

for  $\nu = \sqrt{c/\lambda}$ . From the boundary condition at  $t = 1$  we get  $0 = c_1 \sin(\nu) + c_2 \cos(\nu)$ , so  $c_1/c_2 = -1/\tan(\nu)$ . From the boundary condition at  $t = 0$ , we have

$$0 = \left( -\frac{d}{dt} - \frac{1}{2}\Delta \right) g f = (-g' + \frac{\mu}{2}g) f.$$

We thus require  $-g'(0) + \mu g(0)/2 = 0$ , which is equivalent to  $0 = \frac{\nu}{\tan \nu} + \frac{\mu}{2}$ .

Since  $g f$  needs to be normalised with respect to the  $L^2$ -norm, we require that  $\int_0^1 g(t)^2 dt = 1$ , thus

$$g(t) = \frac{-\frac{1}{\tan(\nu)} \sin(\nu t) + \cos(\nu t)}{\sqrt{\frac{-\frac{1}{\tan(\nu)} + \frac{\nu}{\sin(\nu)^2}}{2\nu}}}.$$

Note that for  $\nu$  satisfying  $\frac{\nu}{\tan(\nu)} + \frac{\mu}{2} = 0$ , we have  $\tan(\nu) = -2\nu/\mu < 0$ . The norming constant is lower bounded by

$$\sqrt{\frac{-\frac{1}{\tan(\nu)} + \frac{\nu}{\sin(\nu)^2}}{2\nu}} \geq \sqrt{0 + \frac{1}{2 \sin(\nu)^2}} \geq \sqrt{\frac{1}{2}}.$$

Thus the second term of  $g$  is bounded. The absolute value of the constant in front of the first term can be bounded using the definitions of  $\nu$  and  $\mu$  as

$$\left| \frac{\frac{1}{\tan(\nu)}}{\sqrt{\frac{-\frac{1}{\tan(\nu)} + \frac{\nu}{\sin(\nu)^2}}{2\nu}}} \right| = \sqrt{\frac{2\nu}{-\tan(\nu) + \frac{\nu}{\cos(\nu)^2}}} = \left( \frac{1}{\mu} + \frac{1}{2 \cos(\nu)^2} \right)^{-\frac{1}{2}} \leq \sqrt{\frac{2}{3}}.$$

We see that  $g$  is uniformly bounded, and thus also the eigenfunctions are uniformly bounded. It remains to show that all eigenfunctions of  $K^\top K$  are of the form (2.37). Note that the  $(f_i(x))_{i \in \mathbb{N}^d}$  form an orthonormal basis of  $L^2$ , and the  $(g_i)_{i \in \mathbb{N}}$  form an orthonormal basis of  $L^2([0, 1])$ , since they are the eigenvalues of a Sturm-Liouville problem. Hence the set of products of these functions forms an orthonormal basis of  $L^2(\mathcal{O} \times [0, 1])$ . Therefore the assumption that  $0 < c < 1$  is justified, as there do not exist more eigenfunctions.

Now consider a fixed  $\mu \in \mathbb{N}$ . Note that  $\frac{d}{dx} \frac{x}{\tan(x)} = \frac{\sin(x) \cos(x) - x}{\sin(x)^2} < 0$  for all  $x > 0$  such that  $\tan(x) \neq 0$ . Furthermore, the function  $\frac{x}{\tan(x)}$  has asymptotes for  $x \rightarrow k\pi, k \in \mathbb{N}_+$ . Thus for each  $k \in \mathbb{N}_0$ , there exists a unique  $\nu_k \in [k\pi, (k+1)\pi]$  such that  $\frac{\nu_k}{\tan(\nu)} + \frac{\mu}{2} = 0$ . The eigenvalues are therefore bounded by

$$\frac{1}{(k+1)^2\pi^2 + \mu_i^2/4} \leq \lambda_{k,i} \leq \frac{1}{k^2\pi^2 + \mu_i^2/4},$$

concluding the proof of our statement. □

## 2.D Smoothness norm contraction rates in the Gaussian white noise linear inverse problems

In this section, we extend the results of (Gugushvili, van der Vaart and Yan 2020) to give rates of contraction relative to smoothness norms in the context of the Gaussian white noise linear inverse problem (2.7). We obtain general results for an abstract scale of spaces of numerical sequences. By taking these sequences equal to the coefficients of functions relative to a suitable basis, the scale may be identified with concrete function spaces. For instance, for a suitable wavelet basis, the scale can (partially) be identified with the scale of Sobolev spaces and the general result will give contraction relative to Sobolev norms of positive order. By the Sobolev embedding this implies consistency (at a rate) for other norms such as the uniform norm. Alternatively, the sequences can be taken equal to coefficients relative to the eigenbasis of an operator of interest, giving a contraction rate to a norm intrinsic to the problem.

### 2.D.1 Smoothness scale

Fix a sequence  $1 \leq \rho_1 \leq \rho_2 \leq \dots \uparrow \infty$ , and for  $s \in \mathbb{R}$  consider the Hilbert spaces

$$S^s = \left\{ g = (g_i) : \mathbb{N} \rightarrow \mathbb{R} : \sum_{i \in \mathbb{N}} \rho_i^{2s} g_i^2 < \infty \right\},$$

with the inner product and norm defined by

$$\langle g, h \rangle_{S^s} = \sum_{i \in \mathbb{N}} \rho_i^{2s} g_i h_i, \quad \|g\|_{S^s} = \sqrt{\sum_{i \in \mathbb{N}} \rho_i^{2s} g_i^2}.$$

The spaces are nested: if  $s \geq t$ , then  $S^s \subset S^t$ . By the Cauchy-Schwarz inequality, the inner product and norm satisfy, for  $s \geq 0$  and every  $t \in \mathbb{R}$ ,

$$\|g\|_{S^{t-s}} = \sup_{g \in S^{t+s} : \|g\|_{S^{t+s}} \leq 1} \langle g, h \rangle_t, \quad g \in S^{t-s}. \quad (2.45)$$

The map defined by  $(C_s g)_i = \rho_i^s g_i$  is an isometry  $C_s : S^{t+s} \rightarrow S^t$ , and  $\langle g, h \rangle_{S^s} = \langle C_s g, C_s h \rangle_{S^0}$ . The subspace  $V_j = \{g \in S_0 : g_i = 0 \text{ for } i > j\}$  has the properties, for  $s < t$ , for  $P_j$  the orthogonal projection onto  $V_j$ ,

$$\begin{aligned} \|g - P_j g\|_{S^s}^2 &= \sum_{i > j} \rho_i^{2s} g_i^2 \leq \rho_j^{-2(t-s)} \|g\|_{S^t}^2, \\ \|g\|_{S^t}^2 &\leq \rho_j^{2(t-s)} \|g\|_{S^s}^2, \quad \text{if } g \in V_j. \end{aligned}$$

Thus it satisfies the conditions of a smoothness scale as in (Gugushvili, van der Vaart and Yan 2020).

We consider a prior distribution on the sequences given as the law of the random sequence  $\mathbb{G} = (\rho_i^{-\alpha} i^{-\frac{1}{2}} c_i Z_i)$ , for  $Z_i \stackrel{\text{i.i.d.}}{\sim} N(0, 1)$  and a sequence  $(c_i)$  such that  $c_\pi \leq c_i \leq C_\pi$ , for some constants  $0 < c_\pi < C_\pi$ , for all  $i$ . Then

$$\mathbb{E} \|\mathbb{G}\|_{S^s}^2 = \sum_{i \in \mathbb{N}} \rho_i^{-2(\alpha-s)} i^{-1} c_i^2.$$

If  $\rho_i \gtrsim i^r$ , for some  $r > 0$ , then this is finite for every  $s < \alpha$ , and hence the prior concentrates on the space  $S^s$ , for  $s < \alpha$ : it is regular of nearly order  $\alpha$ .

**Example 2.D.1.** Given a wavelet basis  $\{\psi_{j,k} : j \in \mathbb{N}_0, k = 1, \dots, 2^{jd}\}$  for functions on a bounded domain in  $\mathbb{R}^d$ , we can consider the sequences of coefficients  $(h_{j,k})$  of expansions of given functions  $h$  in the basis, linearly ordered by level and



location as  $h_{0,1}, h_{1,1}, \dots, h_{1,2^d}, h_{2,1}, \dots, h_{2,2^{2d}}, h_{3,1}, \dots$ . The ordering defines the map

$$(j, k) \mapsto i := 1 + 2^d + \dots + 2^{(j-1)d} + k = \frac{2^{jd} - 1}{2^d - 1} + k.$$

Set  $\rho_i = 2^j$  and  $\bar{h}_i = h_{j,k}$  if  $i \in \mathbb{N}$  is the image of  $(j, k)$  under this map. Then the square norm  $\sum_{j=1}^{\infty} \sum_k 2^{js} h_{j,k}^2$  is equal to  $\sum_{i=1}^{\infty} \rho_i^{2s} \bar{h}_i^2$  and hence the smoothness space with norm the square root of  $\sum_{j=1}^{\infty} \sum_k 2^{js} h_{j,k}^2$  corresponds to the space  $S^s$ .

The wavelet basis can be chosen so that the scale of spaces  $S^s$ , with  $s$  in a bounded interval, corresponds (partially) to the scale of Sobolev spaces.

By the preceding display  $i \asymp 2^{jd}$ , as  $j \rightarrow \infty$ , and hence  $\rho_i \asymp i^{1/d}$ . The Gaussian prior  $\sum_{j=1}^{\infty} \sum_k 2^{-j(\alpha+d/2)} Z_{j,k} \psi_{j,k}$ , for  $Z_{j,k} \stackrel{i.i.d.}{\sim} N(0, 1)$ , in the doubly indexed sequence space, corresponds to the Gaussian prior  $(\rho_i^{-\alpha-d/2} \bar{Z}_i)$ , for  $\bar{Z}_i \stackrel{i.i.d.}{\sim} N(0, 1)$  in the sequence space. This is the prior  $\mathbb{G}$  for  $\rho_i^{-d/2} = i^{-\frac{1}{2}} c_i$ , i.e. the constants  $c_i$  account for the discrepancy between  $i^{1/d}$  and  $\rho_i$ .

From now on we consider a sequence  $\rho_i$  such that  $c_0 i^{1/d} \leq \rho_i \leq C_0 i^{1/d}$ , for constants  $C_0 \geq c_0 > 0$  and some  $d > 0$ . The power  $1/d$  is included so that the index  $s$  of  $S^s$  can be interpreted as “smoothness” of a function on a  $d$ -dimensional domain, but is irrelevant for the results. On the other hand, the polynomial increase of  $\rho_i$  will be important for the form of the results. Allowing  $\rho_i \asymp i^{1/d}$  rather than setting exactly  $\rho_i = i^{1/d}$  and including the constants  $c_i$  in the prior ensure that the construction applies to the “exact” weights  $2^j$  in the wavelet basis, but is inessential for the results.

## 2.D.2 Inverse problem

Let  $L$  be an arbitrary separable Hilbert space, with inner product  $\langle \cdot, \cdot \rangle_L$  and norm  $\| \cdot \|_L$ . Let  $K : S^0 \rightarrow L$  be a linear operator such that, for some given  $p \geq 0$  and constants  $C \geq c > 0$ ,

$$c \|f\|_{S^{-p}} \leq \|Kf\|_L \leq C \|f\|_{S^{-p}}, \quad f \in S^0. \quad (2.46)$$

For a standard white noise process  $\xi$  in  $L$ , we observe

$$Y_N = Kf + \frac{1}{\sqrt{N}} \xi.$$

We put the  $\alpha$ -regular prior  $\mathbb{G}$  as indicated on  $f$  and denote the corresponding posterior distribution of  $f$  by  $\Pi_N(f \in \cdot \mid Y_N)$ .

**Theorem 2.D.2.** *Assume that  $c_0 i^{1/d} \leq \rho_i \leq C_0 i^{1/d}$ , for constants  $C_0 \geq c_0 > 0$  and  $d > 0$ . Let  $K : S^0 \rightarrow L$  be a linear operator satisfying Equation (2.46) for every  $f \in S^0$ . If  $f_0 \in S^\beta$  and  $0 < \delta < \alpha \wedge \beta$ , then there exists a constant  $M$  such that  $\mathbb{E}_{f_0} \Pi_N(f : \|f - f_0\|_{S^\delta} \leq MN^{-(\alpha \wedge \beta - \delta)/(2\alpha + 2p + d)} \mid Y_N) \rightarrow 1$ .*

*Proof.* For positive constants to  $c_1, c_2$  be determined later in the proof, and a sequence of constants  $c_{3,N}$ , set

$$\epsilon_N = c_1 \left( \frac{1}{N} \right)^{\frac{\alpha \wedge \beta + p}{2\alpha + 2p + d}}, \quad \eta_N = c_2 \left( \frac{1}{N} \right)^{\frac{\alpha \wedge \beta - \delta}{2\alpha + 2p + d}}, \quad \rho_{j_N} = c_{3,N} N^{\frac{1}{2\alpha + 2p + d}}.$$

By the assumption that  $\rho_i \asymp i^{1/d}$ , we can find  $j_N \in \mathbb{N}$  so that the corresponding  $c_{3,N}$  is bounded from below and above by positive constants.

The smoothing property of  $K$  implies that  $K$  is injective. Therefore its restriction  $K : V_j \rightarrow L$  to the finite-dimensional subspace  $V_j \subset S^0$  is invertible on its range  $KV_j$ . Let  $R_j : L \rightarrow V_j$  be the map  $R_j = K^{-1}Q_j$  formed by first applying the orthogonal projection  $Q_j : L \rightarrow KV_j$  of the Hilbert space  $L$  onto  $KV_j$  and next applying the inverse  $K^{-1}$  (the map  $R_j K : S^0 \rightarrow V_j \subset S^0$  is called the Galerkin solution to  $Kf$ ).

Let  $p_f^{(N)}$  be a density of  $Y_N$  relative to a suitable dominating measure. The proof is based on the following claims: we can choose  $c_1$  and  $c_2$  so that

- (1) The events  $A_N := \left\{ \int p_f^{(N)} / p_{f_0}^{(N)}(Y_N) d\Pi(f) \geq e^{-2N\epsilon_N^2} \right\}$  satisfy asymptotically  $P_{f_0}^{(N)}(A_N^c) \rightarrow 0$ .
- (2)  $\Pi(\|Kf - Kf_0\|_L < \epsilon_N) \geq e^{-N\epsilon_N^2}$ .
- (3)  $\Pi(\|R_{j_N}Kf - f\|_{S^\delta} > \eta_N) \leq e^{-4N\epsilon_N^2}$ .
- (4) There exist  $M > 0$  and tests  $\tau_N$  such that  $P_{f_0}^{(N)}\tau_N \rightarrow 0$  and  $P_f^{(N)}(1 - \tau_N) \leq e^{-4N\epsilon_N^2}$ , for every  $f \in \mathcal{F}_N$ , where

$$\mathcal{F}_N = \{f : \|f - f_0\|_{S^\delta} > M\eta_N, \|R_{j_N}Kf - f\|_{S^\delta} < \eta_N\}.$$

Claim (1) is the usual evidence lower bound, as in Lemma 8.21 of (Ghosal and van der Vaart 2017). Its validity follows from the prior mass condition (2) and the fact that the Kullback-Leibler divergence and variation between  $P_f^{(N)}$  and  $P_{f_0}^{(N)}$  are equal to  $N\|Kf - Kf_0\|_L^2/2$  and twice this quantity, respectively (e.g. Lemma 8.30 in the same reference or Lemma 9.1 in (Gugushvili, van der Vaart

and Yan 2020)). Claim (2) is the content of Lemma 2.D.4. Claims (3) and (4) are proved at the end of this proof.

Given Claims (1)–(4) the proof of the theorem can be finished as follows. We bound  $\Pi_N(f : \|f - f_0\|_{S^\delta} > M\eta_N \mid Y^{(N)})$  by

$$\tau_N + 1_{A_N^c} + \Pi_N(\|R_{j_N}Kf - f\|_{S^\delta} \geq \eta_N \mid Y_N) + \Pi_N(\mathcal{F}_N \mid Y_N)1_{A_N}(1 - \tau_N).$$

The expectation under  $P_{f_0}^{(N)}$  of the first two terms tends to zero by the first part of Claim (4) and Claim (1). The expectation of the third term tends to zero by Claims (3) and (2), in view of the remaining mass condition (see Theorem 8.20(iii') in (Ghosal and van der Vaart 2017)). To bound the fourth term, we use Bayes' formula to write the posterior probability of  $\mathcal{F}_N$  as the quotient of  $\int_{\mathcal{F}_N} p_f^{(N)}/p_{f_0}^{(N)}(Y_N) d\Pi(f)$  and  $\int p_f^{(N)}/p_{f_0}^{(N)}(Y_N) d\Pi(f)$ , and bound the denominator on the event  $A_N$  below by  $e^{-2N\epsilon_N^2}$ . This shows that the expectation of the fourth term is bounded above by

$$e^{2N\epsilon_N^2} P_{f_0}^{(N)} \int_{\mathcal{F}_N} \frac{p_f^{(N)}}{p_{f_0}^{(N)}} d\Pi(f) (1 - \tau_N) \leq e^{2N\epsilon_N^2} \sup_{f \in \mathcal{F}_N} P_f^{(N)} (1 - \tau_N),$$

which tends to zero by the second part of (4).

*Proof of Claim (3).* We apply Lemma 2.D.5 with  $j = j_N$  as specified at the beginning of the proof. We choose  $t$  such that  $bt^2 j_N \rho_{j_N}^{2(\alpha-\delta)} = 4N\epsilon_N^2$ , thus reducing the right side of the lemma to  $e^{-4N\epsilon_N^2}$ . We always have that  $\eta_N = c_2(\rho_{j_N}/c_{3,N})^{-(\alpha\wedge\beta-\delta)} \geq c_2(\rho_{j_N}/c_{3,N})^{-(\alpha-\delta)}$ . Therefore it suffices to show that also  $\eta_N \gtrsim t = 2\sqrt{N}\epsilon_N j_N^{-\frac{1}{2}} \rho_{j_N}^{\delta-\alpha}/\sqrt{b}$ . This follows, since  $N\epsilon_N^2 = c_1^2 N^{(2\alpha-2(\alpha\wedge\beta)+d)/(2\alpha+2p+d)}$  and  $j_N \asymp N^{d/(2\alpha+2p+d)}$ .

*Proof of Claim (4).* To construct the tests in (4), we note that observing  $Y_N$  is equivalent to observing a Gaussian stochastic process  $(Y_N(g) : g \in L)$  with mean  $(\langle Kf, g \rangle_L : g \in L)$  and covariance function  $\mathbb{E}Y_N(g)Y_N(h) = \langle g, h \rangle_L$ . In particular, for a given orthonormal basis  $(\psi_i)_{i \leq j}$  of  $KV_j$ , we observe the variable  $\sum_{i=1}^j Y_N(\psi_i)\psi_i$ , which we shall denote by  $Q_j Y_N$  (it can formally be understood to be the projection of  $Y_N$  onto  $KV_j$ ). It can be represented in distribution as

$$Q_j Y_N = Q_j Kf + \frac{1}{\sqrt{N}} \xi_j,$$

for the Gaussian variable  $\xi_j = \sum_{i=1}^j \langle \xi, \psi \rangle_L \psi$  with mean zero and  $\mathbb{E}\langle \xi_j, g \rangle_L^2 = \|Q_j g\|_L^2$ . Then  $R_j Q_j Y_N = R_j Kf + N^{-\frac{1}{2}} R_j \xi_j$  is a well defined variable with

values in  $V_j \subset S^0$ , with  $R_j \xi_j$  a Gaussian random element in  $V_j \subset S^0$  with strong second moment

$$\begin{aligned} \mathbb{E} \|R_j \xi_j\|_{S^0}^2 &\leq \|R_j\|_{S^0}^2 \mathbb{E} \|\xi_j\|_L^2 \\ &= \|R_j\|_{S^0}^2 \mathbb{E} \sum_{i \leq j} \langle \xi_j, \psi_i \rangle_L^2 \\ &= \|R_j\|_{S^0}^2 j \\ &\lesssim \rho_j^{2p} j, \end{aligned}$$

and weak second moment

$$\begin{aligned} \sup_{\|f\|_{S^0} \leq 1} \mathbb{E} \langle R_j \xi_j, f \rangle_{S^0}^2 &= \sup_{\|f\|_{S^0} \leq 1} \mathbb{E} \langle \xi_j, R_j^* f \rangle_L^2 \\ &= \sup_{\|f\|_{S^0} \leq 1} \|Q_j R_j^* f\|_L^2 \leq \|R_j^*\|_{S^0}^2 \\ &\lesssim \rho_j^{2p}. \end{aligned}$$

In both cases the inequality on the norms  $\|R_j\|_{S^0} = \|R_j^*\|_{S^0}$  of the operators  $R_j : L \rightarrow S^0$  and its adjoint  $R_j^* : S^0 \rightarrow L$  at the right-hand side follow from Equation (2.49). The first inequality implies that the first moment  $\mathbb{E} \|R_j \xi_j\|_{S^0}$  of the variable  $\|R_j \xi_j\|_{S^0}$  is bounded above by  $\rho_j^p \sqrt{j}$ . By Borell's inequality (e.g. Lemma 3.1 in (Ledoux and Talagrand 1991) and subsequent discussion or Proposition A.2.1 in (van der Vaart and Wellner 2023)), applied to the Gaussian random variable  $R_j \xi_j$  in  $S^0$ , we see that there exist constants  $a, b > 0$  such that, for every  $t > 0$ ,

$$\Pr(\|R_j \xi_j\|_{S^0} > a \rho_j^p \sqrt{j} + t) \leq e^{-bt^2/\rho_j^{2p}}.$$

Since  $R_j \xi_j \in V_j$ , we have  $\|R_j \xi_j\|_{S^\delta} \leq \rho_j^\delta \|R_j \xi_j\|_{S^0}$  and hence

$$\Pr(\|R_j \xi_j\|_{S^\delta} > a \rho_j^{p+\delta} \sqrt{j} + \rho_j^\delta t) \leq e^{-bt^2/\rho_j^{2p}}.$$

Choose  $j = j_N$  and  $t^2 = 4n\epsilon_N^2 \rho_{j_N}^{2p}/b$  to reduce the right side to  $e^{-4N\epsilon_N^2}$ . Then  $\rho_{j_N}^\delta t = 2\sqrt{N}\epsilon_N \rho_{j_N}^{p+\delta}/\sqrt{b} \simeq \sqrt{N}\eta_N$  and  $\rho_{j_N}^{p+\delta} \sqrt{j_N} \lesssim \sqrt{N}\eta_N$ . It follows that, for a sufficiently large constant  $c_4$ ,

$$P_f^{(N)}(\|R_{j_N} Y_N - R_{j_N} K f\|_{S^\delta} > c_4 \eta_N) \leq e^{-4N\epsilon_N^2}. \quad (2.47)$$

We define

$$\tau_N = 1\{\|R_{j_N} Y_N - R_{j_N} K f_0\|_{S^\delta} > c_4 \eta_N\}.$$

It is immediate that  $P_{f_0}^{(N)} \tau_N \leq e^{-4N\epsilon_N^2} \rightarrow 0$ . Since  $f_0 \in S^\beta$  and  $\delta < \beta$ , we have

$$\|f_0 - R_{j_N} K f_0\|_{S^\delta} \leq \rho_{j_N}^{-(\beta-\delta)} \|f_0\|_{S^\beta} \leq c_6 \eta_N \|f_0\|_{S^\beta}. \quad (2.48)$$

It follows that  $\|R_{j_N} K f - f\|_{S^\delta} < \eta_N$  implies that

$$\|R_{j_N} K f - R_{j_N} K f_0\|_{S^\delta} \geq \|f - f_0\|_{S^\delta} - \eta_N - \eta_N c_6 \|f_0\|_{S^\beta}.$$

If also  $\|f - f_0\|_{S^\delta} > M \eta_N$ , then the right side is bounded below by  $(M - 1 - c_6 \|f_0\|_{S^\beta}) \eta_N$ , and hence

$$\|R_{j_N} Y_N - R_{j_N} K f\|_{S^\delta} \geq (M - 1 - c - 6 \|f_0\|_{S^\beta}) \eta_N - \|R_{j_N} Y_N - R_{j_N} K f_0\|_{S^\delta}.$$

If  $\tau_N = 0$ , then the norm on the far right is bounded above by  $c_4 \eta_N$ . It follows that  $P_f^{(N)}(1 - \tau_N) \leq \Pr(\|R_{j_N} Y_N - R_{j_N} K f\|_{S^\delta} \geq c_5 \eta_N)$  for  $c_5 = M - 1 - c_6 \|f_0\|_{S^\beta} - c_4$ , whenever  $f \in \mathcal{F}_N$ . For  $M \geq 1 + c_6 \|f_0\|_{S^\beta} + 2c_4$ , the latter probability is bounded above by  $e^{-4N\epsilon_N^2}$ .  $\square$

**Lemma 2.D.3.** *If  $K : S^0 \rightarrow L$  is a bounded linear operator satisfying Equation (2.46) for every  $f \in S^0$ , then the norms of the operators  $R_j : L \rightarrow S^0$  and  $R_j K : S^0 \rightarrow S^0$  and  $R_j K - I : S^t \rightarrow S^s$  satisfy*

$$\|R_j\|_{L \rightarrow S^0} \lesssim \rho_j^p, \quad (2.49)$$

$$\|R_j K\|_{S^s \rightarrow S^s} \lesssim 1, \quad s \geq 0, \quad (2.50)$$

$$\|R_j K - I\|_{S^t \rightarrow S^s} \lesssim \rho_j^{-(t-s)}, \quad t \geq s \geq 0. \quad (2.51)$$

*Proof.* For  $s = 0$  the lemma is the same as Lemma A.2 in (Gugushvili, van der Vaart and Yan 2020). For the extension of Equation (2.50) to general  $s \geq 0$ , we first use the triangle inequality to see that  $\|R_j K f\|_{S^s} \leq \|R_j K f - P_j f\|_{S^s} + \|P_j f\|_{S^s}$ , for  $P_j : S^s \rightarrow V_j \subset S^s$  the orthogonal projection onto  $V_j$ . Since  $P_j$  simply sets the coefficients of index larger than  $j$  to zero, we have  $\|P_j f\|_{S^s} \leq \|f\|_{S^s}$ , for any  $s$ . Because  $R_j K f - P_j f \in V_j$  and  $R_j K$  is the identity on  $V_j$ , we have

$$\begin{aligned} \|R_j K f - P_j f\|_{S^s} &\leq \rho_j^s \|R_j K (I - P_j) f\|_{S^0} \leq \rho_j^s \|R_j\|_{L \rightarrow S^0} \|K (I - P_j) f\|_L \\ &\lesssim \rho_j^s \rho_j^p \|(I - P_j) f\|_{S^{-p}} \leq \|f\|_{S^s}, \end{aligned}$$

where we used Equation (2.49) in the second last step and  $\|(I - P_j) f\|_{S^{-p}} \leq \rho_j^{-(p+s)} \|f\|_{S^s}$  in the last. We conclude that  $\|R_j K f\|_{S^s} \lesssim \|f\|_{S^s}$ , for every  $f$ , which is Equation (2.50).

Because  $(R_j K - I)f = (R_j K - I)(I - P_j)f$ , for every  $f \in S^s$ , we have

$$\|(R_j K - I)f\|_{S^s} \leq (\|R_j K\|_{S^s \rightarrow S^s} + 1)\|f - P_j f\|_{S^s}.$$

The right side is bounded by a multiple of  $\|f - P_j f\|_{S^s} \leq \rho_j^{-(t-s)}\|f\|_{S^t}$ , for  $t \geq s$ . This proves Equation (2.51).  $\square$

**Lemma 2.D.4.** *Let  $\alpha, \beta, p > 0$  and let  $K : S^0 \rightarrow L$  be a bounded linear operator satisfying Equation (2.46) for every  $f \in S^0$ . Let  $\mathbb{G} = (\rho_i^{-\alpha} i^{-\frac{1}{2}} c_i Z_i)$ , for  $Z_i \stackrel{i.i.d.}{\sim} N(0, 1)$  and constants  $c_i$  that are bounded away from 0 and  $\infty$ . Then for every  $f_0 \in S^\beta$  there exists a constant  $C$  so that, for every  $\epsilon < 1$ ,*

$$\log \Pr(\|K\mathbb{G} - Kf_0\|_L < C\epsilon) \geq \left(\frac{1}{\epsilon}\right)^{\frac{2\alpha-2\beta+d}{p+\beta}} + \left(\frac{1}{\epsilon}\right)^{\frac{d}{p+\alpha}}.$$

Consequently, there exists a constant  $c_1$  such that  $\epsilon_N = c_1 N^{-(\alpha \wedge \beta + p)/(2\alpha + 2p + d)}$  satisfies  $\Pr(\|K\mathbb{G} - Kf_0\|_L < \epsilon_N) \geq e^{-N\epsilon_N^2}$ .

*Proof.* By the smoothing property of Equation (2.46), the left side is bounded below by  $\log \Pr(\|\mathbb{G} - f_0\|_{S^{-p}} < C_1\epsilon)$ . This is bounded below by, for  $\mathbb{H}$  the reproducing kernel Hilbert space of  $\mathbb{G}$ ,

$$\inf_{h \in \mathbb{H}: \|h - f_0\|_{S^{-p}} < C_1\epsilon} \|h\|_{\mathbb{H}}^2 - \log \Pr(\|\mathbb{G}\|_{S^{-p}} < C_1\epsilon).$$

The reproducing kernel Hilbert space of the Gaussian variable  $\mathbb{G}$  consists of the sequences  $(\rho_i^{-\alpha} i^{-\frac{1}{2}} c_i w_i)$ , for  $w \in \ell_2$ , with square norm  $\sum_{i=1}^{\infty} w_i^2$ , or equivalently the sequences  $(h_i)$  with square norm  $\sum_{i=1}^{\infty} \rho_i^{2\alpha} i h_i^2 / c_i$ . Since  $i \simeq \rho_i^d$ , this norm is equivalent to the norm of  $S^{\alpha+d/2}$ . Thus the first term in the preceding display is bounded above by a multiple of

$$\inf \left\{ \sum_{i=1}^{\infty} \rho_i^{2\alpha+d} h_i^2 : \sum_{i=1}^{\infty} \rho_i^{-2p} (h_i - f_{0,i})^2 < \epsilon^2 \right\}.$$

If  $\alpha + d/2 \leq \beta$ , we can choose  $h = f_0$  to see that the infimum is bounded above by  $\|f_0\|_{S^{\alpha+d/2}}^2 \leq \|f_0\|_{S^\beta}^2$ . For  $\alpha + d/2 > \beta$ , we choose  $h_i = f_{0,i}$  for  $i \leq J$  and  $h_i = 0$  for  $i > J$  with  $\rho_J \simeq (1/\epsilon)^{1/(p+\beta)}$ . Then

$$\begin{aligned} \|h - f_0\|_{S^{-p}}^2 &= \sum_{i=J+1}^{\infty} \rho_i^{-2p} f_{0,i}^2 \leq \rho_J^{-(2p+2\beta)} \sum_{i=1}^{\infty} \rho_i^{2\beta} f_{0,i}^2 \lesssim \epsilon^2, \\ \|h\|_{S^{\alpha+d/2}}^2 &= \sum_{i=1}^J \rho_i^{2\alpha+d} f_{0,i}^2 \leq \rho_J^{2\alpha+d-2\beta} \sum_{i=1}^{\infty} \rho_i^{2\beta} f_{0,i}^2 \lesssim \left(\frac{1}{\epsilon}\right)^{\frac{2\alpha-2\beta+d}{p+\beta}}. \end{aligned}$$

This gives the first term of the lower bound.

The centred small probability can be bounded using direct computation or from entropy bound of the unit ball of  $S^{\alpha+\delta/2}$  in  $S^{-p}$ .  $\square$

**Lemma 2.D.5.** *Let  $\mathbb{G} = (\rho_i^{-\alpha} i^{-\frac{1}{2}} c_i Z_i)$ , for  $Z_i \stackrel{i.i.d.}{\sim} N(0, 1)$  and constants  $c_i$  that are bounded away from 0 and  $\infty$ . For all  $\delta < \alpha$ , there exist constants  $a, b > 0$  so that, for every  $j \in \mathbb{N}$  and  $t > 0$ ,*

$$\Pr(\|R_j K - I\|_{S^\delta} > a \rho_j^{-(\alpha-\delta)} + t) \leq e^{-bt^2 j \rho_j^{2(\alpha-\delta)}}.$$

*Consequently, for every  $b_1 > 0$ , there exists  $a_1$  such that  $\Pr(\|R_j K - I\|_{S^\delta} > a_1 \rho_j^{-(\alpha-\delta)}) \leq e^{-b_1 j}$ , for every  $j \in \mathbb{N}$ .*

*Proof.* By Borell's inequality, a centred Gaussian variable  $G$  in a separable Banach space satisfies  $\Pr(\|G\| \gtrsim \mathbb{E}\|G\| + t) \leq e^{-bt^2/\sigma^2}$ , for every  $t > 0$ , where  $\sigma^2 = \sup_{b^*} \mathbb{E}(b^* G)^2$  is the supremum of the second moments of the variables  $b^* G$  if  $b^*$  ranges over the unit ball of the dual space of the Banach space. The variable  $G = (R_j K - I)\mathbb{G}$  is centred Gaussian in  $S^\delta$ . Thus it suffices to show that  $\mathbb{E}\|(R_j K - I)\mathbb{G}\|_{S^\delta} \lesssim \rho_j^{-(\alpha-\delta)}$  and  $\mathbb{E}\langle (R_j K - I)\mathbb{G}, g \rangle_{S^\delta}^2 \lesssim j^{-1} \rho_j^{-2(\alpha-\delta)} \|g\|_{S^\delta}^2$ , for every  $g \in S^\delta$ . Here we can bound the first strong moment by the root of the second strong moment.

We have  $\langle \mathbb{G}, f \rangle_{S^\delta} = \sum_{i \in \mathbb{N}} \rho_i^{-\alpha} i^{-\frac{1}{2}} c_i Z_i f_i \rho_i^{2\delta}$  and hence

$$\mathbb{E}\langle \mathbb{G}, f \rangle_{S^\delta}^2 = \sum_{i \in \mathbb{N}} \rho_i^{-2\alpha+4\delta} i^{-1} c_i^2 f_i^2 = \langle f, C_{2(\delta-\alpha)} D^2 f \rangle_{S^\delta},$$

for  $(Df)_i = f_i i^{-\frac{1}{2}} c_i$ . Therefore, as a map into  $S^\delta$  the variable  $(R_j K - I)\mathbb{G}$  has covariance operator  $(R_j K - I) C_{2(\delta-\alpha)} D^2 (R_j K - I)^* = S S_\delta^*$ , for  $S = (R_j K - I) C_{\delta-\alpha} D$ , and  $(R_j K - I)_\delta^*$  and  $S_\delta^*$  the adjoint operators of  $(R_j K - I) : S^\delta \rightarrow S^\delta$  and  $S_\delta : S^\delta \rightarrow S^\delta$ . (The operators  $C_s : S^\delta \rightarrow S^\delta$ , for  $s \leq 0$ , and  $D^2 : S^\delta \rightarrow S^\delta$  are self-adjoint, and commute, with roots  $C_{s/2}$  and  $D$ .)

The strong second moment is equal to the trace of the covariance operator. Therefore, for any orthonormal basis  $(\phi_i)$  of  $S^\delta$ ,

$$\begin{aligned} & \mathbb{E}\|(R_j K - I)\mathbb{G}\|_{S^\delta}^2 \\ &= \text{tr}(S S_\delta^*) = \sum_{i \in \mathbb{N}} \langle S S_\delta^* \phi_i, \phi_i \rangle_{S^\delta} = \sum_{i \in \mathbb{N}} \|S_\delta^* \phi_i\|_{S^\delta}^2 \\ &= \sum_{i \in \mathbb{N}} \sum_{j \in \mathbb{N}} \langle S_\delta^* \phi_i, \phi_j \rangle_{S^\delta}^2 = \sum_{i \in \mathbb{N}} \sum_{j \in \mathbb{N}} \langle \phi_i, S \phi_j \rangle_{S^\delta}^2 = \sum_{j \in \mathbb{N}} \|S \phi_j\|_{S^\delta}^2. \end{aligned}$$

We apply this with the orthonormal basis consisting of the sequences  $\phi_i = e_i/\rho_i^\delta$ , for  $e_i = (0, \dots, 0, 1, 0, \dots)$  the sequence with a 1 in the  $i$ th coordinate and zeros elsewhere, to see that this is equal to

$$\begin{aligned} \sum_{i \in \mathbb{N}} \|(R_j K - I)C_{\delta-\alpha} D e_i \rho_i^{-\delta}\|_{S^\delta}^2 &= \sum_{i \in \mathbb{N}} \|(R_j K - I)e_i\|_{S^\delta}^2 \rho_i^{-2\alpha} i^{-1} c_i^2 \\ &\lesssim \sum_{i > j} \|R_j K - I\|_{S^\delta \rightarrow S^\delta}^2 \|e_i\|_{S^\delta}^2 \rho_i^{-2\alpha} i^{-1} \leq \sum_{i > j} \rho_i^{-2(\alpha-\delta)} i^{-1}, \end{aligned}$$

since  $R_j K - I$  vanishes on  $V_j$ , the norm of  $R_j K - I : S^\delta \rightarrow S^\delta$  is bounded by a constant by (2.51), and  $\|e_i\|_{S^\delta} = \rho_i^\delta$ . For  $\delta < \alpha$ , the right side is bounded by a multiple of  $\rho_j^{-2(\alpha-\delta)}$ .

The weak second moment is the supremum over  $\|g\|_{S^\delta} \leq 1$  of

$$\begin{aligned} \mathbb{E} \langle (R_j K - I) \mathbb{G}, g \rangle_{S^\delta}^2 &= \langle g, S S_\delta^* g \rangle_{S^\delta} = \|S_\delta^* g\|_{S^\delta}^2 = \|C_{\delta-\alpha} D (R_j K - I)_\delta^* g\|_{S^\delta}^2 \\ &\lesssim \|(R_j K - I)_\delta^* g\|_{S^{2\delta-\alpha-d/2}}^2 = \sup_{\|f\|_{S^{\alpha+d/2}} \leq 1} \langle f, (R_j K - I)_\delta^* g \rangle_{S^\delta}^2, \end{aligned}$$

by Equation (2.45), applied with  $t - s = \delta + (\delta - \alpha - d/2)$  and  $t + s = \delta - (\delta - \alpha - d/2) = \alpha + d/2$ . The right side is equal to

$$\begin{aligned} &\sup_{\|f\|_{S^{\alpha+d/2}} \leq 1} \langle (R_j K - I) f, g \rangle_{S^\delta}^2 \\ &\leq \sup_{\|f\|_{S^{\alpha+d/2}} \leq 1} \|(R_j K - I) f\|_{S^\delta}^2 \|g\|_{S^\delta}^2 \\ &\leq \sup_{\|f\|_{S^{\alpha+d/2}} \leq 1} \|R_j K - I\|_{S^{\alpha+d/2} \rightarrow S^\delta}^2 \|f\|_{S^{\alpha+d/2}}^2 \|g\|_{S^\delta}^2 \\ &\leq \rho_j^{-2(\alpha+d/2-\delta)} \|g\|_{S^\delta}^2, \end{aligned}$$

by Equation (2.51). Therefore the weak second moment is bounded from above by  $\rho_j^{-2(\alpha+d/2-\delta)} \lesssim \rho_j^{-2(\alpha-\delta)} j^{-1}$ .

This concludes the proof of the first probability bound. The second bound follows by choosing  $t = \sqrt{b_1/b} \rho_j^{-(\alpha-\delta)}$  in the first bound.  $\square$

## 2.E Smoothness norm contraction rates in the discrete observation linear inverse problems

In this section, we derive similar results as in the previous section, but in the context of the discrete observational framework. We extend the contraction rate



results derived in Chapter 8 of (Yan 2020) or (Yan and van der Vaart 2023) to smoothness norms.

The setup is the same as in the previous section, and so is the general aim: to recover a function  $f$  from the noisy observation of its image  $Kf$  under a given operator  $K : S^0 \rightarrow L$ . However, presently the Hilbert space  $L$  is a space of functions on a domain  $\mathcal{O}$ , and for fixed points  $x_{N,1}, \dots, x_{N,N}$  in  $\mathcal{O}$ , the observations are  $Y_N(1), \dots, Y_N(N)$  given by, for i.i.d. standard normal variables  $Z_i$ ,

$$Y_N(i) = Kf(x_{N,i}) + Z_i, \quad i = 1, \dots, N, \quad (2.52)$$

This problem will be tackled by putting it in the continuous setup, using an interpolation scheme.

Define the empirical inner product on the design points by

$$\langle f, g \rangle_N := \frac{1}{N} \sum_{i=1}^N g(x_{N,i}) h(x_{N,i}) \quad (2.53)$$

and let  $\|f\|_N := \sqrt{\langle f, f \rangle_N}$  be the induced norm. We assume that, for every  $N \in \mathbb{N}$ , there exists an  $N$ -dimensional subspace  $\mathbb{L}_N \subset L$  such that, for fixed constants  $0 < C_1 < C_2 < \infty$ ,

$$C_1 \|w\|_L \leq \|w\|_N \leq C_2 \|w\|_L, \quad w \in \mathbb{L}_N. \quad (2.54)$$

Furthermore, we assume that for every  $f \in L$  the unique element  $\mathcal{I}_N f$  of  $\mathbb{L}_N$  that interpolates  $f$  at the design points, i.e.  $\mathcal{I}_N f(x_i) = f(x_i)$ , for every  $i = 1, \dots, N$ , satisfies, for every  $s$  in some interval  $(s_d, S_d)$ , and some sequence  $\delta_{N,s} = o(1)$ ,

$$\|\mathcal{I}_N f - f\|_L \lesssim \delta_{N,s} \|f\|_{S^s}. \quad (2.55)$$

Fix an arbitrary orthonormal basis  $e_{1,N}, \dots, e_{N,N}$  of  $\mathbb{L}_N$  relative to  $\langle \cdot, \cdot \rangle_N$  and given the observations  $Y_N(1), \dots, Y_N(N)$  as in (2.52), define

$$Y^{(N)} = \sum_{i=1}^N \frac{1}{N} \sum_{j=1}^N Y_N(j) e_{N,i}(x_{N,j}) e_{N,i}. \quad (2.56)$$

This variable contains the same information as the original observational data  $Y_N(1), \dots, Y_N(N)$  (and thus will lead to the same posterior distribution), and embeds the discrete data as a “continuous signal” into the space  $\mathbb{L}_N \subset L^2$ . The

variable can be decomposed as

$$\begin{aligned} Y^{(N)} &= \sum_{i=1}^N \langle Kf, e_{N,i} \rangle_N e_{N,i} + \frac{1}{N} \sum_{j=1}^N Z_j \sum_{i=1}^N e_{N,i}(x_{N,j}) e_{N,i} \\ &= \mathcal{I}_N Kf + \frac{1}{\sqrt{N}} \xi^{(N)}, \end{aligned} \quad (2.57)$$

where  $\xi^{(N)}$  is a Gaussian random variable with values in the space in  $\mathbb{L}_N$ . This formulation connects the regression model to the Gaussian white noise model and allows the application of the techniques developed for this idealised model.

To simplify the notation, we write  $K_N = \mathcal{I}_N K$ . Consider the same random sequence prior  $\mathbb{G} = (\rho_i^{-\alpha} i^{-\frac{1}{2}} c_i Z_i)$  as in Section 2.D, with  $Z_i \stackrel{\text{i.i.d.}}{\sim} N(0, 1)$  and constants  $c_i$  such that  $c_\pi \leq c_i \leq C_\pi$ , for some constants  $0 < c_\pi < C_\pi$ .

**Theorem 2.E.1.** *Assume that  $c_0 i^{1/d} \leq \rho_i \leq C_0 i^{1/d}$ , for constants  $C_0 \geq c_0 > 0$  and  $d > 0$ . Let  $K : S^0 \rightarrow L$  be a linear operator satisfying Equation (2.46) for every  $f \in S^0$ . If  $f_0 \in S^\beta$ , for  $d/2 < \beta < B$  and  $\delta < \alpha \wedge \beta$ , and (2.55)–(2.56) hold with  $\delta_{N,\beta} = c_2 N^{-\beta/(2\alpha+2p+d)}$ , then there exists a constant  $M$  such that  $\mathbb{E}_{f_0} \Pi_N(f : \|f - f_0\|_{S^\delta} \leq M N^{-(\alpha \wedge \beta - \delta)/(2\alpha+2p+d)} \mid Y^{(N)}) \rightarrow 1$ .*

*Proof.* Similarly to the proof of of Theorem 2.D.2 we show below that we can choose  $c_1$  and  $c_2$  so that for  $\varepsilon_N = N^{-(\alpha \wedge \beta - \delta)/(2\alpha+2p+d)}$

- (1\*) The events  $A_N := \left\{ \int p_f^{(N)} / p_{f_0}^{(N)}(Y^{(N)}) d\Pi(f) \geq e^{-2N\epsilon_N^2} \right\}$  satisfy asymptotically  $P_{f_0}^{(N)}(A_N^c) \rightarrow 0$ .
- (2\*) There exists  $M_1 > 0$  such that  $\Pi(\|Kf - Kf_0\|_N < M_1 \epsilon_N) \geq e^{-N\epsilon_N^2}$ .
- (3\*) There exists  $M_2 > 0$  such that  $\Pi(\|R_{j_N} K_N f - f\|_{S^\delta} > M_2 \eta_N) \leq e^{-4N\epsilon_N^2}$ .
- (4\*) There exist  $M > 0$  and tests  $\tau_N$  such that  $P_{f_0}^{(N)} \tau_N \rightarrow 0$  and  $P_f^{(N)}(1 - \tau_N) \leq e^{-4N\epsilon_N^2}$ , for every  $f \in \mathcal{F}_N$ , where

$$\mathcal{F}_N = \{f : \|f - f_0\|_{S^\delta} > M \eta_N, \|R_{j_N} K_N f - f\|_{S^\delta} < M_2 \eta_N\}.$$

Then we note that  $\Pi_N(f : \|f - f_0\|_{S^\delta} > M \eta_N \mid Y^{(N)})$  is bounded by

$$\begin{aligned} &\tau_N + 1_{A_N^c} + \Pi(\|R_{j_N} K_N f - f\|_{S^\delta} \geq M_2 \eta_N \mid Y^{(N)}) \\ &+ \Pi_N(\mathcal{F}_N \mid Y^{(N)}) 1_{A_N}(1 - \tau_N). \end{aligned}$$

Then in view of (1\*)–(4\*) the  $P_{f_0}^{(N)}$ -expectation of each term in the preceding display tends to zero following standard arguments, as in the proof of Theorem 2.D.2.

It remained to verify the above four assumptions. The Kullback-Leibler divergence and variation between the (multivariate-normal) distributions of  $(Y_1, \dots, Y_N)$  under two functions  $f$  and  $f_0$  are given by  $N\|Kf - Kf_0\|_N^2/2$  and twice this quantity, respectively. Hence the evidence lower bound (1\*) holds by the standard arguments as in Lemma 8.21 of (Ghosal and van der Vaart 2017).

Next, we verify Claim (2\*). In view of Lemma 2.E.2,  $\|Kf - Kf_0\|_N \leq C_2\|K_N f - K_N f_0\|_L$ , for some  $C_2 > 0$ . By several applications of the triangle inequality, since  $d/2 < \beta + p < B + p$  and  $\|Kf\|_{S^{\beta+p}} \asymp \|f\|_{S^\beta}$ , and using Assumptions (2.55) and (2.46)

$$\begin{aligned} & \|K_N f - K_N f_0\|_L \\ & \leq \|K_N f - Kf\|_L + \|K_N f_0 - Kf_0\|_L + \|Kf - Kf_0\|_L \\ & \leq \delta_{N,\beta+p}(\|Kf\|_{S^{\beta+p}} + \|Kf_0\|_{S^{\beta+p}}) + \|Kf - Kf_0\|_L \\ & \leq C_1 \delta_{N,\beta+p}(\|f\|_{S^\beta} + \|f_0\|_{S^\beta}) + \|Kf - Kf_0\|_L \leq C_3 \epsilon_N + \|Kf - Kf_0\|_L, \end{aligned}$$

for some  $C_2, C_3 > 0$ , since  $\delta_{N,\beta+p} = o(\epsilon_N)$ . Therefore, the left-hand side of (2\*) can be bounded from below using triangle inequality and Claim (2) from Theorem 2.D.2 as

$$\Pi(\|Kf - Kf_0\|_N < M_1 \epsilon_N) \geq \Pi(\|Kf - Kf_0\|_L < (\frac{M_1}{C_2} - C_3) \epsilon_N) \geq e^{-4N\epsilon_N^2},$$

for  $M_1 \geq C_2(C_3 + 1)$ .

Then, we prove assertion (3\*). First, by the triangle inequality, we get

$$\|R_{j_N} K_N f - f\|_{S^\delta} \leq \|R_{j_N} Kf - f\|_{S^\delta} + \|R_{j_N} Kf - R_{j_N} K_N f\|_{S^\delta}.$$

Then in view of assertion (2.49), the second term on the right-hand side can be further bounded from above by

$$\begin{aligned} C_0 \rho_{j_N}^\delta \|R_{j_N} Kf - R_{j_N} K_N f\|_{S^0} & \leq C_1 \rho_{j_N}^\delta \|R_{j_N}\|_{L \mapsto S^0} \|K_N f - Kf\|_L \\ & \leq C_2 \rho_{j_N}^{\delta+p} \delta_{N,\beta+p} \|Kf\|_{S^{\beta+p}} \leq C_3 \eta_N, \end{aligned}$$

for some constants  $C_i > 0, i = 0, \dots, 3$ , since  $\rho_{j_N}^{\delta+p} \delta_{N,\beta+p} \lesssim N^{\frac{\delta+p-\beta-p}{1+2\alpha+2p}} \leq \eta_N$ . Therefore, in view of Claim (3) from Theorem 2.D.2, the probability on the left-hand side of (3\*) is bounded from above, for  $M_2 \geq C_3 + 1$  by

$$\begin{aligned} \Pi(\|R_{j_N} K_N f - f\|_{S^\delta} > M_2 \eta_N) & \leq \Pi(\|R_{j_N} Kf - f\|_{S^\delta} > (M_2 - C_3) \eta_N) \\ & \leq e^{-4N\epsilon_N^2}, \end{aligned}$$

concluding the proof of the theorem.

Finally, we show that the test

$$\tau_N = 1 \{ \|R_{j_N} Y^{(N)} - R_{j_N} K_N f_0\|_{S^\delta} \geq M_0 \eta_N \}, \quad (2.58)$$

where  $M_0 > 0$  is a given constant, to be determined later, satisfies Claim (4\*). For  $Y^{(N)}$  given in (2.57) we have

$$R_j Y^{(N)} = R_j K_N f + \frac{1}{\sqrt{N}} R_j \xi^{(N)}. \quad (2.59)$$

The variable  $R_j \xi^{(N)} = R_j Q_j \xi^{(N)}$  is a centred Gaussian random element in  $V_j$  with strong and weak second moments,

$$\begin{aligned} \mathbb{E} \|R_j \xi^{(N)}\|_{S^0}^2 &\leq \|R_j\|_{L \mapsto S^0}^2 \mathbb{E} \|Q_j \xi^{(N)}\|_L^2 \lesssim \rho_j^{2p} j, \\ \sup_{\|f\|_L \leq 1} \mathbb{E} \langle R_j \xi^{(N)}, f \rangle_{S^0}^2 &= \sup_{\|f\|_L \leq 1} \mathbb{E} \langle \xi^{(N)}, R_j^* f \rangle_L^2 \\ &\lesssim \sup_{\|f\|_L \leq 1} \|R_j^* f\|_L^2 \leq \|R_j^*\|_{L \mapsto S^0}^2 \lesssim \rho_j^{2p}. \end{aligned}$$

In both cases the inequality on  $\|L \mapsto R_j\|_{S^0} = \|R_j^*\|_{L \mapsto S^0}$  at the right-hand side follows from Equation (2.49), and we also use that, the covariance operator of  $\xi^{(N)}$  is bounded above by a multiple of the projection onto  $\mathbb{L}_N$ , and hence the identity, so that the covariance operator of  $Q_j \xi^{(N)}$  is bounded above by a multiple of  $Q_j$ .

The first inequality shows that the first moment  $\mathbb{E} \|R_j \xi^{(N)}\|_{S^0}$  of the variable  $\|R_j \xi^{(N)}\|_{S^0}$  is bounded above by  $\sqrt{j} \rho_j^p$ . Since  $R_j \xi^{(N)} \in V_j$ , the  $S^\delta$ -norm version of Borell's inequality, see (2.47), applied to the Gaussian random variable  $R_j \xi^{(N)}$  in  $S^0$ , implies that there exist positive constants  $a$  and  $b$  such that, for every  $t > 0$ ,

$$\Pr(\|R_j \xi_j\|_{S^\delta} > a \rho_j^{p+\delta} \sqrt{j} + \rho_j^\delta t) \leq e^{-bt^2/\rho_j^{2p}}.$$

Choosing  $j = j_N$  and  $t^2 = 4N\epsilon_N^2 \rho_{j_N}^{2p}/b$ , the right side to  $e^{-4N\epsilon_N^2}$ . Furthermore, as in the computations above Equation (2.47), both terms on the right-hand side in the probability are bounded from above by  $\sqrt{N} \eta_N$ , resulting in for a sufficiently large constant  $M_0$  that

$$P_f^{(N)}(\|R_{j_N} Y^{(N)} - R_{j_N} K f\|_{S^\delta} > M_0 \eta_N) \leq e^{-4N\epsilon_N^2} = o(1).$$

Next, we give an upper bound for the type-2 error. First note that in view of assertion (2.49) and assumption (2.55),

$$\begin{aligned} \|R_{j_N} K f_0 - R_{j_N} K_N f_0\|_{S^\delta} &\leq \rho_{j_N}^\delta \|R_{j_N} K f_0 - R_{j_N} K_N f_0\|_{S^0} \\ &\leq \rho_{j_N}^\delta \|R_{j_N}\|_{L \mapsto S^0} \|K_N f_0 - K f_0\|_L \\ &\lesssim \rho_{j_N}^{\delta+p} \delta_{N,\beta+p} \|K f_0\|_{S^{\beta+p}} \lesssim \eta_N. \end{aligned}$$

Then by triangle inequality and (2.48) there exists a large enough  $c_0 > 0$  such that

$$\|f_0 - R_{j_N} K_N f_0\|_{S^\delta} \leq \|f_0 - R_{j_N} K f_0\|_{S^\delta} + \|R_{j_N} K f_0 - R_{j_N} K_N f_0\|_{S^\delta} \leq c_0 \eta_N.$$

It follows that  $\|R_{j_N} K f - f\|_{S^\delta} < M_2 \eta_N$  and  $\|f - f_0\|_{S^\delta} > M \eta_N$  imply

$$\|R_{j_N} K f - R_{j_N} K f_0\|_{S^\delta} \geq \|f - f_0\|_{S^\delta} - M_2 \eta_N - c_0 \eta_N \geq (M - M_2 - c_0) \eta_N.$$

Then, for  $\tau_N = 0$ ,

$$\begin{aligned} \|R_{j_N} Y^{(N)} - R_{j_N} K f\|_{S^\delta} &\geq (M - M_2 - c_0) \eta_N - \|R_{j_N} Y^{(N)} - R_{j_N} K f_0\|_{S^\delta} \\ &\geq (M - M_2 - c_0 - M_0) \eta_N. \end{aligned}$$

It follows that  $P_f^{(N)}(1 - \tau_N) \leq \Pr(\|R_{j_N} Y^{(N)} - R_{j_N} K f\|_{S^\delta} \geq c_5 \eta_N)$ , for  $c_5 = M - M_2 - c_0 - M_0$ , whenever  $f \in \mathcal{F}_N$ . For  $M \geq M_2 + c_0 + 2M_0$ , the latter probability is bounded above by  $e^{-4N\epsilon_N^2}$ , concluding the proof of the theorem.  $\square$

The next lemma is Lemma 8.3 of (Yan 2020) or Lemma 2.9 of (Yan and van der Vaart 2023).

**Lemma 2.E.2.** *Suppose that Equation (2.17) holds. Then*

- (i) *For every  $g \in L$  the function  $\mathcal{I}_N g$  is the orthogonal projection of  $g$  onto  $\mathbb{L}_N \subset L$  relative to the discrete inner product  $\langle \cdot, \cdot \rangle_N$ .*
- (ii)  *$C_1 \|\mathcal{I}_N g\|_L \leq \|g\|_N \leq C_2 \|\mathcal{I}_N g\|_L$ , for every  $g \in L$ .*
- (iii) *The Gram matrix  $(\langle e_{N,i}, e_{N,j} \rangle_N)_{i,j=1,\dots,N}$  of any basis  $e_{1,N}, \dots, e_{N,N}$  of  $\mathbb{L}_N$  that is orthonormal relative to the discrete inner product  $\langle \cdot, \cdot \rangle_N$  is bounded below and above by the identity, up to multiplicative constants that do not depend on  $N$ .*

# Chapter 3

## Distributed Inverse Problems

### 3.1 Introduction

In statistical inverse problems, one does not observe a noisy version of a parameter of interest, but only a noise-corrupted transformation of this parameter. The infinite-dimensional nature of the statistical model, combined with the smoothing property of inverse problems, creates difficulties in recovering the truth. Bayesian statistics provides a well-founded method of handling these problems, since regularization can be handled by taking a well-chosen prior, and the inverse nature is naturally carried over from the prior into the posterior. Priors coming from Gaussian processes have been extensively tested in machine learning and applied statistical research. As the theoretical analysis of these priors is comparatively easy, many advancements have been made in recent decades for both regression and inverse problem settings. However, inverse problems in particular often require a large number of observations. This can slow down the computation time for Gaussian posterior distributions, as the explicit expression for the posterior involves the inversion of a large matrix. Distributed methods are an attractive way of combatting these issues, as they are straightforward to implement and have shown good results in other settings (Szabó and van Zanten 2019; Hadji 2023).

We assume that the observed data is of the form

$$Y_i = (K f_0)(X_i) + \epsilon_i, \quad \epsilon_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), X_1, \dots, X_N \stackrel{\text{i.i.d.}}{\sim} U(\mathcal{O}). \quad (3.1)$$

The noise  $\epsilon_i$  and the design points  $X_1, \dots, X_N$  are assumed to be independent. Here  $f_0$  is an element of the  $L^2(\mathcal{O})$ -space on a  $d$ -dimensional domain  $\mathcal{O}$ , and  $K$  is a known mapping  $K : L^2 \rightarrow L^2$ . Using the data  $\mathbb{D}_N := \{(X_i, Y_i) : i =$

$1, \dots, N\}$ , we wish to recover  $f_0$ . The goal of this chapter is to quantify how well a Bayesian procedure can recover the truth in a distributed setting.

Although the analytically more tractable white noise model is statistically equivalent to the discrete observational model as stated in Equation (3.1), in practice one only observes the latter data. In Bayesian computations, this is often a real concern, as the posterior can only be analytically computed by performing a matrix inversion. When the data consists of  $N$  data points, this computation scales as  $O(N^3)$ . This quickly becomes impractical even in regression settings, and becomes worse when applied to inverse problems. In inverse problems, often large datasets are required to get a decent recovery of the parameters of interest. Therefore, besides statistical considerations, practitioners are often interested in the computational performance of the method in question. A natural way of combatting this problem is by splitting the observed data up into smaller packets, and allocating these datasets to different machines. One can then do independent, parallel computations on smaller datasets, and when all machines have finished, aggregate these into a final statistical estimator using a central computing unit. This is the fundamental idea of distributed computing and it has gained significant attention in the past few years. Since the commonly applied algorithms used for matrix-inversion are of the order  $O(N^3)$ , by using  $m$  machines sequentially, the computation time can be cut down to  $mO(m^{-3}N^3)$ , allowing for a great reduction of computation time even when one does not use parallel computation. When  $m$  is chosen too big compared to  $N$ , the contraction rate of the posterior deteriorates polynomially below the minimax-rate. Therefore, a proper choice of  $m$  is paramount to recovering the truth in an optimal way.

In recent years, non-linear inverse problems have gained a lot of attention due to their importance in many real-world applications. In particular non-linear partial differential equations have been analysed on a case by case basis, each requiring significant analysis of PDE theory and its connection to statistical theory. Moreover, the methods often do not allow for easy computation of the posterior or posterior mean. This means that often one is restricted to using some sort of MCMC scheme, which computationally can be very slow to converge. We apply the results in this chapter to the techniques that apply methods for linear inverse problems to solve non-linear PDEs, see Chapter 2.

The chapter starts with the mathematical setup of the problem in Section 3.2. In Section 3.3.1, we state the contraction rates for inverse problems for operators with polynomially decaying eigenvalues. The main result regarding the uncertainty quantification is placed in Section 3.3.2. The connection to non-linear inverse problems can be found in Section 3.4. Then in Section 3.5 the previous

parts are illustrated by a numerical simulation study and computation time analysis. Finally, some derivations of important mappings are put in Section 3.6, and the proofs of the contraction rates and uncertainty quantification are given in Section 3.7 and Section 3.10 respectively. The technical lemmas can be found in Section 3.8.

## 3.2 Setup

Let  $\mathcal{O} \subset \mathbb{R}^d$ ,  $d \in \mathbb{N}$ , be an arbitrary measurable domain of unit volume. Let  $K : L^2(\mathcal{O}) \rightarrow L^2(\mathcal{O})$  be a compact, injective, self-adjoint, positive linear operator. Let  $X_i \stackrel{\text{i.i.d.}}{\sim} U(\mathcal{O})$  be uniformly distributed over the domain. We observe  $(X_1, Y_1), \dots, (X_N, Y_N)$  as defined in Equation (3.1) and wish to recover  $f_0 \in L^2(\mathcal{O})$ . The assumptions on  $K$  imply that  $K$  has an orthonormal eigenbasis consisting of eigenfunctions  $(\phi_j)_{j \in \mathbb{N}} \subset L^2(\mathcal{O})$  with corresponding eigenvalues  $\kappa_j > 0$ , satisfying  $K\phi_j = \kappa_j\phi_j$ . We assume that the eigenvalues asymptotically decay as  $\kappa_j \asymp j^{-\frac{p}{d}}$  for some  $p > 0$ . This corresponds to a mildly ill-posed inverse problem as defined in (Cavalier 2008). Since the spatial domain  $\mathcal{O}$  remains fixed throughout the remainder of the chapter, it is often dropped from the notation.

The smoothness of the parameter is qualified by the Sobolev space norm with respect to the  $L^2$ -basis  $(\phi_i)_{i \in \mathbb{N}}$ , so

$$\|f\|_{S^\gamma}^2 = \sum_{j=1}^{\infty} f_j^2 j^{\frac{2\gamma}{d}},$$

with  $\gamma > 0$  arbitrary. The space consisting of elements of  $L^2$  that have finite  $\|\cdot\|_{S^\gamma}$  norm is denoted by  $S^\gamma$ :

$$S^\gamma := \left\{ f \in L^2 : \sum_{j=1}^{\infty} f_j^2 j^{\frac{2\gamma}{d}} < \infty \right\}.$$

For the orthonormal eigenbasis of the Laplacian, elements of  $S^\gamma$  can be thought of as a subset of the elements of  $L^2$  that have  $\gamma$  weak-derivatives in  $\mathcal{O}$  (Taylor 2011). A ball of radius  $B \geq 0$  centred at  $0 \in S^\gamma$  is denoted by

$$S^\gamma(B) := \{f \in S^\gamma : \|f\|_{S^\gamma} \leq B\}.$$

In the rest of the chapter, it is assumed that  $f_0 \in S^\beta(B)$  for some fixed  $\beta > \frac{d}{2}$  and  $B > 0$ .  $L^2$  and  $S^\gamma$  can be canonically connected through the isomorphism



$G_\gamma : S^\gamma \rightarrow L^2$ , which is defined by

$$G_\gamma \sum_{j=1}^{\infty} f_j \phi_j := \sum_{j=1}^{\infty} j^{\frac{\gamma}{d}} f_j \phi_j. \quad (3.2)$$

We will assume some conditions on the eigenfunctions. Define the Orlicz norm  $\|\cdot\|_{\psi_2}$  of a random variable  $Z$  as

$$\|Z\|_{\psi_2} := \inf \left\{ C > 0 : \mathbb{E}_Z(e^{Z^2/C^2}) \leq 2 \right\}.$$

Throughout the rest of the chapter it is assumed that the eigenfunctions possess the following properties.

**Assumption 3.2.1.** *There exists a constant  $C > 0$  such that the eigenfunctions  $(\phi_j)_{j \in \mathbb{N}}$  satisfy*

$$\limsup_{N \rightarrow \infty} \sup_{\ell, j \in \mathbb{N}} \left| \frac{1}{N} \sum_{i=1}^n \phi_\ell(X_i) \phi_j(X_i) \right| < C \quad \text{a.s.}, \quad (3.3)$$

$$\sup_{\ell, j \in \mathbb{N}} \|\phi_j(X) \phi_\ell(X)\|_{\psi_2} < \infty, \quad (3.4)$$

where  $X \sim \text{Unif}(\mathcal{O})$ . The first inequality is assumed to hold  $\otimes_{i=1}^{\infty} P_{X_i}$ -outer almost surely.

In particular, a sufficient condition on the eigenfunctions is that there exists a constant  $C_\phi > 0$  such that for all  $j \in \mathbb{N}$ , we have  $\|\phi_j\|_\infty \leq C_\phi$ . Assumption (3.3) can be slightly weakened by replacing the supremum over  $\ell, j \in \mathbb{N}$  by the supremum over  $\ell, j \in \{k \in \mathbb{N} : k > N^c\}$ , for some fixed constant  $c > 0$  small enough (see the proof of Lemma 3.7.3).

In the non-distributed setting, we equip  $f$  with a Gaussian prior of the form  $\Pi_\alpha = GP(0, C_\alpha)$ , where  $C_\alpha$  is a positive definite kernel given by

$$C_\alpha(s, t) = \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha)} \phi_j(s) \phi_j(t)$$

for some  $\alpha > 0$ . The smoothness of the prior is denoted by  $\alpha$ , so that prior variance decays as  $\kappa_j^{\frac{1}{p}(d+2\alpha)} \asymp j^{-1-\frac{2\alpha}{d}}$ . Throughout the chapter it is assumed that the smoothness  $\beta$  of  $f \in S^\beta$  is lower bounded by  $\beta > \frac{d}{2}$  and  $\gamma < \frac{d}{2} + 2\alpha + p$ .

### 3.2.1 Distributed computing

In the distributed setting, we denote the number of local machines by  $m \in \{1, \dots, N\}$ , each of which we give a dataset with  $n = \frac{N}{m} \in \mathbb{N}$  data points, which is denoted by  $\mathbb{X}_n^{(k)} := \{X_1^{(k)}, \dots, X_n^{(k)}\}$ ,  $\mathbb{Y}_n^{(k)} := \{Y_1^{(k)}, \dots, Y_n^{(k)}\}$  and  $\mathbb{D}_n^{(k)} = \{(X_i^{(k)}, Y_i^{(k)}) : i = 1, \dots, n\}$ . The parts  $\mathbb{D}_n^{(k)}$ ,  $k = 1, \dots, m$ , form a disjoint partition of the data  $\mathbb{D}_N$ , so that  $\cup_{k=1}^m \mathbb{D}_n^{(k)} = \mathbb{D}_N$ . To compensate for the fact that each local machine only has access to partial data, following (Scott et al. 2016), the covariance kernel on the local machines are multiplied by a factor of  $m$ . This scaling is equivalent to raising the local prior density to the power  $1/m$ , which can be interpreted as keeping the total prior information in the distributed setup equal to the total prior information in the undistributed setting. On each local machine  $k$ , we compute the local posterior distribution  $\Pi_{n,m}^{(k)}$ , with the prior as described and its corresponding dataset. The local posterior mean is denoted by  $\hat{f}_{n,m}^{(k)}$  and its covariance by  $\hat{C}_\alpha^{(k)}$ , so that  $\Pi_{n,m}^{(k)} = N(\hat{f}_{n,m}^{(k)}, \hat{C}_\alpha^{(k)})$ . If for each  $k = 1, \dots, m$  one has a draw  $\tilde{f}_{n,m}^{(k)}$  from the local posterior  $\Pi_{n,m}^{(k)}$ , a draw from the aggregated posterior  $\Pi_{n,m}^\dagger$  is defined as the average  $\frac{1}{m} \sum_{k=1}^m \tilde{f}_{n,m}^{(k)}$  of the draws from the local posteriors. Therefore, the aggregated posterior has a Gaussian distribution as well, with mean  $\hat{f}_{n,m}^\dagger := \frac{1}{m} \sum_{k=1}^m \hat{f}_{n,m}^{(k)}$  and covariance matrix  $\frac{1}{m^2} \sum_{k=1}^m \hat{C}_\alpha^{(k)}$ . Earlier results applying a combination of scaling local priors and aggregating them by averaging can be found in (Shang and Cheng 2015), (Scott 2017), (Szabó and van Zanten 2019) and (Hadji, Hesselink and Szabó 2022).

## 3.3 Main results

### 3.3.1 Contraction rates

The statistical method as described in Section 3.2 is consistent for the unknown parameter  $f_0$  as  $N \rightarrow \infty$ , as is shown in the first theorem below. The proof can be found in Section 3.7.1.

**Theorem 3.3.1.** *Assume that  $\kappa_j \asymp j^{-\frac{p}{d}}$  for some  $p > 0$  and that the eigenfunctions satisfy Assumptions (3.3) and (3.4). Let  $\alpha = \beta$ ,  $0 \leq \gamma < \beta$  be arbitrary, and  $m \asymp N^\zeta$ , for some  $0 < \zeta < \frac{1}{3}$  and let  $f_0 \in S^\beta(B)$  be such that  $\beta \geq 9(d + p)$ . Then for any sequence  $M_N \rightarrow \infty$ ,*

$$\sup_{f_0 \in S^\beta(B)} \mathbb{E}_0(\Pi_{n,m}^\dagger(\|f - f_0\|_{S^\gamma}^2 > M_N N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}} \mid \mathbb{D}_N)) \rightarrow 0. \quad (3.5)$$

When  $\gamma = 0$  this becomes the  $L^2$ -norm convergence rate and is equal to the minimax rate, as can be seen in (Cavalier 2008). The rate is also comparable to non-distributed settings in the white noise model (Knapik, van der Vaart and van Zanten 2011) and (Knapik, Szabó et al. 2016). In order to have a concise condition for attaining minimax contraction rates, we require the sufficient condition  $\beta \geq 9(d + p)$ . However, weaker conditions on  $\beta$  can be found in the proof in Section 3.7.

### 3.3.2 Uncertainty quantification

The posterior distribution provides a canonical way of measuring uncertainty in the estimation of the parameter of interest. While posterior contraction rates show that the posterior is consistent, it does not show whether the posterior spread is an accurate measure of the uncertainty. The standard way of gauging the correctness of the uncertainty as captured by the posterior distribution is by analysing the frequentist coverage of credible sets coming from this posterior distribution. One natural way of defining these credible sets is to take balls with respect to the  $\|\cdot\|_{L^2}$  norm that are centred at the posterior mean. For  $r \geq 0$ , define the ball centred at the posterior mean  $\hat{f}_{n,m}^\dagger$  as

$$B_r(L) := \{f \in L^2 : \|f - \hat{f}_{n,m}^\dagger\|_{L^2} \leq Lr\}.$$

Then for  $\tau \in (0, 1)$ , define  $r_{n,m,\tau}$  as the radius of the credible set  $B_{r_{n,m,\tau}}(1)$  such that it has a posterior mass equal to  $1 - \tau$ , hence  $r_{n,m,\tau}$  satisfies

$$\Pi(B_{r_{n,m,\tau}}(1) \mid \mathbb{D}_N) = 1 - \tau.$$

We are interested in the frequentist coverage of the credible sets  $B_{r_{n,m,\tau}}(L)$ , which is defined as

$$P_0(f_0 \in B_{r_{n,m,\tau}}(L)). \quad (3.6)$$

The main question is whether this probability is at least equal to  $1 - \tau$  when the sample size  $N$  grows to infinity, and the true parameter  $f_0$  stays fixed. In order to prove that these credible balls are also confidence sets, they are enlarged with a multiplication factor  $L > 1$ . The following theorem tells us that centred posterior balls, when slightly enlarged, asymptotically do provide good coverage. The proof can be found in Section 3.10

**Theorem 3.3.2.** *Let  $f_0 \in S^\beta(B)$  such that  $\beta \geq 9(d + p)$ . Take  $m \asymp N^\zeta$ , for some  $0 < \zeta < \frac{1}{3}$  and  $\alpha = \beta$ . Assume that there exists a constant  $C > 0$  such that  $\sup_{i \in \mathbb{N}} \|\phi_i\|_\infty \leq C$ . Then for any  $\tau \in (0, 1)$  there exists a constant  $L \geq 0$  such that*

$$\liminf_{N \rightarrow \infty} \inf_{f_0 \in S^\beta(B)} P_0(f_0 \in B_{r_{n,m,\tau}}(L)) \geq 1 - \tau.$$

Since the asymptotic probability that these blown-up credible sets contain the true parameter is not smaller than  $1 - \tau$  uniformly over  $S^\beta$ , these credible sets are honest confidence sets. A more precise analysis of the posterior credible sets is required to determine whether  $B_{r_{n,m,\tau}}(1)$  is a confidence set that contains the  $f_0$  with asymptotic probability  $1 - \tau$ .

### 3.4 Non-linear problems

We now extend the results in Section 3.3 to a class of non-linear inverse problems in the field of partial differential equations. We adapt the technique derived in Chapter 2 to a specific model. For a more comprehensive discussion of the method, we refer to the above chapter.

Consider a non-linear PDE

$$\begin{cases} \mathcal{L}u_f = c(f, u_f) & \text{on } \mathcal{O}, \\ u_f = g & \text{on } \Gamma \subseteq \partial\mathcal{O}. \end{cases} \quad (3.7)$$

The linear differential operator  $\mathcal{L}$ , the function  $c : \mathbb{R}^2 \rightarrow \mathbb{R}$  and the boundary function  $g : \Gamma \rightarrow \mathbb{R}$  are considered to be known. It is assumed that there exists a map  $e : L^2 \rightarrow L^2$  such that  $e(\mathcal{L}u_f) = f$ .

The observational data consists of

$$Y_{N,i} = u_{f_0}(X_i) + \frac{1}{\sqrt{N}}\epsilon_i, \epsilon_i \stackrel{\text{i.i.d.}}{\sim} N(0, \sigma^2), X_i \stackrel{\text{i.i.d.}}{\sim} U(\mathcal{O}), \quad (3.8)$$

for  $i = 1, \dots, N$  and some fixed function  $f_0 : \mathcal{O} \rightarrow \mathbb{R}$ . The noise  $\epsilon_i$  is assumed to be independent of the design variables  $X_i$ .

The standard approach in a Bayesian setting for non-linear inverse problems is to make use of an MCMC scheme to generate draws from the posterior. This method requires calculating the forward map multiple times between consecutive posterior draws and can therefore be computationally very demanding.

To speed up the computations, in Chapter 2 we propose the following linear inversion technique. Using the map  $e$ , one method of recovering  $f$  is by recovering  $\mathcal{L}u_{f_0}$  first, and then applying the map  $e$  to  $\mathcal{L}u_{f_0}$ . Under certain (continuity)

conditions on the mapping  $e$ , one can then find a good approximation of  $f_0$ , given that one has a good estimate for  $\mathcal{L}u_{f_0}$ . In order to do this, we assume that we can write Equation (3.8) as

$$Y_{N,i} = K\mathcal{L}u_{f_0}(X_i) + \tilde{g}(X_i) + \frac{1}{\sqrt{N}}\epsilon_i,$$

for  $i = 1, \dots, N$ . Here  $K$  is defined as

$$\begin{cases} \mathcal{L}Ku &= u & \text{on } \mathcal{O}, \\ Ku &= 0 & \text{on } \Gamma. \end{cases} \quad (3.9)$$

Note that in view of the linearity of the operator  $\mathcal{L}$ , the operator  $K$  is also linear and can be considered to be the inverse of  $\mathcal{L}$  under Dirichlet boundary conditions. The function  $\tilde{g} : \mathcal{O} \rightarrow \mathbb{R}$  is known, and encapsulates the boundary conditions of the non-linear PDE given in Equation (3.7). It satisfies

$$\begin{cases} \mathcal{L}\tilde{g} &= 0 & \text{on } \mathcal{O}, \\ \tilde{g} &= g & \text{on } \Gamma. \end{cases} \quad (3.10)$$

Now defining  $\tilde{Y}_{N,i} := Y_{N,i} - \tilde{g}(X_i)$  results in a linear inverse-problem.  $\mathcal{L}u_{f_0}$  is endowed with a prior as proposed in Section 3.3.

Similar to the linear inverse problem, the data of the non-linear inverse problem is distributed to  $m$  local machines. Each of these machines computes the local posterior for the  $\mathcal{L}u_f$ , which are then combined on a central computer into an aggregated posterior for  $\mathcal{L}u_f$ . Finally, through the map  $e$  this is then transformed into a posterior for  $f$ .

#### 3.4.1 Schrödinger equation

As a proof of concept, we consider the time-independent Schrödinger equation, which is studied in (Nickl 2018). Let the spatial domain be a compact set  $\mathcal{O} \subset \mathbb{R}^d$ , which is assumed to be  $C^\infty$ . Let  $u_f$  be the solution to

$$\begin{cases} \frac{1}{2}\Delta u_f = u_f & \text{on } \mathcal{O}, \\ u_f = g & \text{on } \partial\mathcal{O}. \end{cases} \quad (3.11)$$

We assume that  $f \geq 0$ . It can be proven that  $\inf_{x \in \mathcal{O}} u_f(x) > 0$  when  $g$  is such that  $\inf_{x \in \partial\mathcal{O}} g(x) > 0$  (see Theorem 1 in Chapter 6.4 in (Evans 2010)). Following

the method described in the previous section, we take  $\mathcal{L} = -\Delta$  as the Laplacian on  $\mathcal{O}$  with homogeneous Dirichlet boundary conditions. Therefore  $K$  is the inverse of minus the Laplacian on  $\mathcal{O}$  with homogeneous Dirichlet boundary conditions, and  $\tilde{g}$  is a function satisfying  $\Delta\tilde{g} = 0$  on  $\mathcal{O}$  and  $\tilde{g}|_{\partial\mathcal{O}} = g$ . The existence of  $K$  and  $\tilde{g}$  is proven in Lemma 2.4.1. The inversion mapping now becomes

$$e(\Delta u_f) = \frac{\Delta u_f}{2(K\Delta u_f + \tilde{g})}, \quad (3.12)$$

satisfying  $e(\Delta u_f) = f$ .

In order to control the supremum norm of the denominator in the previous display, we make use of a Sobolev embedding theorem. For  $\gamma \in \mathbb{N}$ , let  $H^\gamma$  be the space of  $\gamma$ -times weakly differentiable functions, and for  $\gamma \notin \mathbb{N}$  define  $H^\gamma$  by interpolation (see (Taylor 2011) for more details). Then by using the Sobolev embedding and the fact that  $S^\gamma$  can be continuously embedded in  $H^\gamma$  for any  $\gamma \geq 0$ , the map  $e$  can be shown to be Lipschitz as  $\|e(v) - e(\mathcal{L}u_f)\|_{L^2} \lesssim \|v - \mathcal{L}u_f\|_{L^2}$  on a subset of  $L^2$  where  $\|u_f - (K\mathcal{L}v + \tilde{g})\|_\infty$  is small enough. Therefore the contraction rate of  $\|f - f_0\|_{L^2}$  under the posterior is determined by the contraction rate of  $\|v - \mathcal{L}u_{f_0}\|_{L^2}$ . The full proof can be found in Section 3.7.2.

**Theorem 3.4.1.** *Let  $f_0 \in S^\beta$  with  $\beta \geq 9(d+2)$  be arbitrary. Take  $m \asymp N^\zeta$ , for some  $0 < \zeta < \frac{1}{3}$  and  $\alpha = \beta$ . Then for any sequence  $M_N \rightarrow \infty$ ,*

$$\mathbb{E}_0 \left( \Pi_{n,m}^\dagger (\|f - f_0\|_{L^2}^2 > M_N N^{-\frac{2\beta}{d+2\beta+4}} \mid \mathbb{D}_N) \right) \rightarrow 0.$$

## 3.5 Simulations

The theoretical results are illustrated by a simulation study various partial differential equations. We conduct a numerical study of the time-independent Schrödinger equation from Section 3.4.1, and a simulation study of the exponentiated Volterra operator as discussed in Section 2.7. First, we will discuss the general setup of the numerical study, which is equal for both inverse problems.

The data is distributed into  $m$  machines, where the  $k$ th machine gets the observations  $\mathbb{D}_n^{(k)} = \{(X_i^{(k)}, \tilde{Y}_i^{(k)}) : i = 1, \dots, n\}$ . For each  $k = 1, \dots, m$  and  $i = 1, \dots, I$ , let  $\hat{v}_i^{(k)} \in L^2$  denote the linear interpolation of an  $n$ -dimensional vector drawn from the local posterior distribution  $\Pi_{n,m}^{(k)}$  for  $u''_{f_0}$ . These draws are generated by using the explicit expression for the conditional Gaussian posterior distribution, given by  $N(\hat{f}_{n,m}^{(k)}, \hat{C}^{(k)})$ . The draws  $\hat{v}_i$  from the aggregated posterior distribution  $\Pi_{n,m}^\dagger$  for  $u''_f$  are then defined by taking the average over

draws  $\hat{v}_i = \frac{1}{m} \sum_{k=1}^m \hat{v}_i^{(k)}$  for  $i = 1, \dots, I$ . Then the draws from the posterior distribution for  $f$  have the form  $\hat{w}_i := e(\hat{v}_i)$ , for  $i = 1, \dots, I$ . The posterior credible bands are defined as pointwise 95% credible sets for each  $x \in [0, 1]$ , computed from the  $I$  draws of the aggregated posterior  $\tilde{\Pi}_{n,m}$  for  $f$ .

All the statistical methods are implemented in MATLAB and are performed on a computer with an AMD Ryzen 5 3600 processor. The computation of the local posteriors in the distributed method is performed sequentially.

### 3.5.1 Schrödinger equation

First, we consider the non-linear inverse problem in the time-independent Schrödinger equation, discussed in Section 3.4.1. Let  $u_f : [0, 1] \rightarrow \mathbb{R}$  be the solution to  $\frac{1}{2}u_f'' = fu_f$ , with  $u_f(0) = 3$  and  $u_f(1) = 4$ . For an unknown  $f_0 \in L^2((0, 1))$ , the observations consist of

$$Y_i = u_{f_0}(X_i) + \sigma\epsilon_i, \quad i = 1, \dots, N,$$

where  $\sigma > 0$  fixed and known,  $\epsilon_i \stackrel{\text{i.i.d.}}{\sim} N(0, 1)$  and  $X_i \stackrel{\text{i.i.d.}}{\sim} U(\mathcal{O})$ . The goal is to recover  $f_0$  from  $(X_1, Y_1), \dots, (X_N, Y_N)$ .

Take  $\mathcal{O} = (0, 1)$ ,  $\Gamma = \{0, 1\}$ ,  $\mathcal{L} = -\frac{d^2}{dx^2}$  and define the linear operator  $K : L^2((0, 1)) \rightarrow L^2((0, 1))$  by Equation (3.9). The operator  $K$  then satisfies

$$(Kv)(x) = - \int_0^x \int_0^t v(s) ds dt + x \int_0^1 \int_0^t v(s) ds dt,$$

for all  $v \in L^2((0, 1))$  and  $x \in [0, 1]$ . The map  $e$  is the inversion formula given in Equation (3.12). This method transforms the problem into a linear inverse problem for which we can use the results from Section 3.3. The eigenfunctions  $(\phi_j)_{j \in \mathbb{N}}$  of  $K$  are given by  $\phi_j(x) = \sqrt{2} \sin(j\pi x)$ , with corresponding eigenvalues  $\kappa_j = \frac{1}{\pi^2 j^2}$ . Then by using the framework from Section 3.4, it is possible to write the observational model as

$$\tilde{Y}_i := K(u_{f_0}'')(X_i) - \tilde{g}(X_i) + \sigma\epsilon_i, \quad (3.13)$$

with  $\tilde{g}(x) = 3 + x$ ,  $x \in [0, 1]$ .

We take  $f_0 = 10x(1-x) + \frac{1}{4} \sin(3\pi x)$ ,  $\sigma = 0.07$ ,  $\alpha = 0.3$  and  $I = 500$ . The total amount of observations is taken equal to  $N = 5000$ . We consider both the undistributed case with  $m = 1$  and the distributed case with  $m = 10$ . These priors are compared to draws from an undistributed wavelet prior. To generate the draws for the wavelet prior, a Metropolis-Hastings algorithm was used with

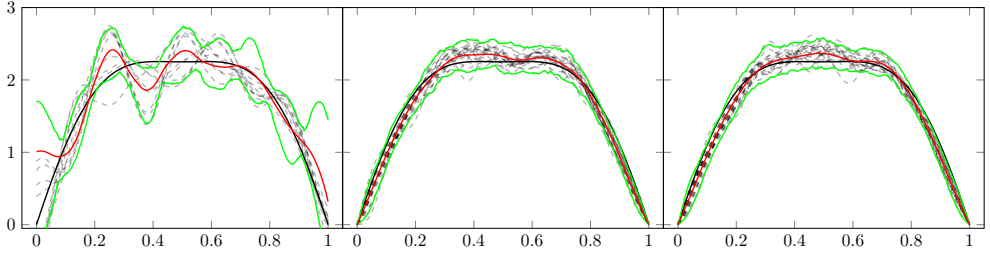


Figure 3.5.1: Inverse problem for the Schrödinger equation. The true function  $f_0$  (black), the posterior mean (red), 95% credible bands (green), and 20 draws from the posterior (dashed grey) for  $n = 2500$ . Left: wavelet prior with  $m = 1$ . Middle: undistributed method with  $m = 1$ . Right: distributed method with  $m = 10$ .

100.000 iterations. In Figure 3.5.1 the posterior draws (dashed grey), pointwise credible bands (green), and posterior mean (red) are plotted along with the true parameter  $f_0$  (black).

It can be seen that the recovery and uncertainty quantification is nearly identical for all the configurations, and that the posterior mean provides a reasonable estimator for  $f_0$ . On the boundary, the recovery of  $f_0$  by the wavelet prior is much worse than the prior based on the singular value decomposition of  $f_0$ , as the latter prior automatically ensures that the homogeneous boundary conditions on  $f_0$  are satisfied. The computation time for the wavelet prior was approximately 2.5 days, the computation time for  $m = 1$  with the prior as described in Section 3.2.1 was equal to 3 minutes, and the computation time for the same prior with  $m = 10$  was equal to 17 seconds. This shows that the distributed method scales roughly with a factor of  $m$ . Another speed-up of factor  $m$  can be achieved by executing the computations in parallel on  $m$  processors instead of sequentially on a single computer.

### 3.5.2 Exponentiated Volterra operator

Although the Volterra operator is not a self-adjoint operator, we now apply the same techniques to the exponentiated Volterra operator (c.f. Section 2.7). For an  $f \in L^2((0, 1))$ , define  $u_f$  as the solution to

$$\begin{cases} u'_f = f u_f & \text{on } (0, 1), \\ u_f = g & \text{at } 0, \end{cases}$$



for a known  $g > 0$ . The observations are equal to

$$Y = u_{f_0}(X_i) + \sigma \epsilon_i, \quad i = 1, \dots, N,$$

for an unknown  $f_0 : [0, 1] \rightarrow \mathbb{R}$ ,  $\epsilon_i \stackrel{\text{i.i.d.}}{\sim} N(0, 1)$  and a known  $\sigma > 0$ . With  $\mathcal{L} = \frac{d}{dx}$ , the operator  $K$  is the Volterra operator, which is defined by

$$Kv(x) = \int_0^x v(s) ds$$

for any function  $v \in L^2([0, 1])$ . The eigenfunctions of  $K^\top K$  are the functions from the cosine basis on  $[0, 1]$ , which is given by  $\phi_j(x) = \sqrt{2} \cos((j - 1/2)\pi x)$  for  $j \in \mathbb{N}$ . The corresponding eigenvalues equal  $\kappa_j^2 = \frac{1}{\pi^2(j-1/2)^2}$ . The inverse map  $e$  is the same as in Equation (2.29). Using this, we can define the equivalent linear problem as

$$\hat{Y} = Y - g = K(u'_{f_0}) + \sigma \epsilon_i, \quad i = 1, \dots, N.$$

We take  $f_0(x) = \cos(\frac{1}{2}\pi x) + \frac{1}{4}\cos(\frac{9}{2}\pi x)$ ,  $g = 3$ ,  $\sigma = 0.1$ ,  $\alpha = 0.5$  and  $I = 500$ . The total number of observations available is equal to  $N = 5000$ . We compare the undistributed setting with  $m = 1$ , to the distributed setting with  $m = 10$ . The results are plotted in Figure 3.5.2, where the true function (black curve), is compared to the draws from the posterior (dashed grey curves), the posterior mean (red curve), and the credible bands (green curves). It can be seen that both methods perform comparably, and in both cases, the credible bands capture the true function  $f_0$ . The recovery of the function is very accurate at the boundary point  $x = 1$ , but the uncertainty grows larger near the boundary point  $x = 0$ . The undistributed case with  $m = 1$  took 466 seconds, compared to the 46 seconds for the distributed case with  $m = 10$ .

## 3.6 Analysis score function

Let us first consider the undistributed problem with  $m = 1$ . Denote the RKHS of the prior by  $\mathcal{H} = \left\{ \sum_{j=1}^{\infty} h_j \phi_j : \sum_{j=1}^{\infty} \kappa_j^{-\frac{1}{p}(d+2\alpha)} h_j^2 < \infty \right\}$ , which becomes a Hilbert space with the inner product  $\langle g, h \rangle_{\mathcal{H}} = \sum_{j=1}^{\infty} \kappa_j^{-\frac{1}{p}(d+2\alpha)} g_j h_j$ . The posterior mean  $\hat{f}_N$  satisfies (see Theorem 11.61 in (Ghosal and van der Vaart 2017))

$$\hat{f}_N = \operatorname{argmin}_{f \in \mathcal{H}} \sum_{i=1}^N (Y_i - Kf(X_i))^2 + \sigma^2 \|f\|_{\mathcal{H}}^2. \quad (3.14)$$

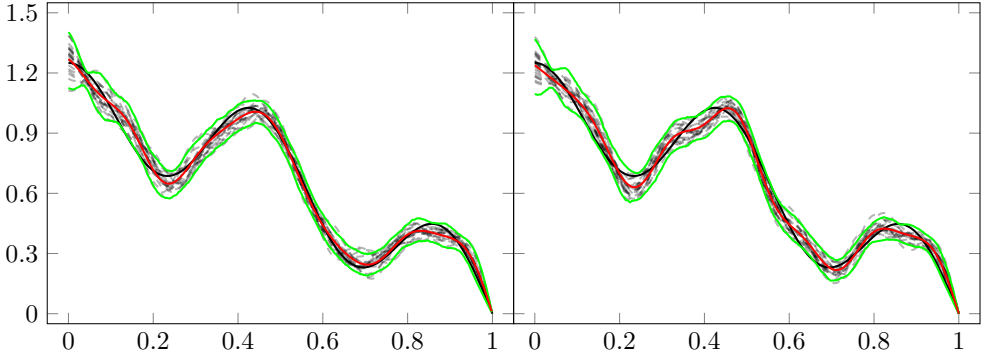


Figure 3.5.2: Inverse problem for the exponentiated Volterra operator. The true function  $f_0$  (black), the posterior mean (red), 95% credible bands (green), and 20 draws from the posterior (dashed grey) for  $n = 5000$ . Left: wavelet prior with  $m = 1$ . Middle: undistributed method with  $m = 1$ . Right: distributed method with  $m = 10$ .

Denote by  $C_x \in \mathcal{H}$  the function  $C_x : \mathcal{O} \rightarrow \mathbb{R}$  defined by  $y \mapsto C(x, y)$ , i.e. for all  $y \in \mathcal{O}$ , it satisfies  $C_x(y) = \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha)} \phi_j(x) \phi_j(y)$ . For  $f \in \mathcal{H}$  we can write  $f(x) = \langle f, C_x \rangle_{\mathcal{H}}$ , so for all  $f \in \mathcal{H}$ , the penalized square error  $\ell_N : \mathcal{H} \rightarrow \mathbb{R}$  is, up to constant terms, equal to

$$-\ell_N(f) := \sum_{i=1}^N (Y_i - \langle Kf, C_{X_i} \rangle_{\mathcal{H}})^2 + \sigma^2 \langle f, f \rangle_{\mathcal{H}}.$$

The Fréchet derivative of the first term of the log-density with respect to  $\|\cdot\|_{\mathcal{H}}$  is given by

$$Ah := -2 \sum_{i=1}^N (Y_i - \langle Kf, C_{X_i} \rangle_{\mathcal{H}}) \langle Kh, C_{X_i} \rangle_{\mathcal{H}},$$

since for any  $h \in \mathcal{H}$  we have

$$\begin{aligned} & (Y_i - \langle K(f+h), C_{X_i} \rangle_{\mathcal{H}})^2 - (Y_i - \langle Kf, C_{X_i} \rangle_{\mathcal{H}})^2 - Ah \\ &= \langle Kh, C_{X_i} \rangle_{\mathcal{H}}^2 \leq \|h\|_{\mathcal{H}}^2 \|K_{\mathcal{H}}^* C_{X_i}\|_{\mathcal{H}}^2, \end{aligned}$$

where in the last line the Cauchy-Schwarz inequality was applied. It now suffices to note that the last term is  $O(\|h\|_{\mathcal{H}}^2)$ . Here  $K_{\mathcal{H}}^*$  is the adjoint of the function  $K_{\mathcal{H}} : \mathcal{H} \rightarrow \mathcal{H}$ , which is the restriction of  $K$  to  $\mathcal{H} \subset L^2$ . Thus the Fréchet

derivative of the penalized square error is given by

$$2 \sum_{i=1}^N (\langle Kf, C_{X_i} \rangle_{\mathcal{H}} - Y_i) \langle h, K_{\mathcal{H}}^* C_{X_i} \rangle_{\mathcal{H}} + 2\sigma^2 \langle h, f \rangle_{\mathcal{H}}. \quad (3.15)$$

Since  $K \sum_{j=1}^{\infty} f_j \phi_j = \sum_{j=1}^{\infty} \kappa_j f_j \phi_j$ ,  $K_{\mathcal{H}}^*$  is equal to  $K$ . Multiplying by  $-\frac{1}{2N}$  and keeping the functions of the second argument of the inner products, we define the adjusted score function  $\hat{S}_N : \mathcal{H} \rightarrow \mathcal{H}$  by

$$\hat{S}_N(f) := \frac{1}{N} \left( \sum_{i=1}^N (Y_i - (Kf)(X_i)) K C_{X_i} - \sigma^2 f \right). \quad (3.16)$$

Since  $\hat{f}_N$  minimizes  $-\ell_N$ , its Fréchet derivatives must equal zero for all directions  $h \in \mathcal{H}$ . Therefore the score must satisfy  $\hat{S}_N(\hat{f}_N) = 0$ . The expectation of  $\hat{S}_N$  is

$$\begin{aligned} S_N(f) &:= \mathbb{E}_0 \hat{S}_N(f) \\ &= \mathbb{E}_X [(Kf_0(X) - Kf(X)) K C_X] - \frac{\sigma^2}{N} f \\ &= \int_{\mathcal{O}} (Kf_0(x) - Kf(x)) K C_x dx - \frac{\sigma^2}{N} f, \end{aligned}$$

since the  $X_i$ 's have a uniform distribution on the domain.

Consider the operator  $\tilde{K} = KC : L^2 \rightarrow L^2$  given by

$$\tilde{K}(f)(t) = (KCf)(t) = \int_{\mathcal{O}} f(s) (KC_s)(t) ds.$$

Now we have

$$\begin{aligned} \tilde{K}(f) &= \int_{\mathcal{O}} f(s) (KC_s) ds \\ &= \int_{\mathcal{O}} \sum_{i=1}^{\infty} f_i \phi_i(s) K \left( \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha)} \phi_j(s) \phi_j \right) ds \\ &= \int_{\mathcal{O}} \sum_{i=1}^{\infty} f_i \phi_i(s) \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha)+1} \phi_j(s) \phi_j ds \\ &= \sum_{j=1}^{\infty} f_j \kappa_j^{\frac{1}{p}(d+2\alpha+p)} \phi_j \\ &= K^{\frac{1}{p}(d+2\alpha+p)}(f). \end{aligned}$$

We see that we can also write  $S_N(f) = \tilde{K}(Kf_0 - Kf) - \frac{\sigma^2}{N}f$  and thus

$$\begin{aligned} S_N(f) &= K^{\frac{1}{p}(d+2\alpha+2p)}(f_0 - f) - \frac{\sigma^2}{N}f \\ &= \sum_{j=1}^{\infty} \left( \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} f_{0,j} - \frac{\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N}{N} f_j \right) \phi_j. \end{aligned}$$

Thus the solution to  $S_N(f) = 0$  is given by the function  $f = \sum_{j=1}^{\infty} \nu_j f_{0,j} \phi_j$ , with

$$\nu_j = \frac{\kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N}{\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N}.$$

Define the index-wise multiplication with these  $\nu_j$ 's as the operator  $\tilde{V}_N : L^2 \rightarrow L^2$  by  $\tilde{V}_N(f) = \sum_{j=1}^{\infty} \nu_j f_j \phi_j$ , and  $\tilde{P}_N := \text{id} - \tilde{V}_N$ . We can then write

$$S_N(f) = \tilde{K}(Kf_0) - \tilde{K}\tilde{V}_N^{-1}(Kf) = \tilde{K}(Kf_0 - \tilde{V}_N^{-1}(Kf)).$$

Define  $\Delta \hat{f}_N := \hat{f}_N - \tilde{V}_N(f_0)$  and the operator  $\tilde{R}_N := K^{-1}\tilde{V}_N\tilde{K}^{-1}$ . Then we have

$$\Delta \hat{f}_N = -\tilde{R}_N(S_N(\hat{f}_N)), \quad (3.17)$$

since

$$\begin{aligned} \tilde{R}_N(S_N(\hat{f}_N)) &= \sum_{j=1}^{\infty} \phi_j \kappa_j^{-1} \nu_j \kappa_j^{-\frac{1}{p}(d+2\alpha+p)} \\ &\quad \times \left( \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} f_{0,j} - \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} \frac{\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N}{\kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N} \kappa_j \hat{f}_{N,j} \right) \\ &= \sum_{j=1}^{\infty} \nu_j f_{0,j} \phi_j - \hat{f}_{N,j} \phi_j = \tilde{V}_N(f_0) - \hat{f}_N = -\Delta \hat{f}_N. \end{aligned}$$

Next we consider the local problem in a distributed setting with  $m$  machines as proposed in Section 3.2.1. At the  $k$ th machine,  $k \in \{1, \dots, m\}$ , the local posterior mean  $\hat{f}_n^{(k)}$  is equal to

$$\hat{f}_n^{(k)} = \underset{f \in \mathcal{H}}{\operatorname{argmin}} \sum_{i=1}^n (Y_i^{(k)} - Kf(X_i^{(k)}))^2 + \frac{\sigma^2}{m} \|f\|_{\mathcal{H}}^2. \quad (3.18)$$

The sample score function of the  $k$ th local machine has the form

$$\begin{aligned}\hat{S}_{n,m}^{(k)}(f) &= \frac{1}{n} \left( \sum_{i=1}^n (Y_i^{(k)} - Kf(X_i^{(k)})) KC_{X_i^{(k)}} - \frac{\sigma^2}{m} f \right) \\ S_{n,m}^{(k)}(f) &= \int_{\mathcal{O}} (Kf_0(x) - Kf(x)) K_{\mathcal{H}}^* C_x dx - \frac{\sigma^2}{N} f.\end{aligned}\quad (3.19)$$

These satisfy  $\hat{S}_{n,m}^{(k)}(\hat{f}_n^{(k)}) = 0$  and  $S_{n,m}^{(k)} = S_N$ , so  $S_{n,m}^{(k)}(\tilde{V}_N(f_0)) = 0$ . Because the aggregated posterior is an average of the local posteriors, the aggregated posterior mean equals  $\hat{f}_{n,m}^\dagger = \frac{1}{m} \sum_{k=1}^m \hat{f}_n^{(k)}$ . Define  $\Delta \hat{f}_n^{(k)} = \hat{f}_n^{(k)} - \tilde{V}_N(f_0)$  and  $\Delta \hat{f}_{n,m}^\dagger = \hat{f}_{n,m}^\dagger - \tilde{V}_N(f_0) = \frac{1}{m} \sum_{k=1}^m \Delta \hat{f}_n^{(k)}$ .

## 3.7 Proofs

The following proofs are based on techniques introduced in (Yang, Bhattacharya and Pati 2017) and are an extension of the direct non-parametric regression problem studied in (Hadji, Hesselink and Szabó 2022) to the inverse problem setting, with general smoothness norms.

### 3.7.1 Proof of Theorem 3.3.1

*Proof of Theorem 3.3.1.* Since

$$\mathbb{E}_0 \|\hat{f}_{n,m}^\dagger - f_0\|_{S^\gamma}^2 \leq 2\|f_0 - \tilde{V}_N(f_0)\|_{S^\gamma}^2 + 2\mathbb{E}_0 \|\hat{f}_{n,m}^\dagger - \tilde{V}_N(f_0)\|_{S^\gamma}^2,$$

and the second term of the right-hand side is bounded by Lemma 3.7.1, we get that there exists a  $C > 0$  such that for any  $\gamma \geq 0$ , and  $f_0 \in S^\beta(B)$ , up to a negligible  $o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}})$  term we have

$$\mathbb{E}_0 \|\hat{f}_{n,m}^\dagger - f_0\|_{S^\gamma}^2 \leq C \left( \|\tilde{P}_N(f_0)\|_{S^\gamma}^2 + \left( \frac{1}{N} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \right) (B^2 + \sigma^2) \right). \quad (3.20)$$

Next, by Markov's inequality and the bound  $(a+b)^2 \leq 2(a^2+b^2)$  we can bound the expectation of the posterior mass by

$$\begin{aligned}\mathbb{E}_0 \Pi_{n,m}^\dagger(f : \|f - f_0\|_{S^\gamma}^2 \geq M_N \epsilon_N \mid \mathbb{D}_N) \\ \leq 2 \frac{\mathbb{E}_0 \Pi_{n,m}^\dagger[\|f - \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 \mid \mathbb{D}_N] + \mathbb{E}_0 \|\hat{f}_{n,m}^\dagger - f_0\|_{S^\gamma}^2}{M_N \epsilon_N}.\end{aligned}$$

Since the condition on  $m$  of Lemma 3.9.1 are satisfied by the assumption on  $\beta$ , the posterior  $\|\cdot\|_{S^\gamma}$ -variance on a local machine is bounded by

$$\mathbb{E}_0(\Pi_{n,m}^{(k)}(\|f - \hat{f}_{n,m}^{(k)}\|_{S^\gamma}^2 \mid \mathbb{D}_n^{(k)})) \lesssim \frac{\sigma^2}{n} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}}.$$

Thus the variance of the aggregated posterior with respect to the  $S^\gamma$ -norm is bounded by

$$\begin{aligned} \mathbb{E}_0(\Pi_{n,m}^\dagger(\|f - \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 \mid \mathbb{D}_N)) &= \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0(\Pi_{n,m}^{(k)}(\|f - \hat{f}_{n,m}^{(k)}\|_{S^\gamma}^2 \mid \mathbb{D}_n^{(k)})) \\ &\lesssim \frac{\sigma^2}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}}. \end{aligned}$$

The  $S^\gamma$ -norm of  $\tilde{P}_N(f_0)$  can be bounded from above by

$$\|\tilde{P}_N(f_0)\|_{S^\gamma}^2 = \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} \frac{\sigma^4}{(\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\beta+2p)} N)^2} f_{0,j}^2 \leq C \|f_0\|_{S^\gamma}^2$$

for some fixed constant  $C > 0$  large enough. Combining the previous displays now shows

$$\begin{aligned} &\mathbb{E}_0(\Pi_{n,m}^\dagger(f : \|f - f_0\|_{S^\gamma} \geq M_N \epsilon_N \mid \mathbb{D}_N)) \\ &\leq 2 \frac{\mathbb{E}_0(\Pi_{n,m}^\dagger(\|f - \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 \mid \mathbb{D}_N)) + \mathbb{E}_0\|\hat{f}_{n,m}^\dagger - f_0\|_{S^\gamma}^2}{M_N \epsilon_N} \\ &\lesssim \frac{\frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}} + \|\tilde{P}_N(f_0)\|_{S^\gamma}^2 + \left(\frac{1}{N} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2-\frac{2\gamma}{p}}\right)(B + \sigma^2)}{M_N \epsilon_N} \\ &\quad + (M_N \epsilon_N)^{-1} o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}) \\ &\asymp \frac{\frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}} + \|\tilde{P}_N(f_0)\|_{S^\gamma}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}})}{M_N \epsilon_N}. \end{aligned}$$

By taking  $\epsilon_N = \frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}} + \|\tilde{P}_N(f_0)\|_{S^\gamma}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}})$  the right-hand side of the last display goes to zero for  $N \rightarrow \infty$ .

We will now derive the asymptotic convergence rates of the three terms in  $\epsilon_N$ . We write

$$\begin{aligned} \frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}} &= \sum_{j=1}^{\infty} \frac{\kappa_j^{\frac{1}{p}(d+2\beta-2\gamma)}}{\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\beta+2p)} N} \\ &\asymp \sum_{j=1}^{\infty} \frac{j^{-\frac{1}{d}(d+2\beta-2\gamma)}}{\sigma^2 + j^{-\frac{1}{d}(d+2\beta+2p)} N}. \end{aligned}$$

We apply Lemma 3.10.7 with  $t = \frac{1}{d}(d+2\beta-2\gamma)$ ,  $u = \frac{1}{d}(d+2\beta+2p)$ ,  $q = -\frac{1}{2}$  and  $v = 1$ . The sum is asymptotically of the order

$$\frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}} \asymp N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}. \quad (3.21)$$

Note that

$$\|\tilde{P}_N(f_0)\|_{S^\gamma}^2 = \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} (1 - \nu_j)^2 f_{0,j}^2 \asymp \sigma^4 \sum_{j=1}^{\infty} \frac{j^{\frac{2\gamma}{d}}}{(\sigma^2 + N j^{-\frac{1}{d}(d+2\beta+2p)})^2} f_{0,j}^2.$$

The conditions of Lemma 3.10.6 are satisfied with  $u = \frac{1}{d}(d+2\beta+2p)$ ,  $v = 2$ ,  $q = \frac{\beta}{d}$  and  $t = -\frac{2\gamma}{d}$ , since  $t = -\frac{2\gamma}{d} > \frac{-2\beta}{d} = -2q$ . It follows that

$$\|\tilde{P}_N(f_0)\|_{S^\gamma}^2 \asymp \sum_{j=1}^{\infty} \frac{j^{\frac{2\gamma}{d}}}{(\sigma^2 + N j^{-\frac{1}{d}(d+2\beta+2p)})^2} f_{0,j}^2 \lesssim N^{-\left(\frac{2\beta-2\gamma}{d+2\beta+2p} \wedge 2\right)},$$

where the constant depends on  $\|f_0\|_{S^\gamma}^2$ .

We find that the posterior contraction rate is asymptotically upper bounded by  $N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}$ .  $\square$

### 3.7.2 Proof of Theorem 3.4.1

*Proof of Theorem 3.4.1.* For  $v \in L^2$ , define  $f_v := e(v)$ . We can write  $f_v - f_0 = e(v) - e(\mathcal{L}u_{f_0})$ , from which we see by elementary algebra, that its  $L^2$ -norm is bounded by (c.f. Chapter 2)

$$\begin{aligned} \|f_v - f_0\|_{L^2} &\lesssim \left( \inf_x u_{f_0}(x) - \|u_{f_0} - Kv - \tilde{g}\|_\infty \right)^{-1} \\ &\quad \times [(\|v\|_{L^2} + \|v - \mathcal{L}u_{f_0}\|_{L^2}) \|Kv - K\mathcal{L}u_{f_0}\|_{L^2} + \|v - \mathcal{L}u_{f_0}\|_{L^2}]. \end{aligned}$$

We know that  $\|v\|_{L^2} \leq \|\mathcal{L}u_{f_0}\|_{L^2} + \|v - \mathcal{L}u_{f_0}\|_{L^2}$ . Furthermore, due to the continuity of  $K$  we have the bound  $\|Kv - K\mathcal{L}u_{f_0}\|_{L^2} \leq \|K\|\|v - \mathcal{L}u_{f_0}\|_{L^2}$ . Lastly, the Sobolev embedding shows that

$$\|u_{f_0} - Kv - \tilde{g}\|_\infty = \|Kv - K\mathcal{L}u_{f_0}\|_\infty \lesssim \|Kv - K\mathcal{L}u_{f_0}\|_{H^{\frac{d}{2}+\epsilon}}$$

for any  $\epsilon > 0$ . The norm on the right can be bounded by  $\|Kv - K\mathcal{L}u_{f_0}\|_{S^{\frac{d}{2}+\epsilon}}$  by Equation (A.18) in (Taylor 2011). We know from Theorem 3.3.1 with  $\gamma = \frac{d}{2} + \epsilon < \beta$  for sufficiently small  $\epsilon > 0$ , that this goes to zero under the posterior. Hence under the posterior, the multiplicative factor on the right-hand side of the first display of the proof is bounded above by a fixed constant for large enough  $N$ . Therefore the expectation of the posterior of  $f$

$$\mathbb{E}_0 \Pi_{n,m}^\dagger (\|f - f_0\|_{L^2}^2 > M_N N^{-\frac{2\beta}{d+2\beta+4}} \mid \mathbb{D}_N)$$

is bounded above by

$$\mathbb{E}_0 \Pi_{n,m}^\dagger (\|v - v_0\|_{L^2}^2 > CM_N N^{-\frac{2\beta}{d+2\beta+4}} \mid \mathbb{D}_N)$$

for some fixed constant  $C > 0$ , with  $v_0 = \mathcal{L}u_{f_0}$ . The last expectation goes to zero by Theorem 3.3.2.  $\square$

### 3.7.3 Proof of the lemmas

**Lemma 3.7.1.** *Assume that  $\alpha = \beta$ ,  $\beta \geq 9(d+p)$  and  $0 < \zeta < \frac{1}{3}$ . Then there exists a constant  $C > 0$  such that for all  $f_0 \in S^\beta$ , asymptotically for  $m \rightarrow \infty$ ,*

$$\mathbb{E}_0 \|\hat{f}_{n,m}^\dagger - \tilde{V}_N(f_0)\|_{S^\gamma}^2 \leq \frac{C}{N} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} (\|\tilde{P}_N(f_0)\|_{S^\beta}^2 + \sigma^2) + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}).$$

*Proof.* In view of the inequality  $(a+b)^2 \leq 2a^2 + 2b^2$  we bound  $\|\Delta \hat{f}_n^{(k)}\|_{S^\gamma}^2$  from above by

$$2\|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 + 2\|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2.$$

From Lemma 3.7.2 we find that the expectation of the first term is bounded by

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 \leq \frac{C'}{m} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}) \quad (3.22)$$



for some  $C' > 0$  independent of  $f_0$ . Combining the two previous displays and taking  $m$  large enough, shows that we have

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{S^\gamma}^2 \lesssim \frac{1}{m} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 + \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}).$$

When  $m \rightarrow \infty$ , this results in the bound

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{S^\gamma}^2 \leq C \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}), \quad (3.23)$$

with some  $C > 0$  independent of  $f_0$ . By first applying the inequality (3.22), and then upper bounding  $\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{S^\gamma}^2$  with the use of Equation (3.23), we see that

$$\begin{aligned} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 & \quad (3.24) \\ & \leq \frac{C'}{m} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}) \\ & \lesssim \frac{1}{m} (\mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}})) + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}) \\ & \lesssim \frac{1}{m} \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}). \end{aligned}$$

We have the inequality

$$\begin{aligned} \|\Delta \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 & \lesssim \left\| \frac{1}{m} \sum_{k=1}^m \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2 \\ & \quad + \left\| \Delta \hat{f}_{n,m}^\dagger - \frac{1}{m} \sum_{k=1}^m \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2. \end{aligned} \quad (3.25)$$

Note that  $\mathbb{E}_0 \hat{S}_{n,m}^{(k)} \tilde{V}_N(f_0) = S_{n,m}(\tilde{V}_N(f_0)) = 0$ , so  $\mathbb{E}_0 \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) = 0$ . Hence, by the independence of the local experiments,

$$\mathbb{E}_0 \left\| \frac{1}{m} \sum_{k=1}^m \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2 = \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2.$$

We apply the previous display to the first term of Equation (3.25), and use the definition of  $\Delta \hat{f}_{n,m}^\dagger$  for the second term of Equation (3.25). This results in the bound

$$\begin{aligned} \mathbb{E}_0 \|\Delta \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 & \lesssim \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 \\ & \quad + \mathbb{E}_0 \left\| \frac{1}{m} \sum_{k=1}^m (\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))) \right\|_{S^\gamma}^2. \end{aligned}$$

For any functions  $g_1, \dots, g_m \in S^\gamma$ , we can bound the squared  $S^\gamma$ -norm of their sum by  $\|\sum_{k=1}^m g_k\|_{S^\gamma}^2 \leq m \sum_{k=1}^m \|g_k\|_{S^\gamma}^2$ . We apply this to the second term of the last display and find that

$$\begin{aligned} \mathbb{E}_0 \|\Delta \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 &\lesssim \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 \\ &\quad + \frac{1}{m} \sum_{k=1}^m \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2. \end{aligned}$$

Now using the upper bound from Equation (3.24) for the second term of the last display, we see that

$$\begin{aligned} \|\Delta \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 &\lesssim \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 \\ &\quad + \frac{1}{m} \sum_{k=1}^m \frac{1}{m} \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}) \\ &\lesssim \frac{1}{m^2} \sum_{k=1}^m \mathbb{E}_0 \|\tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}). \end{aligned}$$

An application of Lemma 3.7.4 to the expectations inside the sum shows that the previous display is bounded by a multiple of

$$\|\Delta \hat{f}_{n,m}^\dagger\|_{S^\gamma}^2 \lesssim \frac{1}{m} \left( \frac{1}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \right) (\|\tilde{P}_N(f_0)\|_{S^\beta}^2 + \sigma^2) + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}),$$

where the constant is independent of  $f_0$ , finishing the proof.  $\square$

**Lemma 3.7.2.** *Assume that  $\alpha = \beta$ ,  $\beta > 9(d+p)$  and  $0 < \zeta < \frac{1}{3}$ . Then there exists a constant  $C > 0$  such that for all  $f_0 \in S^\beta(B)$ ,*

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 \leq \frac{C}{m} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 + o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}).$$

*Proof.* By Equation (3.17), we have  $\Delta \hat{f}_n^{(k)} = -\tilde{R}_N(S_{n,m}^{(k)}(\hat{f}_n^{(k)}))$ . Applying the inverse operator  $\tilde{R}_N^{-1}$  to both sides shows that  $S_{n,m}^{(k)}(\hat{f}_n^{(k)}) = -\tilde{R}_N^{-1}(\Delta \hat{f}_n^{(k)})$ . Using the fact that  $\hat{f}_n^{(k)}$  and  $\tilde{V}_N(f_0)$  are the roots of  $\hat{S}_{n,m}^{(k)}$  and  $S_{n,m}^{(k)}$ , respectively, shows that

$$\begin{aligned} &(\hat{S}_{n,m}^{(k)}(\hat{f}_n^{(k)}) - S_{n,m}^{(k)}(\hat{f}_n^{(k)})) - (\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)) - S_{n,m}^{(k)}(\tilde{V}_N(f_0))) \\ &= \tilde{R}_N^{-1}(\Delta \hat{f}_n^{(k)}) - \hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)). \end{aligned}$$

Moreover, by Equation (3.19) we have

$$\begin{aligned} \hat{S}_{n,m}^{(k)}(f) - S_{n,m}^{(k)}(f) \\ = \frac{1}{n} \sum_{i=1}^n (Y_i^{(k)} - Kf(X_i^{(k)})) KC_{X_i^{(k)}} - \int_{\mathcal{O}} (Kf_0(x) - Kf(x)) K_{\mathcal{H}}^* C_x dx. \end{aligned}$$

Taking  $f = \hat{f}_n^{(k)}$  and  $f = \tilde{V}_N(f_0)$  in the preceding display, we arrive at

$$\begin{aligned} (\hat{S}_{n,m}^{(k)}(\hat{f}_n^{(k)}) - S_{n,m}^{(k)}(\hat{f}_n^{(k)})) - (\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)) - S_{n,m}^{(k)}(\tilde{V}_N(f_0))) \\ = -\frac{1}{n} \sum_{i=1}^n K(\Delta \hat{f}_n^{(k)})(X_i^{(k)}) KC_{X_i^{(k)}} + \int_{\mathcal{O}} K(\Delta \hat{f}_n^{(k)})(x) KC_x dx. \end{aligned}$$

The conditions of Lemma 3.7.3 with  $\hat{g}^{(k)} = \Delta \hat{f}_n^{(k)}$  are satisfied due to Lemma 3.8.3, so that for some fixed  $C > 0$  independent of  $f$ , up to a negligible term we have, for any  $\mathcal{I} \lesssim N^c$  for arbitrary  $c > 0$ ,

$$\begin{aligned} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)} - \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)))\|_{S^\gamma}^2 & \quad (3.26) \\ & = \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n K(\Delta \hat{f}_n^{(k)})(X_i^{(k)}) KC_{X_i^{(k)}} - \int_{\mathcal{O}} K(\Delta \hat{f}_n^{(k)})(x) KC_x dx \right) \right\|_{S^\gamma}^2 \\ & \leq C \left( \frac{\log N \sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 \right. \\ & \quad \left. + \mathbb{E}_0 \|\Delta \hat{f}_{n,\mathcal{I}^c}^{(k)}\|_{\mathcal{H}}^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell > \mathcal{I}}^{\infty} \kappa_\ell^{\frac{1}{p}(d+2\alpha+2p)} + N^{-c} \right) \\ & =: T_1 + T_2. \end{aligned}$$

Here  $\Delta \hat{f}_{n,\mathcal{I}^c}^{(k)}$  denotes the tail of the  $\Delta \hat{f}_n^{(k)}$ , so  $\Delta \hat{f}_{n,\mathcal{I}^c}^{(k)} = \sum_{j=\mathcal{I}+1}^{\infty} (\Delta \hat{f}_n^{(k)})_j \phi_j$ . Note that

$$\sup_{f_0 \in S^\beta(B)} \mathbb{E}_0 \|\Delta \hat{f}_{n,\mathcal{I}^c}^{(k)}\|_{\mathcal{H}}^2 \leq \sup_{f_0 \in S^\beta(B)} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{\mathcal{H}}^2 \lesssim N$$

by Lemma 3.8.2. We take  $\mathcal{I}$  to be the largest  $\mathcal{I} \in \mathbb{N}$  that satisfies  $\mathcal{I} \leq N^Q$  for some  $Q > 0$  and

$$\sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2 \leq \frac{n}{m \log N} \left( \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \right)^{-1}.$$

Then the term  $T_1$  on the right-hand side of Equation (3.26) is  $O\left(\frac{1}{m} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2\right)$ .

We now show that  $T_2$  in Equation (3.26) is  $o(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}})$ . In order to do this, we first derive an explicit bound for  $\mathcal{I}$ . We have

$$\begin{aligned} \sigma^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2-\frac{2\gamma}{p}} &= \sigma^2 \sum_{j=1}^{\infty} \frac{\kappa_j^{\frac{2}{p}(d+2\alpha+p-\gamma)} N^2}{(\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N)^2} \\ &\asymp N^2 \sum_{j=1}^{\infty} \frac{j^{-\frac{2}{d}(d+2\alpha+p-\gamma)}}{(\sigma^2 + N j^{-\frac{1}{d}(d+2\alpha+2p)})^2}. \end{aligned}$$

We can apply Lemma 3.10.7, with  $q = -\frac{1}{2}$ ,  $t = \frac{2}{d}(d+2\alpha+p-\gamma)$ ,  $v = 2$ ,  $u = \frac{1}{d}(d+2\alpha+2p)$  and  $\mathcal{S} \equiv 1$ , since  $q = -\frac{1}{2} > -\frac{1}{d}(d+2\alpha+p-\gamma) = -\frac{t}{2}$ . From this lemma, we see that

$$\sigma^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2-\frac{2\gamma}{p}} \asymp N^2 N^{-\frac{d+4\alpha+2p-2\gamma}{d+2\alpha+2p}} = N^{\frac{d+2p+2\gamma}{d+2\alpha+2p}}. \quad (3.27)$$

Note that by definition,  $n \asymp N^{1-\zeta}$  and  $m \asymp N^\zeta$ , so that  $\mathcal{I}$  needs to satisfy

$$\log(N) \sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2 \leq N^{-\frac{d+2p+2\gamma}{d+2\alpha+2p}+1-2\zeta}.$$

We now first consider  $2p \geq d$ . In the case with  $\mathcal{I} \asymp N^Q$  we can bound the sum as  $\log N \sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2 \lesssim \log N \sum_{\ell=1}^{N^Q} j^{-1} \asymp Q(\log N)^2$ . Thus we can take  $Q$  arbitrarily large and  $\zeta$  such that  $\zeta < \frac{\alpha-\gamma}{d+2\alpha+2p}$ , showing that  $T_2$  is negligible. The second case is  $2p < d$ , for which we have

$$\sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2 \asymp \sum_{\ell=1}^{\mathcal{I}} \ell^{-\frac{2p}{d}} \asymp \int_1^{N^Q} x^{-\frac{2p}{d}} dx \asymp N^{Q(1-\frac{2p}{d})}.$$

Thus we take  $Q = \frac{d-s}{d-2p} \left( \frac{2\alpha-2\gamma}{d+2\alpha+2p} - 2\zeta \right)$  for  $s > 0$  arbitrarily small. We can now bound the term as

$$\begin{aligned} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_\ell^{\frac{1}{p}(d+2\alpha+2p)} &= \sum_{j > cN^{\frac{d-s}{d-2p} \left( \frac{2\alpha-2\gamma}{d+2\alpha+2p} - 2\zeta \right)}} j^{-\frac{1}{d}(d+2\alpha+2p)} \\ &\lesssim \int_{N^{\frac{d-s}{d-2p} \left( \frac{2\alpha-2\gamma}{d+2\alpha+2p} - 2\zeta \right)}}^{\infty} x^{-\frac{1}{d}(d+2\alpha+2p)} dx \\ &\lesssim N^{-\frac{d-s}{d-2p} \left( \frac{2\alpha-2\gamma}{d+2\alpha+2p} - 2\zeta \right) \frac{2p+2\alpha}{d}}. \end{aligned}$$

By applying the previous bound, we can upper bound  $T_2$  by

$$\begin{aligned}
 N \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_{\ell}^{\frac{1}{p}(d+2\alpha+2p)} &\asymp N^{1+\frac{d+2p+2\gamma}{d+2\alpha+2p}-\frac{d}{d-2p}\left(\frac{2\alpha-2\gamma}{d+2\alpha+2p}-2\zeta\right)\frac{2p+2\alpha}{d}+\tilde{s}} \\
 &= N^{\frac{2(-2\alpha^2+\alpha(d-4p)+(d-2p)(d+\gamma+2p))}{(d-2p)(d+2\alpha+2p)}+4\zeta\frac{p+\alpha}{d-2p}+\tilde{s}}
 \end{aligned} \tag{3.28}$$

for some arbitrarily small  $\tilde{s} > 0$ . The exponent being smaller than  $-\frac{2\beta-2\gamma}{d+2\beta+2p}$ , with  $\alpha = \beta$ , is equivalent to the polynomial in  $\beta$  satisfying

$$(-4 + 8\zeta)\beta^2 + (4d - 12p + \zeta(16p + 4d))\beta + 2d^2 - 8p^2 + 4\zeta(dp + 2p^2) < 0.$$

Since  $\zeta < \frac{1}{2}$ , it suffices that

$$\beta > \left( \frac{8}{4-8\zeta} + \frac{\sqrt{10}}{\sqrt{4-8\zeta}} \right) d + \left( \frac{4}{4-8\zeta} + \frac{\sqrt{2}}{\sqrt{4-8\zeta}} \right) p.$$

Since  $\zeta < \frac{1}{3}$ , these quantities can be lower bounded by 9, resulting in the sufficient condition

$$\beta > 9(d + p),$$

showing that

$$T_2 = o\left(N^{-\frac{2\beta-2\gamma}{d+2\beta+2p}}\right). \tag{3.29}$$

This concludes the proof.  $\square$

**Lemma 3.7.3.** *Consider an  $\mathcal{I} \in \mathbb{N}$  with  $\mathcal{I} \leq N^Q$  for some  $Q > 0$ . Then there exists a constant  $C' > 0$  such that for all functions of the form  $\hat{g}^{(k)}(x) := g(x, \mathbb{D}_n^{(k)})$  for some measurable map  $g$  on  $\mathcal{O} \times \mathbb{R}^n \times \mathcal{O}^n$  that is measurable and satisfies, for any  $\sigma(X_1^{(k)}, \dots, X_n^{(k)})$ -measurable event  $A$ , the bound  $\mathbb{E}_0 \|\hat{g}^{(k)}\|_{L^2}^2 1_A \lesssim N^Q \mathbb{P}(A)^{\frac{1}{2}}$ , we have*

$$\begin{aligned}
 &\mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n K(\hat{g}^{(k)})(X_i^{(k)}) K C_{X_i^{(k)}} - \mathbb{E}_X(K\hat{g}^{(k)}(X) K C_X) \right) \right\|_{S^{\gamma}}^2 \\
 &\leq C' \left( \frac{\log N \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell}^2}{n} \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} \nu_j^2 \kappa_j^{-2} \mathbb{E}_0 \|\hat{g}^{(k)}\|_{L^2}^2 \right. \\
 &\quad \left. + \mathbb{E}_0 \|\hat{g}_{\mathcal{I}^c}^{(k)}\|_{\mathcal{H}}^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_{\ell}^{\frac{1}{p}(d+2\alpha+2p)} + N^{-C_0} \right),
 \end{aligned}$$

with  $C_0 > 0$  arbitrarily large. Here  $\hat{g}_{\mathcal{I}^c}^{(k)} = \sum_{j=\mathcal{I}+1}^{\infty} g_j^{(k)} \phi_j$ .

*Proof.* Let  $\mathcal{I} = \mathcal{I}(N) \in \mathbb{N}$  be such that  $\mathcal{I} \leq N^Q$  for some  $Q > 0$  and let  $k \in \{1, \dots, m\}$  be fixed. For some  $t \geq 0$  to be determined later, define

$$\mathcal{A}_{\mathcal{I},j} := \left\{ (X_1^{(k)}, \dots, X_n^{(k)}) \in \mathcal{O}^n : \left( \frac{1}{n} \sum_{i=1}^n \phi_j(X_i^{(k)}) \phi_\ell(X_i^{(k)}) - \delta_{j,\ell} \right)^2 \leq \frac{t \log N}{n}, \forall \ell \leq \mathcal{I} \right\}.$$

In view of assumption (3.4), we use the union bound followed by Hoeffding's inequality for sub-Gaussian random variables to see that for some fixed  $c > 0$ ,

$$\begin{aligned} \mathbb{P}(\mathcal{A}_{\mathcal{I},j}^c) &\leq \sum_{\ell=1}^{\mathcal{I}} \mathbb{P} \left( \left| \sum_{i=1}^n \phi_j(X_i^{(k)}) \phi_\ell(X_i^{(k)}) - \delta_{j,\ell} \right| > \sqrt{nt \log N} \right) \\ &\leq \sum_{\ell=1}^{\mathcal{I}} 2 \exp \left( - \frac{cnt \log N}{n \|\phi_j \phi_\ell - \delta_{j,\ell}\|_{\psi_2}^2} \right) \\ &\leq \mathcal{I} N^{-\frac{ct}{\sup_{\ell,j \in \mathbb{N}} \|\phi_j \phi_\ell - \delta_{j,\ell}\|_{\psi_2}^2}}. \end{aligned}$$

On the right-hand side,  $\|\cdot\|_{\psi_2}$  denotes the Orlicz norm. This norm is bounded by assumption (3.4), thus for any  $\tilde{C} > 0$ , there exists a  $t(\tilde{C}) > 0$  such that  $\mathbb{P}(\mathcal{A}_{\mathcal{I},j}^c) \lesssim \mathcal{I} N^{-6\tilde{C}}$ . Therefore

$$\begin{aligned} \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n K(\hat{g}^{(k)})(X_i^{(k)}) K C_{X_i^{(k)}} - \mathbb{E}_X(K(\hat{g}^{(k)})(X) K C_X) \right) \right\|_{S^\gamma}^2 \\ = \mathbb{E}_0 \left\| \tilde{R}_N \left( \sum_{j,\ell=1}^{\infty} \kappa_\ell \hat{g}_\ell^{(k)} \kappa_j^{\frac{1}{p}(d+2\alpha+p)} \left( \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i^{(k)}) \phi_j(X_i^{(k)}) - \delta_{\ell,j} \right) \right) \phi_j \right\|_{S^\gamma}^2 \\ = \mathbb{E}_0 \left\| \sum_{j,\ell=1}^{\infty} \kappa_\ell \hat{g}_\ell^{(k)} \nu_j \kappa_j^{-1} \left( \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i^{(k)}) \phi_j(X_i^{(k)}) - \delta_{\ell,j} \right) \phi_j \right\|_{S^\gamma}^2. \end{aligned}$$

We split the sum into  $\ell \leq \mathcal{I}$  and  $\ell > \mathcal{I}$ , so that the right-hand side of the last display is bounded by

$$\begin{aligned} 2\mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \left( \sum_{\ell=1}^{\mathcal{I}} \kappa_\ell \hat{g}_\ell^{(k)} \left( \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i) \phi_j(X_i) - \delta_{\ell,j} \right) \right)^2 \\ + 2\mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \left( \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_\ell \hat{g}_\ell^{(k)} \left( \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i) \phi_j(X_i) - \delta_{\ell,j} \right) \right)^2. \end{aligned}$$

By assumption (3.3), the empirical sum in the second term is  $O(1)$  almost surely. Then by applying Cauchy-Schwarz, we see that the preceding display is bounded by a multiple of

$$\begin{aligned}
 & \mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \left( \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell} \hat{g}_{\ell}^{(k)} \left( \frac{1}{n} \sum_{i=1}^n \phi_{\ell}(X_i) \phi_j(X_i) - \delta_{\ell,j} \right) \right)^2 \\
 & + \mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \left( \sum_{\ell=\mathcal{I}+1}^{\infty} |\kappa_{\ell} \hat{g}_{\ell}^{(k)}| \right)^2 \\
 & \leq \mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell}^2 \sum_{\ell=1}^{\mathcal{I}} (\hat{g}_{\ell}^{(k)})^2 \left( \frac{1}{n} \sum_{i=1}^n \phi_{\ell}(X_i) \phi_j(X_i) - \delta_{\ell,j} \right)^2 \\
 & + \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_{\ell}^{\frac{1}{p}(d+2\alpha+2p)} \mathbb{E}_0 \sum_{\ell=\mathcal{I}+1}^{\infty} (\kappa_{\ell} \hat{g}_{\ell}^{(k)})^2 \kappa_{\ell}^{-\frac{1}{p}(d+2\alpha+2p)}.
 \end{aligned}$$

Now conditioning on either  $\mathcal{A}_{\mathcal{I},j}$  or  $\mathcal{A}_{\mathcal{I},j}^c$ , again using the bound on the empirical sum, and Lemma 3.8.3 shows that the right-hand side is upper bounded by a constant multiple of

$$\begin{aligned}
 & \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell}^2 \mathbb{E}_0 \sum_{\ell=1}^{\mathcal{I}} (\hat{g}_{\ell}^{(k)})^2 \left( \frac{t \log N}{n} + 1_{\mathcal{A}_{\mathcal{I},j}^c} \right) \\
 & + \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_{\ell}^{\frac{1}{p}(d+2\alpha+2p)} \mathbb{E}_0 \|\hat{g}_{\mathcal{I}^c}^{(k)}\|_{\mathcal{H}}^2 \\
 & \lesssim \frac{\log N \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell}^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \mathbb{E}_0 \|\hat{g}^{(k)}\|_{L^2}^2 \\
 & + \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=1}^{\mathcal{I}} \kappa_{\ell}^2 N^Q \mathcal{I}^{\frac{1}{2}} N^{-3\tilde{C}} \\
 & + \mathbb{E}_0 \|\hat{g}_{\mathcal{I}^c}^{(k)}\|_{\mathcal{H}}^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_{\ell}^{\frac{1}{p}(d+2\alpha+2p)}.
 \end{aligned}$$

The first four factors in the second term on the right-hand side grow polynomially in  $N$ , so we can choose  $\tilde{C}$  large enough so that the term goes to zero arbitrarily fast for  $N \rightarrow \infty$ .  $\square$

**Lemma 3.7.4.** *There exists a constant  $C > 0$  such that for all  $f_0 \in S^\beta$ ,*

$$\mathbb{E}_0 \left\| \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2 \leq C \left( \frac{1}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}} \right) (\|\tilde{P}_N(f_0)\|_{S^\beta}^2 + \sigma^2).$$

*Proof.* Since  $\tilde{R}_N$  is linear, and  $\|f + g\|_{L^2}^2 \leq 2\|f\|_{L^2}^2 + 2\|g\|_{L^2}^2$ , and

$$\begin{aligned} & \hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)) \\ &= \frac{1}{n} \sum_{i=1}^n \epsilon_i^{(k)} K C_{X_i^{(k)}} + \frac{1}{n} \sum_{i=1}^n (K \tilde{P}_N(f_0))(X_i^{(k)}) K C_{X_i^{(k)}} - \frac{1}{N} \sigma^2 \tilde{V}_N(f_0), \end{aligned}$$

we have, using  $S_{n,m}^{(k)}(\tilde{V}_N(f_0)) = 0$ , that

$$\begin{aligned} & \mathbb{E}_0 \left\| \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2 \\ &= \mathbb{E}_0 \left\| \tilde{R}_N(\hat{S}_{n,m}^{(k)}(\tilde{V}_N(f_0)) - S_{n,m}^{(k)}(\tilde{V}_N(f_0))) \right\|_{S^\gamma}^2 \\ &= \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n \epsilon_i^{(k)} K C_{X_i^{(k)}} + \frac{1}{n} \sum_{i=1}^n (K \tilde{P}_N(f_0))(X_i^{(k)}) K C_{X_i^{(k)}} \right. \right. \\ &\quad \left. \left. - \mathbb{E}_X(K \tilde{P}_N(f_0)(X) K C_X) \right) \right\|_{S^\gamma}^2. \end{aligned}$$

From the fact that  $(a + b)^2 \leq 2a^2 + 2b^2$ , the right-hand side is bounded by two times

$$\begin{aligned} & \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n \epsilon_i^{(k)} K C_{X_i^{(k)}} \right) \right\|_{S^\gamma}^2 \\ &+ \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n (K \tilde{P}_N(f_0))(X_i^{(k)}) K C_{X_i^{(k)}} - \mathbb{E}_X(K(\tilde{P}_N f_0)(X) K C_X) \right) \right\|_{S^\gamma}^2 \\ &=: T_1 + T_2. \end{aligned}$$

In view of the independence of the noise, the first term  $T_1$  can be reformulated as

$$\begin{aligned} T_1 &= \frac{2\sigma^2}{n^2} \sum_{i=1}^n \mathbb{E}_0 \left( \left\| \tilde{R}_N(K C_{X_i^{(k)}}) \right\|_{S^\gamma}^2 \right) \\ &= \frac{2\sigma^2}{n^2} \sum_{i=1}^n \mathbb{E}_0 \left\langle \sum_{\ell=1}^{\infty} \kappa_\ell^{-\frac{\gamma}{p}} \kappa_\ell^{-1} \nu_\ell \phi_\ell(X_i) \phi_\ell, \sum_{j=1}^{\infty} \kappa_j^{-\frac{\gamma}{p}} \kappa_j^{-1} \nu_j \phi_j(X_i) \phi_j \right\rangle_2 \\ &= \frac{2\sigma^2}{n} \sum_{j=1}^{\infty} \kappa_j^{-2} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}}. \end{aligned}$$



For  $T_2$ , we use Lemma 3.8.1 with  $g = \tilde{P}_N(f_0)$  to see that

$$T_2 \lesssim \frac{\|\tilde{P}_N(f_0)\|_{S^\beta}^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}}.$$

since  $\tilde{P}_N(f_0) \in S^\beta$  and  $p + \beta > \beta > \frac{d}{2}$ . □

### 3.8 Technical lemmas

**Lemma 3.8.1.** *Assume  $t \geq 0$  is such that  $p + t > \frac{d}{2}$ . Then there exists a constant  $C > 0$ , depending only on  $K$ , such that for any  $\gamma \geq 0$  and  $g \in S^t$ , we have*

$$\begin{aligned} \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) KC_{X_i^{(k)}} - \mathbb{E}_X((Kg)(X)KC_X) \right) \right\|_{S^\gamma}^2 \\ \leq \frac{C \|g\|_{S^t}^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}}. \end{aligned} \quad (3.30)$$

*Proof.* By expanding both  $g$  and  $C_X$  into sums of eigenfunctions, we see that  $(Kg)(X)KC_X = \sum_{\ell=1}^{\infty} \kappa_\ell g_\ell \phi_\ell(X) \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha+p)} \phi_j(X) \phi_j$ . In view of this equality, we have

$$\mathbb{E}_X((Kg)(X)KC_X) = \sum_{j=1}^{\infty} g_j \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} \phi_j.$$

Hence we get the two identities

$$\begin{aligned} \tilde{R}_N(\mathbb{E}_X((Kg)(X)KC_X)) &= \sum_{j=1}^{\infty} g_j \nu_j \phi_j, \\ \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) KC_{X_i^{(k)}} \right) &= \frac{1}{n} \sum_{i=1}^n \sum_{\ell,j=1}^{\infty} \kappa_\ell g_\ell \phi_\ell(X_i^{(k)}) \kappa_j^{-1} \nu_j \phi_j(X_i^{(k)}) \phi_j. \end{aligned}$$

The left-hand side of Equation (3.30) is therefore equal to

$$\begin{aligned}
 & \mathbb{E}_0 \left\| \tilde{R}_N \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) KC_{X_i^{(k)}} - \mathbb{E}_X((Kg)(X)KC_X) \right) \right\|_{S^\gamma}^2 \\
 &= \mathbb{E}_0 \left\| \sum_{j=1}^{\infty} \nu_j \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) \kappa_j^{-1}(X_i^{(k)}) - g_j \right) \phi_j \right\|_{S^\gamma}^2 \\
 &= \mathbb{E}_0 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}} \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) \kappa_j^{-1} \phi_j(X_i^{(k)}) - g_j \right)^2.
 \end{aligned}$$

Since the expectation of the sum inside equals  $g_j$ , this is equal to the variance

$$\begin{aligned}
 & \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}} \text{Var} \left( \frac{1}{n} \sum_{i=1}^n (Kg)(X_i^{(k)}) \kappa_j^{-1} \phi_j(X_i^{(k)}) \right)^2 \\
 & \leq \frac{1}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}} \mathbb{E}_X((Kg)(X) \kappa_j^{-1} \phi_j(X))^2 \\
 & \lesssim \frac{1}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}} \kappa_j^{-2} \left( \sum_{i=1}^{\infty} \kappa_i^{1+\frac{t}{p}} \kappa_i^{-\frac{t}{p}} |g_i| \right)^2 \sup_{j \in \mathbb{N}} \mathbb{E}_0 \phi_j(X)^4.
 \end{aligned}$$

By assumption (3.4), we know that  $\sup_{j \in \mathbb{N}} \mathbb{E}_0 \phi_j(X)^4 < \infty$ . Thus from the Cauchy-Schwarz inequality, we see that the right-hand side of the last display is bounded up to a constant by

$$\frac{1}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-\frac{2\gamma}{p}} \kappa_j^{-2} \sum_{i=1}^{\infty} \kappa_i^{2+2\frac{t}{p}} \sum_{i=1}^{\infty} \kappa_i^{-\frac{2t}{p}} g_i^2 = \frac{\|g\|_{S^t}^2}{n} \sum_{i=1}^{\infty} \kappa_i^{2+2\frac{t}{p}} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \kappa_j^{-\frac{2\gamma}{p}}.$$

□

**Lemma 3.8.2.** *There exists a universal  $C > 0$  such that for all  $f_0 \in S^\beta(B)$  we have*

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{\mathcal{H}}^2 \leq CN.$$

*Proof.* By definition,

$$\|\Delta \hat{f}_n^{(k)}\|_{\mathcal{H}}^2 = \|\hat{f}_n^{(k)} - \tilde{V}_N(f_0)\|_{\mathcal{H}}^2 \leq 2\|\hat{f}_n^{(k)}\|_{\mathcal{H}}^2 + 2\|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2.$$

In view of the definition of  $\tilde{V}_N$ , the RKHS  $\mathcal{H}$  and Lemma 3.10.6 with  $q = \frac{\beta}{d}$ ,  $t = \frac{1}{d}(d + 2\alpha + 4p)$ ,  $u = \frac{1}{d}(d + 2\alpha + 2p)$  and  $v = 2$ , the second term is bounded

by

$$\begin{aligned}
 \|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2 &= \sum_{j=1}^{\infty} \kappa_j^{-\frac{1}{p}(d+2\alpha)} \nu_j^2 f_{0,j}^2 \\
 &= N^2 \sum_{j=1}^{\infty} \frac{\kappa_j^{\frac{1}{p}(d+2\alpha+4p)}}{(\sigma^2 + \kappa_j^{\frac{1}{p}(d+2\alpha+2p)} N)^2} f_{0,j}^2 \\
 &\asymp N^{1-2\frac{p+\beta}{d+2\alpha+2p}} \|f_0\|_{S^\beta}.
 \end{aligned} \tag{3.31}$$

Furthermore, since  $\hat{f}_n^{(k)}$  minimizes Equation (3.18), we have

$$\begin{aligned}
 \sigma^2 \|\hat{f}_n^{(k)}\|_{\mathcal{H}}^2 &\leq m \sum_{i=1}^n ((K\hat{f}_n^{(k)})(X_i^{(k)}) - Y_i^{(k)})^2 + \sigma^2 \|\hat{f}_n^{(k)}\|_{\mathcal{H}}^2 \\
 &\leq m \sum_{i=1}^n ((K\tilde{V}_N(f_0))(X_i^{(k)}) - (Kf_0)(X_i^{(k)}) - \epsilon_i^{(k)})^2 + \sigma^2 \|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2 \\
 &\leq 2m \sum_{i=1}^n (K\tilde{P}_N(f_0))(X_i^{(k)})^2 + 2m \sum_{i=1}^n (\epsilon_i^{(k)})^2 + \sigma^2 \|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2.
 \end{aligned} \tag{3.32}$$

Since  $n \asymp N^{1-\zeta}$ , the expectation of this then provides the bound

$$\begin{aligned}
 \mathbb{E}_0 \sigma^2 \|\hat{f}_n^{(k)}\|_{\mathcal{H}}^2 &\leq 2m \sum_{i=1}^n \mathbb{E}_0 (K\tilde{P}_N f_0)(X_i^{(k)})^2 + 2nm\sigma^2 + \sigma^2 \|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2 \\
 &= O(N).
 \end{aligned}$$

□

**Lemma 3.8.3.** *Let  $f_0 \in S^\beta$  such that  $p + \beta > \frac{d}{2}$ . Then for any event  $A$  that is independent of  $\sigma(\epsilon_1^{(k)}, \dots, \epsilon_n^{(k)})$ , we have*

$$\sup_{f_0 \in S^\beta(B)} \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 1_A \lesssim N^2 \mathbb{P}(A)^{\frac{1}{2}},$$

*Proof.* We can bound the expectation of  $\|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 1_A$  by Equation (3.32) as

$$\begin{aligned}
 \mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 1_A &\leq \mathbb{E}_0 2 \|\hat{f}_n^{(k)}\|_{\mathcal{H}}^2 1_A + 2 \mathbb{E}_0 \|\tilde{V}_N(f_0)\|_{L^2}^2 1_A \\
 &\lesssim \mathbb{E}_0 \left( m \sum_{i=1}^n ((K\tilde{P}_N f_0)(X_i^{(k)})^2 + (\epsilon_i^{(k)})^2) + \sigma^2 \|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2 \right) 1_A \\
 &\quad + \|\tilde{V}_N(f_0)\|_{L^2}^2 \mathbb{P}(A).
 \end{aligned}$$

For the first term inside the expectation, multiple applications of the Cauchy-Schwarz inequality show that for any indices  $j_1, j_2, j_3, j_4 \in \mathbb{N}$ ,

$$\mathbb{E}_0 |\phi_{j_1}(X) \phi_{j_2}(X) \phi_{j_3}(X) \phi_{j_4}(X)| \leq \sup_{i \in \mathbb{N}} \mathbb{E}_0 \phi_i(X)^4.$$

The right-hand side is finite by assumption (3.4). Thus we can bound the expectation  $\mathbb{E}_0 \sum_{i=1}^n K \tilde{P}_N(f_0)(X_i^{(k)})^2 1_A$  from above by

$$\begin{aligned} \mathbb{E}_0 (K \tilde{P}_N f_0)(X_i^{(k)})^2 1_A &\leq \sqrt{\mathbb{E}_0 \left( (K \tilde{P}_N f_0)(X_i^{(k)}) \right)^4} \mathbb{E}_0 1_A \\ &= \mathbb{P}(A)^{\frac{1}{2}} \left( \sum_{\substack{j_1, j_2, \\ j_3, j_4=1}}^{\infty} \prod_{\ell=1}^4 (\kappa_{j_\ell} (1 - \nu_{j_\ell}) f_{0,j_\ell}) \mathbb{E}_0 \prod_{\ell=1}^4 \phi_{j_\ell}(X) \right)^{\frac{1}{2}} \\ &\leq \mathbb{P}(A)^{\frac{1}{2}} \left( \sup_{i \in \mathbb{N}} \mathbb{E}_0 \phi_i(X)^4 \right)^{\frac{1}{2}} \left( \sum_{j=1}^{\infty} \kappa_j (1 - \nu_j) |f_{0,j}| \right)^2 \\ &\lesssim \mathbb{P}(A)^{\frac{1}{2}} \sum_{j=1}^{\infty} \kappa_j^{2+\frac{2\beta}{p}} \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\beta}{p}} f_{0,j}^2. \end{aligned}$$

The first sum is finite by the assumption on  $\beta$ , the second is finite since  $f_0 \in S^\beta$ . Note that by Equation (3.31) the RKHS-norm of  $\tilde{V}_N(f_0)$  is bounded by

$$\|\tilde{V}_N(f_0)\|_{\mathcal{H}}^2 \lesssim \frac{N^2}{\sigma^4} \sum_{j=1}^{\infty} f_{0,j}^2.$$

Combining the previous bounds results in

$$\mathbb{E}_0 \|\Delta \hat{f}_n^{(k)}\|_{L^2}^2 1_A \lesssim \mathbb{P}(A)^{\frac{1}{2}} (N + \sigma^2 N + N^2).$$

finishing the proof.  $\square$

### 3.9 Local posterior variance

In this section, we will derive the asymptotic behaviour of the expected local posterior variance with respect to the  $S^\gamma$ -norm. This quantity is given by

$$\mathbb{E}_0 (\Pi_{n,m}^{(k)} (\|f - \hat{f}_{n,m}^{(k)}\|_{S^\gamma} \mid \mathbb{D}_n^{(k)})).$$

To simplify the exposition, the superscript  $(k)$  is dropped in the remainder of the section.

We will first derive an asymptotic upper bound for the posterior variance on a local machine. The main step in the proof is by observing that the posterior mean is the minimizer of the Bayes risk with respect to the squared  $S^\gamma$ -norm. Hence the posterior variance with respect to this norm is upper bounded by the expected squared  $S^\gamma$ -norm of the difference of  $f$  and any  $\sigma(\mathbb{D}_n)$ -measurable random variable. We will thus start with constructing an estimator for which the  $S^\gamma$ -norm is analytically tractable, and afterwards show that this asymptotically gives a sufficiently tight bound.

Denote the local covariance kernel function by

$$C_\alpha(s, t) = m \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\alpha)} \phi_j(s) \phi_j(t).$$

Recall that the prior on a local machine is defined as  $\Pi_m = N(0, C_\alpha)$ , and that the local posterior is denoted by  $\Pi_m(\cdot \mid \mathbb{D}_n)$ . The data is denoted by  $\mathbb{X}_n = \{X_1, \dots, X_n\}$ ,  $\mathbb{Y}_n = (Kf(X_1) + \epsilon_1, \dots, Kf(X_n) + \epsilon_n)^\top$  and  $\mathbb{D}_n = \{(X_i, Y_i) : i = 1, \dots, n\}$ .

We first derive an upper bound for the asymptotic variance of the posterior with respect to the  $S^\gamma$ -norm.

**Lemma 3.9.1.** *Suppose that  $\sup_{j \in \mathbb{N}} \mathbb{E}_X \phi_j(X)^6 < \infty$ . Then, for  $m \lesssim (nm)^{\frac{2\alpha+2p-d}{2\alpha+2p+3d}}$ ,*

$$\mathbb{E}_0(\Pi_m(\|f - \hat{f}_{n,m}\|_{S^\gamma}^2 \mid \mathbb{D}_n)) \lesssim m(nm)^{-\frac{2\alpha-2\gamma}{2\alpha+2p+d}},$$

which is of the same order as  $\frac{1}{n} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}}$ .

*Proof.* The posterior of  $f$  given  $\mathbb{D}_n$  depends only on  $\mathbb{Y}_n$  through the posterior mean, since the posterior covariance is only a function of  $\mathbb{X}_n$  and is independent of  $\mathbb{Y}_n$  (see Section 3.10). Therefore, the posterior distribution of  $f - \hat{f}_{n,m}$  only depends on  $\mathbb{X}_n$  and we can write  $\Pi_m(\|f - \hat{f}_{n,m}\|_{S^\gamma} \mid \mathbb{D}_n) = \Pi_m(\|f - \hat{f}_{n,m}\|_{S^\gamma} \mid \mathbb{X}_n)$ . Since the posterior mean is the minimizer of the Bayes risk for the squared  $S^\gamma$ -norm loss, for any matrix  $\tilde{A}_{\mathbb{X}_n}$  indexed by  $\mathbb{N} \times n$  (possibly depending on  $\mathbb{X}_n$ ), we have  $\Pi_m(\|f - \hat{f}_{n,m}\|_{S^\gamma}^2 \mid \mathbb{X}_n) \leq \Pi_m(\|f - \tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n\|_{S^\gamma}^2 \mid \mathbb{X}_n)$ .

We will now derive a suitable choice for  $\tilde{A}_{\mathbb{X}_n}$ . Define the vector  $\vec{f} = (f_1, f_2, \dots)^\top$ , and the matrix  $\Phi$  indexed by  $n \times \mathbb{N}$  by  $\Phi_{i,\ell} = \phi_i(X_\ell)$  for  $1 \leq i \leq n$  and  $\ell \in \mathbb{N}$ . Furthermore, let  $\vec{K}$  denote the matrix form of the operator  $K$ , so  $\vec{K} \in \mathbb{R}^{\mathbb{N} \times \mathbb{N}}$  with  $\vec{K}_{ij} = \kappa_j \delta_{ij}$ . Define the matrix form of the operator  $\tilde{V}$  as

$\vec{V}_n \in \mathbb{R}^{N \times N}$  by  $\vec{V}_{n,\ell,\ell'} = \nu_\ell \delta_{\ell,\ell'}$ , for  $\ell, \ell' \in \mathbb{N}$ . Lastly, define the matrix  $B$  indexed by  $\mathbb{N} \times \mathbb{N}$  by

$$B_{\ell,\ell'} = \frac{1}{n} \sum_{i=1}^n (\phi_\ell(X_i) \phi_{\ell'}(X_i) - \delta_{\ell,\ell'}),$$

for  $\ell, \ell' \in \mathbb{N}$ .

Section 3.10 shows that one can write  $\hat{f}_{n,m} = A_{\mathbb{X}_n} \mathbb{Y}_n$  for some matrix  $A_{\mathbb{X}_n}$  that is indexed by  $\mathbb{N} \times n$ . By applying the conditional independence of  $f - \hat{f}_{n,m}$  and  $\mathbb{Y}_n$  given  $\mathbb{X}_n$ , and by writing  $\mathbb{Y}_n = \Phi \vec{K} \vec{f} + \epsilon$ , we see that

$$\begin{aligned} 0 &= \mathbb{E}((\vec{f} - \hat{f}_{n,m}) \mathbb{Y}_n^\top \mid \mathbb{X}_n) \\ &= \mathbb{E}(\vec{f} (\Phi \vec{K} \vec{f})^\top \mid \mathbb{X}) - A_{\mathbb{X}_n} \mathbb{E}((\Phi \vec{K} \vec{f} + \epsilon)(\Phi \vec{K} \vec{f} + \epsilon)^\top \mid \mathbb{X}_n). \end{aligned}$$

Here the left-hand side should be interpreted as the null matrix in  $\mathbb{R}^{N \times n}$ . Because the matrix  $\mathbb{E}((\Phi \vec{K} \vec{f} + \epsilon)(\Phi \vec{K} \vec{f} + \epsilon)^\top \mid \mathbb{X}_n)$  is invertible (as it is the sum of  $\sigma^2 I$  and a non-negative matrix), this can be solved to find

$$A_{\mathbb{X}_n} = m \vec{K}^{-1} (\sigma^2 \vec{K}^{-\frac{1}{p}(2\alpha+2p+d)} + m \Phi^\top \Phi)^{-1} \Phi^\top. \quad (3.33)$$

By writing

$$\sigma^2 \vec{K}^{-\frac{1}{p}(2\alpha+2p+d)} + m \Phi^\top = (\sigma^2 \vec{K}^{-\frac{1}{p}(2\alpha+2p+d)} + nmI) + (m \Phi^\top \Phi - nmI),$$

we choose  $\tilde{A}_{\mathbb{X}_n}$  by replacing the inverse matrix in Equation (3.33) by its first order Taylor expansion around  $\sigma^2 \vec{K}^{-\frac{1}{p}(2\alpha+2p+d)} + nmI$ . This means that  $\tilde{A}_{\mathbb{X}_n} = \vec{K}^{-1} (\vec{V}_n - \vec{V}_n B \vec{V}_n) n^{-1} \Phi^\top$ .

We will now show that  $\Pi_m(\|f - \tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n\|_{S^\gamma}^2 \mid \mathbb{X}_n)$  is asymptotically of sufficiently small order for this choice of  $\tilde{A}_{\mathbb{X}_n}$ . Defining the vector  $(C_\ell)_{\ell \in \mathbb{N}}$  by  $C_\ell = \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i) \epsilon_i$ , we see that

$$\left( \frac{1}{n} \Phi^\top \mathbb{Y}_n \right)_\ell = \frac{1}{n} \sum_{i=1}^n \phi_\ell(X_i) Y_i = \sum_{r=1}^{\infty} \kappa_r f_r(B_{\ell,r} + \delta_{\ell,r}) + C_\ell.$$

This can be written as  $n^{-1} \Phi^\top \mathbb{Y}_n = \vec{K} \vec{f} + B \vec{K} \vec{f} + C$ , which shows that

$$\begin{aligned} \vec{f} - \tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n &= (I - \vec{V}_n) \vec{f} - \vec{K}^{-1} \vec{V}_n B (I - \vec{V}_n) \vec{K} \vec{f} + \vec{K}^{-1} \vec{V}_n B \vec{V}_n B \vec{K} \vec{f} \\ &\quad - \vec{K}^{-1} \vec{V}_n (I - B \vec{V}_n) C. \end{aligned} \quad (3.34)$$

Since the prior on  $f$  is centred, we have  $\Pi_m(f - \tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n \mid \mathbb{X}_n) = 0$ . Combined with the mutual independence of the coefficients of  $f$ , we conclude that

$$\Pi_m(\|\vec{f} - \tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n\|_{S^\gamma}^2 \mid \mathbb{X}_n) = \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} \text{Var}(f_j - (\tilde{A}_{\mathbb{X}_n} \mathbb{Y}_n)_j \mid \mathbb{X}_n).$$

We bound the variances of the last display by the sum of the variances of the four terms on the right-hand side of Equation (3.34). This shows that the previous display can be bounded by a constant multiple of

$$\begin{aligned} & \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} \left[ (1 - \nu_j)^2 \text{Var}(f_j) + \frac{\nu_j^2}{\kappa_j^2} \text{Var}\left(\sum_{\ell=1}^{\infty} B_{j,\ell}(1 - \nu_\ell) \kappa_\ell f_\ell \mid \mathbb{X}_n\right) \right. \\ & \left. + \frac{\nu_j^2}{\kappa_j^2} \text{Var}\left(\sum_{\ell=1}^{\infty} \sum_{r=1}^{\infty} B_{j,\ell} \nu_\ell B_{\ell,r} \kappa_r f_r \mid \mathbb{X}_n\right) + \frac{\nu_j^2}{\kappa_j^2} \text{Var}\left(C_j - \sum_{\ell=1}^{\infty} B_{j,\ell} \nu_\ell C_\ell \mid \mathbb{X}_n\right) \right]. \end{aligned} \quad (3.35)$$

Under the prior distribution, the random variables  $f_1, f_2, \dots$  are mutually independent, independent of  $\mathbb{X}_n$  and satisfy  $\text{Var} f_j^2 = m \kappa_j^{\frac{1}{p}(2\alpha+d)}$ . We find that the expectation of the sum of the first three series in the previous display is equal to

$$\begin{aligned} \mathbb{E}_0 \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} & \left[ (1 - \nu_j)^2 m \kappa_j^{\frac{2\alpha+d}{p}} + \frac{\nu_j^2}{\kappa_j^2} \sum_{\ell=1}^{\infty} B_{j,\ell}^2 (1 - \nu_\ell)^2 \kappa_\ell^2 m \kappa_\ell^{\frac{2\alpha+d}{p}} \right. \\ & \left. + \frac{\nu_j^2}{\kappa_j^2} \sum_{r=1}^{\infty} \left( \sum_{\ell=1}^{\infty} B_{j,\ell} \nu_\ell B_{\ell,r} \right)^2 \kappa_r^2 m \kappa_r^{\frac{2\alpha+d}{p}} \right]. \end{aligned}$$

Note that  $\mathbb{E}_0 B_{\ell,j}^2 = \text{Var}_0 B_{\ell,j} \lesssim n^{-1}$ , which implies that  $\mathbb{E}_0 (\sum_{\ell=1}^{\infty} B_{j,\ell} \nu_\ell B_{\ell,r})^2 \lesssim n^{-2} (\sum_{\ell=1}^{\infty} \nu_\ell)^2$  by the Cauchy-Schwarz inequality. Furthermore, since  $p + \alpha > 0$ , we have  $\sum_{r=1}^{\infty} \kappa_r^2 \kappa_r^{\frac{2\alpha+d}{p}} < \infty$ . Combining these observations, shows that the previous display is upper bounded by a constant multiple of

$$\begin{aligned} & m \sum_{j=1}^{\infty} (1 - \nu_j)^2 \kappa_j^{\frac{2\alpha-2\gamma+d}{p}} + \frac{m}{n} \sum_{j=1}^{\infty} \kappa_j^{-2-\frac{2\gamma}{p}} \nu_j^2 \sum_{\ell=1}^{\infty} (1 - \nu_\ell)^2 \kappa_\ell^{\frac{2\alpha+2p+d}{p}} \\ & + \frac{m}{n^2} \sum_{j=1}^{\infty} \kappa_j^{-2-\frac{2\gamma}{p}} \nu_j^2 \left( \sum_{\ell=1}^{\infty} \nu_\ell \right)^2. \end{aligned}$$

Applying Lemma 3.10.7 to the first sum on the right-hand side of the last display, with  $\frac{1}{d}(2\alpha-2\gamma+d)$ ,  $u = \frac{1}{d}(2\alpha+2p+d)$  and  $v = 2$ , shows that it is asymptotically

of the order  $m(nm)^{-\frac{2\alpha-2\gamma}{2\alpha+2p+d}}$ . Using the same lemma applied with  $t = \frac{1}{d}(2\alpha + 2p + d)$ ,  $u = \frac{1}{d}(2\alpha + 2p + d)$  and  $v = 2$ , we have  $\sum_{\ell=1}^{\infty} (1 - \nu_{\ell}^2) \kappa_{\ell}^{\frac{1}{p}(2\alpha+2p+d)} \asymp (nm)^{-\frac{2\alpha+2p}{2\alpha+2p+d}}$ . By the same lemma with  $t = \frac{1}{d}(2\alpha + 2p + d)$ ,  $u = \frac{1}{d}(2\alpha + 2p + d)$  and  $v = 1$ , we find that  $\sum_{\ell=1}^{\infty} (nm)^{\frac{d}{2\alpha+2p+d}}$ . Finally, using the lemma with  $t = \frac{1}{d}(4\alpha + 2p + 2d - 2\gamma)$ ,  $u = \frac{1}{d}(2\alpha + 2p + d)$  and  $v = 2$ , we see that  $\sum_{j=1}^{\infty} \kappa_j^{-2-\frac{2\gamma}{p}} \nu_j^2 \asymp n(nm)^{-\frac{2\alpha-2\gamma}{2\alpha+2p+d}}$ . Thus the right-hand side of the preceding display is of the order

$$m(nm)^{-\frac{2\alpha-2\gamma}{2\alpha+2p+d}} \left[ 1 + m(nm)^{-\frac{2\alpha+2p}{2\alpha+2p+d}} + \frac{m}{n} (nm)^{\frac{2d}{2\alpha+2p+d}} \right].$$

The terms within the square bracket are bounded for  $m \lesssim n^{\frac{2\alpha+2p}{d}} \wedge n^{\frac{2\alpha+2p-d}{2\alpha+2p+3d}}$ . By assumption,  $\alpha + p > \frac{d}{2}$ , hence the first restriction is satisfied.

Since  $\text{Var}(C_j \mid \mathbb{X}_n) = \frac{\sigma^2}{n}(B_{j,j} + 1)$  and  $\text{Cov}(C_{\ell}, C_{\ell'} \mid \mathbb{X}_n) = \frac{\sigma^2}{n} B_{\ell,\ell'}$ , for  $\ell \neq \ell'$ , the expectation of the fourth sum in the bound (3.35) is bounded by a constant multiple of

$$\begin{aligned} \mathbb{E}_0 \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\gamma}{p}} \frac{\nu_j^2}{\kappa_j^2} \frac{\sigma^2}{n} \left( B_{j,j} + 1 + \sum_{\ell=1}^{\infty} \sum_{\ell'=1}^{\infty} B_{\ell,j} B_{\ell',j} \nu_{\ell} \nu_{\ell'} (B_{\ell,\ell'} + \delta_{\ell,\ell'}) \right) \\ \lesssim \sum_{j=1}^{\infty} \kappa_j^{-2-\frac{2\gamma}{p}} \nu_j^2 \frac{1}{n} \left( 1 + \frac{1}{n} \sum_{\ell=1}^{\infty} \nu_{\ell}^2 + \frac{1}{n\sqrt{n}} \left( \sum_{\ell=1}^{\infty} \nu_{\ell} \right)^2 \right), \end{aligned}$$

since  $\mathbb{E}_0 |B_{\ell,j} B_{\ell',j} B_{\ell,\ell'}| \lesssim \max_{\ell,\ell'} \mathbb{E}_0 |B_{\ell,\ell'}|^3 \lesssim n^{-\frac{3}{2}}$ . The first inequality is due to Hölder's inequality, while the second inequality is by the Marcinkiewics-Zygmund inequality applied to the centred random variables  $B_{\ell,\ell'}$ . Lastly, by applying Lemma 3.10.7 with  $t = \frac{1}{d}(4\alpha + 4p + 2d)$ ,  $u = \frac{1}{d}(2\alpha + 2p + d)$  and  $v = 2$ , we find that  $\sum_{\ell=1}^{\infty} \nu_{\ell}^2 \asymp (nm)^{\frac{d}{2\alpha+2p+d}}$ . Therefore, the right-hand side of the last display is of the order

$$m(nm)^{-\frac{2\alpha-2\gamma}{2\alpha+2p+d}} \left[ 1 + \frac{1}{n} (nm)^{\frac{d}{2\alpha+2p+d}} + \frac{1}{n\sqrt{n}} (nm)^{\frac{2d}{2\alpha+2p+d}} \right].$$

The term inside the square brackets is bounded if  $m \lesssim n^{\frac{2\alpha+2p}{d}} \wedge n^{\frac{3\alpha+3p-\frac{d}{2}}{2d}}$ , which is satisfied as well. This shows the first part of the claim.

In order to show that the bound is asymptotically of the same order as  $\frac{1}{n} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2-\frac{2\gamma}{p}}$ , we use Equation (3.21) with  $\beta$  replaced by  $\alpha$ . It shows that



the latter sum is of the order  $\frac{N}{n} N^{-\frac{2\alpha-2\gamma}{d+2\alpha+2p}}$ . The proof is finished by noting that  $N = nm$  by assumption.  $\square$

To establish good frequentist coverage by the credible sets coming from the aggregated posterior, a lower bound on the expected posterior variance is also needed. Establishing the lower bound is accomplished by first deriving an alternative expression for the posterior covariance. Since this new expression is the trace of the inverse of a matrix, it can be lower bounded by the sum of the inverse of the diagonal elements of this matrix. This sum can be easily analysed and provides a sufficiently tight bound. The following lemma provides the complete details of this approach.

**Lemma 3.9.2.** *Let  $\zeta < \frac{1}{2}$  and  $\beta > \frac{d}{2}$  be given and assume that the eigenfunctions are uniformly bounded. Let  $f \sim \Pi_\beta$ . Then for  $n \rightarrow \infty$ ,*

$$\mathbb{E}_0(\Pi_{n,m}(\|f - \hat{f}_{n,m}\|_{L^2}^2 \mid \mathbb{D}_n)) \asymp \frac{\sigma^2}{n} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2}.$$

*Proof.* The asymptotic upper bound was proven in Lemma 3.9.1 for  $\gamma = 0$ . For the lower bound, we follow the proof in (Oppor and Vivaerlli 1999). We first derive an alternative expression for the posterior variance at an arbitrary point  $x \in \mathcal{O}$ .

Let the operator  $\Lambda : L^2 \rightarrow L^2$  be defined by  $\Lambda f = \sum_{j=1}^{\infty} m \kappa_j^{\frac{d}{p}} f_j \phi_j$  for arbitrary  $f \in L^2$ . Then  $f \sim \Pi_\beta$  is equivalent to  $f \sim N(0, \Lambda)$  as a Gaussian process defined on  $S^\beta$ . For fixed  $x_1, \dots, x_n \in \mathcal{O}$ , we see that  $Y \mid f \sim N(\hat{K}f, \sigma^2 I)$  for the operator  $\hat{K} : S^\beta \rightarrow \mathbb{R}^n$  defined by

$$\hat{K}f = \hat{K} \sum_{j=1}^{\infty} f_j \phi_j = \left[ \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_1), \dots, \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_n) \right]^\top.$$

To show that  $K$  is continuous, consider for an arbitrary  $f \in S^\beta$  the squared norm of  $\hat{K}f$

$$\begin{aligned} \|\hat{K}f\|_{\mathbb{R}^n}^2 &= \sum_{i=1}^n \left( \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_i) \right)^2 \\ &\lesssim n \left( \sum_{j=1}^{\infty} \kappa_j^{2+\frac{2\beta}{p}} \right)^2 \left( \sum_{j=1}^{\infty} f_j^2 \kappa_j^{-\frac{2\beta}{p}} \right)^2 \\ &\asymp \left( \sum_{j=1}^{\infty} \kappa_j^{2+\frac{2\beta}{p}} \right)^2 \|f\|_{S^\beta}^2. \end{aligned}$$

Therefore  $\hat{K}$  is a bounded operator whenever

$$\sum_{j=1}^{\infty} \kappa_j^{2+\frac{2\beta}{p}} \asymp \sum_{j=1}^{\infty} j^{-\frac{p}{d}(2+\frac{2\beta}{p})} < \infty.$$

We find that  $\hat{K}$  is continuous if  $\frac{p}{d}(2+\frac{2\beta}{p}) > 1$ , which is equivalent to  $p + \beta > \frac{d}{2}$ .

Since  $f$  is endowed with the prior  $f \sim N(0, \Lambda)$ , we derive from Proposition 3.1 in (Knapik, van der Vaart and van Zanten 2011) that the posterior covariance is equal to

$$W_n = \Lambda - A(\sigma^2 I + \hat{K} \Lambda \hat{K}^\top) A^\top,$$

with

$$A = \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda^{1/2} \hat{K}^\top = \Lambda \hat{K}^\top (\sigma^2 I + \hat{K} \Lambda \hat{K}^\top)^{-1}.$$

We see that

$$\begin{aligned} W_n &= \Lambda - A(\sigma^2 I + \hat{K} \Lambda \hat{K}^\top) A^\top \\ &= \Lambda - \Lambda \hat{K}^\top (\sigma^2 I + \hat{K} \Lambda \hat{K}^\top)^{-1} (\sigma^2 I + \hat{K} \Lambda \hat{K}^\top) (\sigma^2 I + \hat{K} \Lambda \hat{K}^\top)^{-1} \hat{K} \Lambda^\top \\ &= \Lambda - \Lambda \hat{K}^\top (\sigma^2 I + \hat{K} \Lambda \hat{K}^\top)^{-1} \hat{K} \Lambda. \end{aligned}$$

Using the first form of  $A$ , we find

$$\begin{aligned} W_n &= \Lambda - \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda \\ &= \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} [\sigma^2 \Lambda^{1/2} + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda - \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda] \\ &= \sigma^2 \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2}) \Lambda^{1/2} \end{aligned}$$

We now consider the posterior covariance at a single point  $x \in \mathcal{O}$ . Let  $T_x : S^\beta \rightarrow \mathbb{R}$  be the point-evaluation operator

$$T_x f = f(x) = \sum_{j=1}^{\infty} f_j \phi_j(x).$$

Due to the uniform boundedness of the eigenfunctions,  $T_x$  is a continuous linear operator if  $\beta > \frac{d}{2}$ . Note that we have  $T_x^\top s = s \sum_{j=1}^{\infty} \kappa_j^{\frac{2\beta}{p}} \phi_j(x) \phi_j$ . Then the posterior variance of  $f$  at a point  $x \in \mathcal{O}$  is equal to

$$T_x W_n T_x^\top = \sigma^2 T_x \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda^{1/2} T_x^\top.$$

Since this is a scalar, it is equal to its trace

$$\begin{aligned} & \sigma^2 \text{tr} \left( T_x \Lambda^{1/2} (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda^{1/2} T_x^\top \right) \\ &= \sigma^2 \text{tr} \left( (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda^{1/2} T_x^\top T_x \Lambda^{1/2} \right). \end{aligned}$$

The product  $\Lambda^{1/2} T_x^\top T_x \Lambda^{1/2}$  equals

$$\Lambda^{1/2} T_x^\top T_x \Lambda^{1/2} f = m \sum_{j,\ell=1}^{\infty} \kappa_\ell^{\frac{2\beta}{p}} (\kappa_j \kappa_\ell)^{\frac{d}{2p}} f_j \phi_j(x) \phi_\ell(x) \phi_\ell.$$

Therefore its expectation over  $x$  simplifies to

$$\begin{aligned} \mathbb{E}_X \Lambda^{1/2} T_X^\top T_X \Lambda^{1/2} f &= m \sum_{j,\ell=1}^{\infty} \kappa_\ell^{\frac{2\beta}{p}} (\kappa_j \kappa_\ell)^{\frac{d}{2p}} f_j \mathbb{E}_X (\phi_j(X) \phi_\ell(X)) \phi_\ell \\ &= m \sum_{j=1}^{\infty} \kappa_j^{\frac{2\beta}{p}} \kappa_j^{\frac{d}{p}} f_j \phi_j = \Lambda \Gamma f, \end{aligned}$$

with  $\Gamma : L^2 \rightarrow L^2$  the operator defined as  $\Gamma f = \sum_{j=1}^{\infty} \kappa_j^{\frac{2\beta}{p}} f_j$ . We find that the integrated posterior variance equals

$$\begin{aligned} \mathbb{E}_X T_X W_n T_X^\top &= \sigma^2 \text{tr} \left( (\sigma^2 I + \Lambda^{1/2} \hat{K}^\top \hat{K} \Lambda^{1/2})^{-1} \Lambda \Gamma \right) \\ &= \sigma^2 \text{tr} \left( (\sigma^2 \Gamma^{-1} \Lambda^{-1/2} + \Lambda^{-1/2} \hat{K}^\top \hat{K} \Lambda^{-1/2})^{-1} \right). \end{aligned} \quad (3.36)$$

In order to analyse the previous display, we derive an expression for the term  $\Gamma^{-1/2} \hat{K}^\top \hat{K} \Gamma^{-1/2}$ . For arbitrary  $f \in S^\beta$  and  $\mu = [\mu_1, \dots, \mu_n]^\top \in \mathbb{R}^n$  we have

$$\langle \hat{K} f, \mu \rangle_{\mathbb{R}^n} = \sum_{i=1}^n \mu_i \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_i) = \sum_{j=1}^{\infty} f_j \kappa_j \sum_{i=1}^n \mu_i \phi_j(x_i)$$

and

$$\langle f, \hat{K}^\top \mu \rangle_{S^\beta} = \sum_{j=1}^{\infty} \kappa_j^{-\frac{2\beta}{p}} f_j (\hat{K}^\top \mu)_j.$$

Since by the definition of the adjoint operator  $\hat{K}^\top$  the two previous displays must be equal, for any  $j \in \mathbb{N}$  we have

$$(\hat{K}^\top \mu)_j = \kappa_j \kappa_j^{\frac{2\beta}{p}} \sum_{i=1}^n \mu_i \phi_j(x_i).$$

Combining this with the definition of  $\hat{K}$ , we see that for any  $\ell \in \mathbb{N}$

$$\begin{aligned} (\hat{K}^\top \hat{K} f)_\ell &= \left( \hat{K}^\top \left[ \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_1), \dots, \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_n) \right] \right)_\ell \\ &= \kappa_\ell \kappa_\ell^{\frac{2\beta}{p}} \sum_{i=1}^n \left( \sum_{j=1}^{\infty} \kappa_j f_j \phi_j(x_i) \right) \phi_\ell(x_i) \\ &= \kappa_\ell \kappa_\ell^{\frac{2\beta}{p}} \sum_{j=1}^{\infty} \kappa_j f_j \sum_{i=1}^n \phi_j(x_i) \phi_\ell(x_i). \end{aligned}$$

By combining the last display with the definition of  $\Gamma$ , we see that the second operator on the right-hand side of Equation (3.36) equals

$$\begin{aligned} \Gamma^{-1/2} \hat{K}^\top \hat{K} \Gamma^{-1/2} f &= \sum_{\ell=1}^{\infty} \kappa_\ell^{-\frac{\beta}{p}} \left( \kappa_\ell \kappa_\ell^{\frac{2\beta}{p}} \sum_{j=1}^{\infty} \kappa_j \kappa_j^{-\frac{\beta}{p}} f_j \sum_{i=1}^n \phi_j(x_i) \phi_\ell(x_i) \right) \phi_\ell. \end{aligned}$$

Lemma 1 of (Oppor and Vivaerlli 1999) shows that for any real, symmetric matrix  $(A_{ij})_{i,j \in \mathcal{J}}$ ,  $\mathcal{J} \subseteq \mathbb{N}$ , and real-valued, convex function  $\psi$ , the trace of  $\psi(A)$  is lower bounded by  $\text{tr } \psi(A) \geq \sum_{j \in \mathcal{J}} \psi(A_{jj})$ . We apply this lemma with  $\psi : (0, \infty) \rightarrow \mathbb{R}$  equal to the function  $x \mapsto \frac{1}{x}$ , and  $A$  equal to the positive-definite matrix on the right-hand side of Equation (3.36). Combined with the previous display, this shows that the posterior variance is bounded from below by

$$\begin{aligned} \sigma^2 \sum_{j=1}^{\infty} \left( \sigma^2 \kappa_j^{-\frac{2\beta}{p}} m^{-1} \kappa_j^{-\frac{d}{p}} + \kappa_j^2 \sum_{i=1}^n \phi_j(x_i) \phi_j(x_i) \right)^{-1} \\ = \sum_{j=1}^{\infty} \frac{\sigma^2 m \kappa_j^{\frac{1}{p}(d+2\beta)}}{\sigma^2 + m \kappa_j^{\frac{1}{p}(d+2\beta+2p)} \sum_{i=1}^n \phi_j(x_i)^2}. \end{aligned}$$

Since the eigenfunctions are uniformly bounded by assumption, this lower bound is asymptotically equivalent to

$$\sum_{j=1}^{\infty} \frac{\sigma^2 m \kappa_j^{\frac{1}{p}(d+2\beta)}}{\sigma^2 + n m \kappa_j^{\frac{1}{p}(d+2\beta+2p)} C_\phi^2} \asymp \frac{\sigma^2}{n} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2}.$$

The sum on the right-hand side is now independent of  $x_1, \dots, x_n$ , so finally taking the expectation with respect to  $\mathbb{X}_n$  does not change the expression.  $\square$

### 3.10 Proof of Theorem 3.3.2

Given  $\mathbb{X}_n^{(k)} = (X_1^{(k)}, \dots, X_n^{(k)})^\top$  and  $\mathbb{Y}_n^{(k)} = (Y_1^{(k)}, \dots, Y_n^{(k)})^\top$  and the observations  $\mathbb{D}_n^{(k)} = \{(X_1^{(k)}, Y_1^{(k)}), \dots, (X_n^{(k)}, Y_n^{(k)})\}$ , we write

$$\begin{aligned} C_\beta((x, x')^\top, (x, x')) &= \begin{pmatrix} C_\beta(x, x) & C_\beta(x, x') \\ C_\beta(x', x) & C_\beta(x', x') \end{pmatrix} \\ C_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) &= \begin{pmatrix} C_{\beta+p}(X_1^{(k)}, X_1^{(k)}) & \dots & C_{\beta+p}(X_1^{(k)}, X_n^{(k)}) \\ \vdots & \ddots & \vdots \\ C_{\beta+p}(X_n^{(k)}, X_1^{(k)}) & \dots & C_{\beta+p}(X_n^{(k)}, X_n^{(k)}) \end{pmatrix} \\ C_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, (x, x')) &= \begin{pmatrix} C_{\beta+\frac{p}{2}}(X_1^{(k)}, x) & C_{\beta+\frac{p}{2}}(X_1^{(k)}, x') \\ \vdots & \vdots \\ C_{\beta+\frac{p}{2}}(X_n^{(k)}, x) & C_{\beta+\frac{p}{2}}(X_n^{(k)}, x') \end{pmatrix}. \end{aligned}$$

Using these matrices, we know that

$$\begin{pmatrix} Y_1^{(k)} \\ \vdots \\ Y_n^{(k)} \\ f(x) \\ f(x') \end{pmatrix} \sim \mathcal{N}\left(0, \begin{pmatrix} mC_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) + \sigma^2 I & mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, (x, x')) \\ mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, (x, x'))^\top & mC_\beta((x, x')^\top, (x, x')) \end{pmatrix}\right).$$

Using the formula for conditional Gaussian distributions, we can write the local posterior distribution now as

$$(f(x), f(x'))^\top \mid \mathbb{Y}_n^{(k)} \sim \mathcal{N}(\hat{f}_n^{(k)}, \hat{C}_n^{(k)})$$

where the local posterior mean is equal to

$$\hat{f}_n^{(k)}(x) = mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, x)^\top [mC_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) + \sigma^2 I]^{-1} \mathbb{Y}_n^{(k)}$$

and the local posterior covariance matrix  $\hat{C}_n^{(k)}$  has the form

$$\begin{aligned} &mC_\beta((x, x')^\top, (x, x')) \\ &- C_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, (x, x'))^\top [C_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) + \frac{1}{m}\sigma^2 I]^{-1} mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, (x, x')). \end{aligned}$$

Only taking the off-diagonal element of the last matrix gives the local posterior covariance of  $(f(x), f(x'))$ :

$$\begin{aligned} \hat{C}^{(k)}(x, x') &= mC_\beta(x, x') \\ &- mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, x)^\top [C_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) + \frac{1}{m}\sigma^2 I]^{-1} mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, x'). \end{aligned}$$

Since the aggregated posterior is defined as the average of the local posteriors, the aggregated posterior covariance is equal to  $\hat{C}_{n,m}(x, x') = \frac{1}{m^2} \sum_{k=1}^m \hat{C}^{(k)}(x, x')$ . Now write  $\hat{C}^{(k)}(x, x') = mC_\beta(x, x') - \tilde{C}_\beta^{(k)}(x, x')$ , with

$$\begin{aligned} \tilde{C}_\beta^{(k)}(x, x') \\ = mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, x)^\top [mC_{\beta+p}(\mathbb{X}_n^{(k)}, (\mathbb{X}_n^{(k)})^\top) + \sigma^2 I]^{-1} mC_{\beta+\frac{p}{2}}(\mathbb{X}_n^{(k)}, x'). \end{aligned}$$

The local posterior mean is the minimizer of Equation (3.18), so it follows that  $\tilde{C}_\beta^{(k)}(\cdot, x')$  must be the function that satisfies

$$\frac{1}{m} \tilde{C}_\beta^{(k)}(\cdot, x') = \operatorname{argmin}_{g \in \mathcal{H}} \frac{1}{n} \left( \sum_{i=1}^n (C_{\beta+\frac{p}{2}}(X_i^{(k)}, x') - Kg(X_i^{(k)}))^2 + \frac{\sigma^2}{m} \|g\|_{\mathcal{H}}^2 \right). \quad (3.37)$$

The Fréchet derivative of this gives the score function

$$\hat{S}_{x',n}^{(k)}(g) = \frac{1}{n} \sum_{i=1}^n (KC_\beta(X_i^{(k)}, x') - Kg(X_i^{(k)})) KC_{\beta, X_i^{(k)}} - \frac{\sigma^2}{N} g,$$

where  $C_{\beta, x'}(\cdot) := C_\beta(\cdot, x')$ . Note that, since  $\tilde{C}_\beta^{(k)}(\cdot, x')$  can be defined as a minimizer, it satisfies  $\hat{S}_{x',n}^{(k)}(\frac{1}{m} \tilde{C}_\beta^{(k)}(\cdot, x')) = 0$ . The expected value of  $\hat{S}_{x',n}^{(k)}(g)$  is

$$\begin{aligned} S_{x',n}^{(k)}(g) &:= \mathbb{E}(\hat{S}_{x',n}^{(k)}(g)) \\ &= \int_{\mathcal{O}} (KC_{\beta, x'}(x) - Kg(x)) KC_{\beta, x} dx - \frac{\sigma^2}{N} g \\ &= \tilde{K}(KC_{\beta, x'}) - \tilde{K}Kg - \frac{\sigma^2}{N} g \\ &= S_{x',N}(g). \end{aligned}$$

Thus using the definition of  $\nu_j$ 's, we find

$$S_{x',N}(g) = \tilde{K}(KC_{\beta, x'}) - \tilde{K}\tilde{V}_N^{-1}Kg = \tilde{K}(KC_{\beta, x'} - \tilde{V}_N^{-1}Kg) \quad (3.38)$$

$$\Delta \tilde{C}_{\beta, x'}^{(k)} := \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} - \tilde{V}_N C_{\beta, x'} = -\tilde{R}_N(S_{x',N}(\frac{1}{m} \tilde{C}_{\beta, x'}^{(k)})) \quad (3.39)$$

$$\hat{S}_{x',n}^{(k)}(\tilde{V}_N C_{\beta, x'}) = \frac{1}{n} \sum_{i=1}^n \tilde{P}_N(KC_{\beta, x'})(X_i^{(k)}) KC_{\beta, X_i^{(k)}} - \frac{\sigma^2}{nm} \tilde{V}_N C_{\beta, x'}. \quad (3.40)$$

The first equation in this display also shows that  $\tilde{V}_N C_{\beta, x'}$  is a root of  $S_{x', N}$ . Using the roots of  $\hat{S}_{x', n}^{(k)}$  and  $S_{x', n}^{(k)}$ , we find

$$\begin{aligned}
 & \tilde{K} \tilde{V}_N^{-1} K^{-1} (\Delta \tilde{C}_{\beta, x'}^{(k)}) - \hat{S}_{x', n}^{(k)} (\tilde{V}_N (C_{\beta, x'})) \\
 &= -S_{x', n}^{(k)} \left( \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right) - \hat{S}_{x', n}^{(k)} (\tilde{V}_N (C_{\beta, x'})) \\
 &= (-S_{x', n}^{(k)} \left( \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right) + S_{x', n}^{(k)} (\tilde{V}_N C_{\beta, x'})) - (\hat{S}_{x', n}^{(k)} (\tilde{V}_N (C_{\beta, x'})) - \hat{S}_{x', n}^{(k)} \left( \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right)) \\
 &= \int_{\mathcal{O}} (K \left( \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right)(x) - K \tilde{V}_N C_{\beta, x'}(x)) K C_{\beta, x} dx \\
 &\quad - \frac{1}{n} \sum_{i=1}^n \left( K \left( \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right)(X_i^{(k)}) - K \tilde{V}_N C_{\beta, x'}(X_i^{(k)}) \right) K C_{\beta, X_i^{(k)}} \\
 &= \int_{\mathcal{O}} (K \Delta \tilde{C}_{\beta, x'}^{(k)})(x) K C_{\beta, x} dx - \frac{1}{n} \sum_{i=1}^n K \Delta \tilde{C}_{\beta, x'}^{(k)}(X_i^{(k)}) K C_{\beta, X_i^{(k)}}.
 \end{aligned} \tag{3.41}$$

Throughout the following theorems and lemmas, we assume that the eigenfunctions are uniformly bounded, and that  $\beta > \frac{d}{2}$ .

*Proof of Theorem 3.3.2.* Note that we can write

$$\begin{aligned}
 P_0(f_0 \in B_{r_{n, m, \tau}}(L)) &= P_0(\|f_0 - \hat{f}_N\|_{L^2} \leq L r_{n, m, \tau}) \\
 &\geq P_0(\|f_0 - \tilde{V}_N(f_0)\|_{L^2} + \|\tilde{V}_N(f_0) - \hat{f}_N\|_{L^2} \leq L r_{n, m, \tau}).
 \end{aligned}$$

From Lemma 3.10.2, we know that for  $N \rightarrow \infty$ ,

$$P_0 \left( r_{n, m, \tau}^2 \geq C_1 \frac{\sigma^2}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2} \right) \rightarrow 1$$

for some  $C_1 \geq 0$  small enough. In view of Lemma 3.10.6, the  $L^2$ -norm of  $\tilde{P}_N f_0$  is bounded by  $\|\tilde{P}_N f_0\|_{L^2}^2 \lesssim N^{-\frac{2\beta}{d+2\beta+p}} \|f_0\|_{L^2}^2$ . Lemma 3.10.7 with  $q = -\frac{1}{2}$ ,  $t = \frac{1}{d}(d + 2\alpha)$ ,  $u = \frac{1}{d}(d + 2\alpha + 2p)$  and  $v = 1$  shows that

$$\frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2} \asymp N^{-\frac{2\beta}{d+2\beta+2p}}, \tag{3.42}$$

hence

$$\|f_0 - \tilde{V}_N(f_0)\|_{L^2}^2 \leq \frac{C_2 \|f_0\|_{L^2}^2}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2}$$

for  $C_2 \geq 0$  a large enough constant. Due to Lemma 3.10.1 we can choose  $M$  large enough so that the probability of  $\|\Delta \hat{f}_{n,m}^\dagger\|_{L^2}^2 \leq M \frac{\sigma^2}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2}$  tends to a value which is at least  $1 - \tau$  uniformly in  $f_0 \in S^\beta(B)$ . Then we choose  $L > 2 \frac{M\sigma^2 + C_2}{C_1\sigma^2}$  so that  $P_0(f_0 \in B_{r_{n,m,\tau}}(L))$  has probability of at least  $1 - \tau$  for large enough  $N$ .  $\square$

**Lemma 3.10.1.** *There exists a  $C_1 \geq 0$  such that for any  $M \geq 1$ , we have*

$$\liminf_{n \rightarrow \infty} \inf_{f_0 \in S^\beta(B)} P_0 \left( \|\Delta \hat{f}_{n,m}^\dagger\|_{L^2}^2 \leq M \frac{\sigma^2}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2} \right) \geq 1 - \frac{C_1}{M}.$$

*Proof.* Using Markov's inequality and Lemma 3.7.1 with  $\gamma = 0$  we can bound the probability by

$$\begin{aligned} P_0 \left( \|\Delta \hat{f}_{n,m}^\dagger\|_{L^2}^2 \geq M \frac{\sigma^2}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2} \right) &\leq \frac{\mathbb{E}_0(\|\Delta \hat{f}_{n,m}^\dagger\|_{L^2}^2)}{M \frac{\sigma^2}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2}} \\ &\lesssim \frac{\frac{1}{N} \sum_{j=1}^\infty \nu_j^2 \kappa_j^{-2} + o(N^{-\frac{2\beta}{d+2\beta+2p}})}{\frac{M}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2}} \\ &\lesssim \frac{1}{M} + \frac{o(N^{1-\frac{2\beta}{d+2\beta+2p}})}{\sum_{j=1}^\infty \nu_j \kappa_j^{-2}}. \end{aligned}$$

From Equation (3.42) we see that

$$\frac{o(N^{1-\frac{2\beta}{d+2\beta+2p}})}{\sum_{j=1}^\infty \nu_j \kappa_j^{-2}} \rightarrow 0, \quad (3.43)$$

which finishes the proof.  $\square$

**Lemma 3.10.2.** *Let  $\zeta < \frac{1}{2}$ ,  $\alpha = \beta$  and assume  $\beta > 9(d+p)$ . Then we have*

$$\mathbb{P}_0 \left( r_{n,m,\tau}^2 \geq \frac{1}{2C_\phi^2} \frac{\sigma^2}{N} \sum_{j=1}^\infty \nu_j \kappa_j^{-2} \right) \rightarrow 1.$$

*Proof.* The credible set is defined by the radius  $r_{n,m,\tau}$  that satisfies  $\mathbb{P}(\|W_{n,m}\|_{L^2}^2 \leq r_{n,m,\tau}^2 \mid \mathbb{X}_N) = 1 - \tau$ . Here  $W_{n,m}$  is a Gaussian process with zero mean and covariance kernel equal to the covariance kernel of the aggregated posterior. First,



we give a lower bound on the radius of the credible set  $r_{n,m,\tau}$  by Chebyshev's inequality, under the assumption that  $\mathbb{E}\|W_{n,m}\|_{L^2}^2 \geq r_{n,m,\tau}^2$  by

$$\begin{aligned} 1 - \tau &= \mathbb{P}(\|W_{n,m}\|_{L^2}^2 - \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) \leq r_{n,m,\tau}^2 - \mathbb{E}\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) \\ &\leq \mathbb{P}(\|W_{n,m}\|_{L^2}^2 - \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) \geq \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) - r_{n,m,\tau}^2 | \mathbb{X}_N) \\ &\leq \frac{\text{Var}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N)}{(\mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) - r_{n,m,\tau}^2)^2}. \end{aligned}$$

This inequality is equivalent to

$$r_{n,m,\tau}^2 \geq \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) - (1 - \tau)^{\frac{1}{2}} \text{Var}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N)^{\frac{1}{2}}. \quad (3.44)$$

It follows that this last inequality holds for arbitrary values of  $\mathbb{E}\|W_{n,m}\|_{L^2}^2$  and  $r_{n,m,\tau}^2$ . In view of Fubini's theorem, the second moment of  $\|W_{n,m}\|_{L^2}$  conditionally on  $\mathbb{X}_N$  is equal to

$$\begin{aligned} \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) &= \mathbb{E}_{\Pi}(\|f - \hat{f}_{n,m}^{\dagger}\|_{L^2}^2 | \mathbb{D}_N) \\ &= \frac{1}{m^2} \sum_{k=1}^m \int_{\mathcal{O}} \text{Var}_{\Pi}^{(k)}[f(x) | \mathbb{D}_n^{(k)}] dx. \end{aligned}$$

We can lower bound this for  $N \rightarrow \infty$  by  $\frac{\sigma^2}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2}$  due to Lemma 3.9.2. From Equation (7.23) in (Hadji, Hesselink and Szabó 2022) we derive the equality

$$\mathbb{E}(\|W_{n,m}\|_{L^2}^4 | \mathbb{X}_N) = \mathbb{E}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N)^2 + 2 \int_{\mathcal{O}} \|\hat{C}_{n,m}(x, \cdot)\|_{L^2}^2 dx.$$

Then following the same the expected posterior variance can be bounded from above by

$$\begin{aligned} \mathbb{E}_0 \text{Var}(\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N) &= \mathbb{E}_0(\mathbb{E}[\|W_{n,m}\|_{L^2}^4 | \mathbb{X}_N] - \mathbb{E}_0[\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N]^2) \\ &= \mathbb{E}_0\left(\mathbb{E}[\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N]^2 + 2 \int_{\mathcal{O}} \|\hat{C}_{n,m}(x, \cdot)\|_{L^2}^2 dx \right. \\ &\quad \left. - \mathbb{E}_0[\|W_{n,m}\|_{L^2}^2 | \mathbb{X}_N]^2\right) \\ &= 2 \int_{\mathcal{O}} \mathbb{E}_0 \|\hat{C}_{n,m}(x, \cdot)\|_{L^2}^2 dx \\ &= 2 \int_{\mathcal{O}} \mathbb{E}_0 \left\| \frac{1}{m^2} \sum_{k=1}^m \hat{C}^{(k)}(x, \cdot) \right\|_{L^2}^2 dx \\ &\leq 2m^{-3} \sum_{k=1}^m \int_{\mathcal{O}} \mathbb{E}_0 \|\hat{C}^{(k)}(x, \cdot)\|_{L^2}^2 dx. \end{aligned}$$

Lemma 3.10.3 provides an upper bound for the term inside the sum, showing that

$$\begin{aligned}
 \mathbb{E}_0 \text{Var}(\|W_{n,m}\|_{L^2}^2 \mid \mathbb{X}_N) &\lesssim \mathbb{E}_X \|\tilde{P}_N C_{\beta,X}\|_{S^\beta}^2 + o(N^{-\frac{4\beta}{d+2\beta+2p}}) \\
 &= \mathbb{E}_X \sum_{j=1}^{\infty} (1 - \nu_j)^2 \kappa_j^{\frac{2}{p}(d+2\beta)} \phi_j(X)^2 + o(N^{-\frac{4\beta}{d+2\beta+2p}}) \\
 &\leq \frac{\sigma^4}{N^2} \mathbb{E}_X \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \phi_j(X)^2 + o(N^{-\frac{4\beta}{d+2\beta+2p}}) \\
 &\asymp \frac{1}{N^2} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} + o(N^{-\frac{4\beta}{d+2\beta+2p}}).
 \end{aligned}$$

We use Markov's inequality and the bound in the previous display to find

$$\begin{aligned}
 P_0 \left( \text{Var}(\|W_{n,m}\|_{L^2}^2 \mid \mathbb{X}_N) \geq t^2 \frac{\sigma^4}{N^2} \left( \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2} \right)^2 \right) \\
 \lesssim t^{-2} \left( \frac{\sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2}}{(\sum_{j=1}^{\infty} \nu_j \kappa_j^{-2})^2} + \frac{o(N^2 N^{-\frac{4\beta}{d+2\beta+2p}})}{(\sum_{j=1}^{\infty} \nu_j \kappa_j^{-2})^2} \right).
 \end{aligned}$$

From Lemma 3.10.6, we see that  $\sum_{j=1}^{\infty} \nu_j \kappa_j^{-2} \asymp N^{\frac{d+2p}{d+2\beta+2p}}$ , so both first term go to zero as  $N \rightarrow \infty$ . Combining this, for a large enough choice of  $t$ , with Lemma 3.9.2, shows

$$P_0 \left( \mathbb{E}(\|W_{n,m}\|_{L^2}^2 \mid \mathbb{X}_N) - (1 - \tau)^{\frac{1}{2}} \text{Var}(\|W_{n,m}\|_{L^2}^2 \mid \mathbb{X}_N)^{\frac{1}{2}} \gtrsim \frac{1}{N} \sum_{j=1}^{\infty} \nu_j \kappa_j^{-2} \right) \rightarrow 1,$$

finishing the proof when combined with Equation (3.44).  $\square$

**Lemma 3.10.3.** *We have*

$$\mathbb{E}_0 \|\hat{C}^{(k)}(x, \cdot)\|_{L^2}^2 \lesssim m^2 \mathbb{E}_X \|\tilde{P}_N C_{\beta,X}\|_{S^\beta}^2 + o(m^2 N^{-\frac{4\beta}{d+2\beta+2p}}).$$

*Proof.* We write

$$\begin{aligned}
 \|\hat{C}^{(k)}(x, \cdot)\|_{L^2}^2 &= \|mC_\beta(x, \cdot) - \tilde{C}_\beta^{(k)}(x, \cdot)\|_{L^2}^2 \\
 &= \|mC_\beta(x, \cdot) - \tilde{V}_N(mC_\beta(x, \cdot)) + \tilde{V}_N(mC_\beta(x, \cdot)) - \tilde{C}_\beta^{(k)}(x, \cdot)\|_{L^2}^2 \\
 &\leq 2m^2 \|\tilde{P}_N C_\beta(x, \cdot)\|_{L^2}^2 + 2m^2 \|\Delta \tilde{C}_\beta^{(k)}\|_{L^2}^2.
 \end{aligned}$$

The inequality  $(a + b)^2 \leq 2a^2 + 2b^2$  shows that

$$\|\Delta\tilde{C}_{\beta,x}^{(k)}\|_{L^2}^2 \leq 2\|\Delta\tilde{C}_{\beta,x}^{(k)} - \tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}))\|_{L^2}^2 + 2\|\tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}))\|_{L^2}^2.$$

In view of Lemma 3.10.4 this results in

$$\mathbb{E}_0\|\Delta\tilde{C}_{\beta,x}^{(k)} - \tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}))\|_{L^2}^2 \leq o\left(\mathbb{E}_0\|\Delta\tilde{C}_{\beta,x}^{(k)}\|_{L^2}^2\right) + o\left(N^{-\frac{4\beta}{d+2\beta+2p}}\right),$$

which combined with the earlier inequality results in

$$\|\Delta\tilde{C}_{\beta,x}^{(k)}\|_{L^2}^2 \lesssim \|\tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}))\|_{L^2}^2 + o\left(N^{-\frac{4\beta}{d+2\beta+2p}}\right).$$

Using Equation (3.40), we find

$$\begin{aligned} & \mathbb{E}_0\left\|\tilde{R}_N\hat{S}_{x,n}^{(k)}\tilde{V}_N C_{\beta,x}\right\|_{L^2}^2 \\ &= \mathbb{E}_0\left\|\tilde{R}_N\left(\hat{S}_{x,n}^{(k)}\tilde{V}_N C_{\beta,x} - S_{x,n}\tilde{V}_N C_{\beta,x}\right)\right\|_{L^2}^2 \\ &= \mathbb{E}_0\left\|\tilde{R}_N\left(\frac{1}{n}\sum_{i=1}^n K\tilde{P}_N(C_{\beta,x})(X_i^{(k)})KC_{\beta,X_i^{(k)}} - \mathbb{E}_X K\tilde{P}_N(C_{\beta,x})(X)KC_{\beta,X}\right)\right\|_{L^2}^2 \\ &\lesssim \frac{\|\tilde{P}_N C_{\beta,x}\|_{S^\beta}^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} = o(\|\tilde{P}_N C_{\beta,x}\|_{S^\beta}^2). \end{aligned}$$

using Lemma 3.8.1 with  $g = \tilde{P}_N(\tilde{C}_{\beta,x})$ . □

**Lemma 3.10.4.** *Let  $\alpha = \beta$  with  $\beta > 9(d + p)$ . Then we have*

$$\mathbb{E}_0\|\Delta\tilde{C}_{\beta,x}^{(k)} - \tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}))\|_{L^2}^2 \leq o(\mathbb{E}_0\|\Delta\tilde{C}_{\beta,x}^{(k)}\|_{L^2}^2) + o\left(N^{-\frac{4\beta}{d+2\beta+2p}}\right).$$

*Proof.* Note that by Equation (3.41) and Lemma 3.7.3 with  $\hat{g}^{(k)} = \Delta\tilde{C}_{\beta,x}^{(k)}$  and

$\gamma = 0$ , for  $C_0$  arbitrarily large,

$$\begin{aligned}
 & \mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x}^{(k)} - \tilde{R}_N(\hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x})) \right\|_{L^2}^2 \\
 &= \mathbb{E}_0 \left\| \tilde{R}_N \left( F \tilde{V}_N^{-1} K(\Delta \tilde{C}_{\beta,x}^{(k)}) - \hat{S}_{x,n}^{(k)}(\tilde{V}_N C_{\beta,x}) \right) \right\|_{L^2}^2 \\
 &= \mathbb{E}_0 \left\| \tilde{R}_N \left( \mathbb{E}_X(K \Delta \tilde{C}_{\beta,x}^{(k)})(X) K C_{\beta,x} - \frac{1}{n} \sum_{i=1}^n K \Delta \tilde{C}_{\beta,x}^{(k)}(X_i^{(k)}) K C_{\beta,X_i^{(k)}} \right) \right\|_{L^2}^2 \\
 &\lesssim \frac{\log N \sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2}{n} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x}^{(k)} \right\|_{L^2}^2 \\
 &+ \mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x,\mathcal{I}^c}^{(k)} \right\|_{\mathcal{H}}^2 \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_\ell^{\frac{1}{p}(d+2\beta+2p)} + N^{-C_0} \\
 &=: T_1 + T_2 + N^{-C_0}.
 \end{aligned}$$

We take  $\mathcal{I}$  to be the largest  $\mathcal{I} \in \mathbb{N}$  such

$$\sum_{\ell=1}^{\mathcal{I}} \kappa_\ell^2 \leq o \left( \frac{n}{\log N} \left( \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \right)^{-1} \right),$$

so that  $T_1$  is  $o(\mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x}^{(k)} \right\|_{L^2}^2)$ . For  $T_2$  we use Lemma 3.10.5 and compare it to Equations (3.28) and (3.29) with  $\gamma = 0$ . We see that

$$\begin{aligned}
 T_2 &\lesssim \left( \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \right)^2 \sum_{\ell=\mathcal{I}+1}^{\infty} \kappa_\ell^{\frac{1}{p}(d+2\beta+2p)} \\
 &= \frac{1}{N} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} o \left( N^{-\frac{2\beta}{d+2\beta+2p}} \right) \\
 &= o \left( N^{-\frac{4\beta}{d+2\beta+2p}} \right),
 \end{aligned}$$

where the last inequality follows from Equation (3.27). □

**Lemma 3.10.5.** *We have*

$$\mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x,\mathcal{I}^c}^{(k)} \right\|_{\mathcal{H}}^2 \lesssim \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2}.$$

Furthermore, if  $\beta$  is such that  $p + \beta \geq \frac{d}{2}$ , then for any measurable event  $A$ , we have

$$\mathbb{E}_0 \left\| \Delta \tilde{C}_{\beta,x}^{(k)} \right\|_{L^2}^2 1_A \lesssim \mathbb{P}(A)^{\frac{1}{2}} N^2.$$

*Proof.* We have

$$\mathbb{E}_0 \|\Delta \tilde{C}_{\beta,x}^{(k)}\|_{\mathcal{H}}^2 \leq 2\mathbb{E}_0 \|\frac{1}{m} \tilde{C}_{\beta,x}^{(k)}\|_{\mathcal{H}}^2 + 2\|\tilde{V}_N C_{\beta,x}\|_{\mathcal{H}}^2.$$

The second term is equal to

$$\begin{aligned} \|\tilde{V}_N C_{\beta,x}\|_{\mathcal{H}}^2 &= \sum_{j=1}^{\infty} \kappa_j^{-\frac{1}{p}(d+2\beta)} \nu_j^{\frac{2}{p}(d+2\beta)} \phi_j(x)^2 \\ &= \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\beta)} \nu_j^2 \phi_j(x)^2 \\ &\lesssim \sum_{j=1}^{\infty} \nu_j^2. \end{aligned}$$

Because  $\frac{1}{m} \tilde{C}_{\beta,x}^{(k)}$  is a minimizer of Equation (3.37), we can write

$$\begin{aligned} &\mathbb{E}_0 \sigma^2 \|\frac{1}{m} \tilde{C}_{\beta,x}^{(k)}\|_{\mathcal{H}}^2 \tag{3.45} \\ &\leq m \mathbb{E}_0 \sum_{i=1}^n (K C_{\beta,x}(X_i) - K(\frac{1}{m} \tilde{C}_{\beta,x}^{(k)}(X_i)))^2 + \sigma^2 \|\frac{1}{m} \tilde{C}_{\beta,x}^{(k)}\|_{\mathcal{H}}^2 \\ &\leq m \mathbb{E}_0 \sum_{i=1}^n (K C_{\beta,x}(X_i) - K \tilde{V}_N C_{\beta,x}(X_i))^2 + \sigma^2 \|\tilde{V}_N C_{\beta,x}\|_{\mathcal{H}}^2 \\ &= m \sum_{i=1}^n \mathbb{E}_X (K \tilde{P}_N C_{\beta,x})(X)^2 + \sigma^2 \|\tilde{V}_N C_{\beta,x}\|_{\mathcal{H}}^2. \end{aligned}$$

The expectation is equal to

$$\begin{aligned} \mathbb{E}_X (\tilde{P}_N K C_{\beta,x}(X))^2 &= \mathbb{E}_0 \left( \sum_{j=1}^{\infty} (1 - \nu_j) \kappa_j^{\frac{1}{p}(d+2\beta+p)} \phi_j(x) \phi_j(X) \right)^2 \\ &= \frac{\sigma^4}{N^2} \sum_{j=1}^{\infty} \nu_j^2 \kappa_j^{-2} \phi_j(x)^2. \end{aligned}$$

The second part of the proof follows similar lines to the previous proof and Lemma 3.8.3 with  $\Delta \hat{f}_n^{(k)}$  replaced by  $\Delta \tilde{C}_{\beta,x}^{(k)}$ . We can write

$$\mathbb{E}_0 \|\Delta \tilde{C}_{\beta,x}^{(k)}\|_{L^2 1_A}^2 \leq 2\mathbb{E}_0 \|\frac{1}{m} \tilde{C}_{\beta,x'}^{(k)}\|_{\mathcal{H}}^2 1_A + 2\mathbb{E}_0 \|\tilde{V}_N C_{\beta,x'}\|_{L^2 1_A}^2.$$

The first term is bounded due to Equation (3.45) by

$$\mathbb{E}_0 \sigma^2 \left\| \frac{1}{m} \tilde{C}_{\beta, x'}^{(k)} \right\|_{\mathcal{H}}^2 1_A \leq m \sum_{i=1}^n \mathbb{E}_X (K \tilde{P}_N C_{\beta, x})(X)^2 1_A + \sigma^2 \|\tilde{V}_N C_{\beta, x}\|_{\mathcal{H}}^2 \mathbb{P}(A).$$

We bound the term inside the sum by

$$\begin{aligned} & \mathbb{E}_0 (K \tilde{P}_N C_{\beta, x})(X)^2 1_A \\ & \leq \left( \mathbb{E}_0 \left( (K \tilde{P}_N C_{\beta, x})(X) \right)^4 \mathbb{E}_0 1_A \right)^{\frac{1}{2}} \\ & = \left( \mathbb{P}(A) \sum_{\substack{j_1, j_2, \\ j_3, j_4=1}}^{\infty} \prod_{\ell=1}^4 (1 - \nu_{j_\ell}) \kappa_{j_\ell}^{\frac{1}{p}(d+2\beta+p)} \phi_{j_\ell}(x) \mathbb{E}_0 \prod_{\ell=1}^4 \phi_{j_\ell}(X) \right)^{\frac{1}{2}} \\ & \leq \mathbb{P}(A)^{\frac{1}{2}} C_\phi^2 \left( \sum_{j=1}^{\infty} (1 - \nu_j) \kappa_j^{\frac{1}{p}(d+2\beta+p)} \phi_j(x) \right)^2 \\ & \lesssim \mathbb{P}(A)^{\frac{1}{2}} \left( \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\beta+2p)} \right)^2. \end{aligned}$$

Finally, we have

$$\begin{aligned} \|\tilde{V}_N C_{\beta, x}\|_{\mathcal{H}}^2 &= \sum_{j=1}^{\infty} \kappa_j^{-\frac{1}{p}(d+2\beta)} \nu_j^2 \kappa_j^{\frac{2}{p}(d+2\beta)} \phi_j(x)^2 \\ &= \sum_{j=1}^{\infty} \kappa_j^{\frac{1}{p}(d+2\beta)} \frac{\kappa_j^{\frac{2}{p}(d+2\beta+2p)} N^2}{(\sigma^2 + N \kappa_j^{\frac{1}{p}(d+2\beta+2p)})^2} \\ &= N^2 \sum_{j=1}^{\infty} \frac{\kappa_j^{\frac{1}{p}(3d+6\beta+4p)}}{(\sigma^2 + N \kappa_j^{\frac{1}{p}(d+2\beta+2p)})^2} \\ &\lesssim N^2, \end{aligned}$$

since the last sum is finite and decreasing in  $N$ . □

The following lemmas and their proof from (Knapik, van der Vaart and van Zanten 2011) can be found in Lemma 8.1 and Lemma 8.2, respectively.

**Lemma 3.10.6.** *Let  $q \geq 0, t \geq -2q, u > 0$  and  $v \geq 0$  be given. Then for  $N \rightarrow \infty$ ,*

$$\sup_{\|\zeta\|_q \leq 1} \sum_{j=1}^{\infty} \frac{\zeta_j^2 j^{-t}}{(1 + Nj^{-u})^v} \asymp N^{-\left(\frac{t+2q}{u}\right) \wedge v}$$

*holds.*

**Lemma 3.10.7.** *Let  $t, v \geq 0, u > 0$  be given. Let  $(\xi_j)_{j \in \mathbb{N}}$  be such that  $|\xi_j| = j^{-q-\frac{1}{2}} \mathcal{S}(j)$  for some  $q > -\frac{t}{2}$  and a slowly varying function  $\mathcal{S} : (0, \infty) \rightarrow (0, \infty)$ . Then, for  $N \rightarrow \infty$ ,*

$$\sum_{j=1}^{\infty} \frac{\xi_j^2 j^{-t}}{(1 + Nj^{-u})^v} \asymp \begin{cases} N^{-\frac{t+2q}{u}} \mathcal{S}^2(N^{\frac{1}{u}}), & \text{if } (t+2q)/u < v, \\ N^{-v} \sum_{j \leq N^{1/u}} \mathcal{S}^2(j)/j, & \text{if } (t+2q)/u = v, \\ N^{-v}, & \text{if } (t+2q)/u > v. \end{cases}$$

*The sum is asymptotically equivalent to the same sum over  $1 \leq i \leq N^{\frac{1}{u}}$  if and only if  $\frac{t+2q}{u} \geq v$ .*

# Chapter 4

## Misspecified Bernstein-von Mises theorem

### 4.1 Introduction

Hierarchical models are frequently used in various fields, for instance in astronomy (weak lensing (Alsing et al. 2016; Sellentin, Heymans and Harnois-Déraps 2018), cosmic microwave background (CMB) temperature maps (Eriksen et al. 2004), gravitational waves (Cornish and Littenberg 2015)), in statistical physics (the Dyson hierarchical model (Monthus and Garel 2011)), in environmental sciences (predicting the spread of ecological processes (Wikle 2003)) and in medicine (spatial modelling of fMRI data (Bowman et al. 2008)). Hierarchical models specify a data-generating process in multiple layers, which yields great modelling flexibility. The connections between the layers are often non-linear, governed by some (system of) partial differential equations. This complex structure imposes analytical and computational challenges.

The goal is to make inference on aspects of the hierarchical process. The direct application of exact methods for inference is typically computationally overly expensive or even infeasible. Therefore, in practice often a surrogate, approximate likelihood is considered instead of the likelihood of the multilayer data-generating structure. For computational convenience and easy interpretation, misspecified Gaussian likelihoods are popular. However, using the simplified likelihood results in information loss and may lead to inaccurate error bands. Although in practice one is aware of the inaccuracy of the misspecified hierarchical procedures, these are typically used without rigorous theoretical guarantees or understanding of their limitations. This sometimes results in



contradictory results.

This chapter was motivated by a concrete example from astronomy. Cosmic microwave background radiation (a snapshot of the universe approximately  $3 \times 10^5$  years after the big bang) is modelled as a draw from a Gaussian random field (although the Gaussianity is debated (Marinucci 2004)). In our current time, we observe non-linear transformations of this field, which can be described by systems of complex partial differential equations, corrupted with noise. Due to the non-linearity, the observed data are non-Gaussian. The goal is to recover important, parametric aspects of the non-linear operators (and their surrogates). This includes the proportion of dark energy and dark matter, and the Hubble constant (the relative expansion rate of the universe). The article (Ezquiaga and Zumalacárregui 2018) reviews tools for understanding the nature of dark energy using gravitational waves. To simplify the model and the computations, the complex hierarchical model is replaced by a Gaussian approximation. Figure 8 of (Ezquiaga and Zumalacárregui 2018) collects the outcome of a collection of methods for estimating the Hubble constant. One can observe that the resulting error bars of the proposed estimators do not intersect, providing contradictory conclusions about the relative expansion rate of the universe. Understanding the theoretical properties of such, misspecified, hierarchical approaches is important to know whether the non-overlapping confidence intervals are due to using different models or due to statistical error. Here we consider such a result in a misspecified Gaussian framework.

In this chapter, we consider Bayesian methods for inference on parametric aspects of the complex, hierarchical data-generating process. Bayesian methods have gained popularity in many fields, due to their built-in uncertainty quantification and their ability to provide a natural way of incorporating preliminary knowledge into the observational model. In our work, we focus on the theoretical, asymptotic properties of the posterior distribution on the parameter of interest and investigate whether it can achieve optimal recovery and reliable uncertainty quantification. In regular, parametric models this is done typically by deriving a Bernstein-von Mises result, which assures that the posterior distribution is asymptotically Gaussian, centred at an efficient estimator and with variance equal to the Cramér-Rao bound.

We consider a general misspecified model indexed by finite-dimensional parameters and derive a Bernstein-von Mises theorem, giving a Gaussian approximation of the posterior distribution. Bernstein-von Mises type results have been derived in the literature for various models, starting from well-specified regular parametric models (Doob 1949; Le Cam 2012; van der Vaart 1998), and extended to misspecified models; see (Kleijn and van der Vaart 2012) and the

recent review article (Bochkina 2022). Extensions to semi-parametric (Castillo and Rousseau 2015) and non-parametric models for weak norms were considered in (Leahu 2011; Castillo and Nickl 2014), while recently a more accurate, skewed version of the theorem was derived in (Durante, Pozza and Szabó 2023). However, none of these results consider hierarchical models as examples, and the derived misspecified results depend on conditions that are not met by our highly non-linear examples. Thus, first, we derive a new, modified version of the classical misspecified theorem. Next, we apply this to a collection of hierarchical models, starting with the square integral operator as a toy example and next turning to partial differential equation (PDE) constrained inverse problems. We first investigate the time-independent Schrödinger equation, which allows tractable computations and is often used as a platform for investigating more complex PDE-constrained inverse problems. Then we derive asymptotic guarantees for general parabolic PDEs with parametric potential and boundary constraints. We consider observations at different time points, mimicking typical non-i.i.d. observational structures from physics and astronomy.

The chapter is organized as follows. In Section 4.2 we introduce the hierarchical observational model, its Gaussian surrogate and the misspecified Bayesian framework. In Section 4.3 we provide our general misspecified Bernstein-von Mises theorem and in Section 4.3.2 we apply it to hierarchical models. Our main contribution is to cover several interesting, PDE-constrained inverse problems, and is presented in Section 4.4. In Section 4.5 we investigate the numerical behaviour of our limiting misspecified posterior in contrast to the well-specified case and the accuracy of the Gaussian approximation. The proofs of the general results, examples, and additional technical lemmas are deferred to Sections 4.6, 4.7, 4.8, and 4.9.

## 4.2 Description of the problem

We consider non-linear hierarchical models of the form

$$\begin{aligned} f_i &\stackrel{\text{i.i.d.}}{\sim} G, & i = 1, \dots, N, \\ X_i \mid f_i, \theta_0 &\stackrel{\text{ind}}{\sim} N_p(T_{\theta_0}^{(i)} f_i, \Lambda^{(i)}), & i = 1, \dots, N. \end{aligned} \quad (4.1)$$

Here  $G$  is a distribution on some Hilbert space  $(\mathcal{F}, \langle \cdot, \cdot \rangle_{\mathcal{F}})$ ,  $T_{\theta}^{(i)} : \mathcal{F} \mapsto \mathbb{R}^p$  are (possibly) non-linear operators, indexed by an unknown finite-dimensional parameter  $\theta \in \mathbb{R}^d$  of interest, and  $\Lambda^{(i)} \in \mathbb{R}^{p \times p}$  are known, positive-definite covariance matrices. We denote the distribution of  $X_i$  by  $P_{\theta, i}$ , for  $i = 1, \dots, N$ ,

assume that this has a density  $p_{\theta,i}$  relative to some measure  $\lambda$ , and denote the expectation and covariance matrix with respect to this distribution by, for  $i = 1, \dots, N$ ,

$$\mu_{\theta}^{(i)} := \mathbb{E}_{\theta,i} X_i, \quad \Sigma_{\theta}^{(i)} := \text{Cov}_{\theta,i}(X_i). \quad (4.2)$$

We aim at making inference on the unknown model parameter  $\theta_0 \in \Theta \subset \mathbb{R}^d$  of interest. In case the operators  $T_{\theta_0}^{(i)}$  are the same for all  $i = 1, \dots, N$ , we arrive at an i.i.d. model for  $(X_i)_{i=1,\dots,N}$ .

The first layer in (4.1) models a random non-parametric functional parameter  $f_i$ . The second layer describes the actual observations, which depend on some (possibly) non-linear transformation of the functional parameter, and are corrupted with Gaussian noise. In Section 4.4 we consider several specific examples of this form, including elliptic and parabolic PDE-constrained inverse problems. The goal is to recover the unknown parameter  $\theta_0$  of the hierarchical model (4.1). In applications the operators  $T_{\theta}^{(i)}$  can be complex (see the examples in the introduction), but the main goal is typically the same: to recover certain parametric aspects of the model/operator.

In practice simplified, approximate models are used to speed up computations and increase the interpretability of the model. In particular, it is common to replace the otherwise complex likelihood with a Gaussian approximation with mean and covariance matching the corresponding quantities in the original model. In our analysis, we consider the misspecified model  $Q_{\theta,i} = N_p(\mu_{\theta}^{(i)}, \Sigma_{\theta}^{(i)})$  for  $X_i$ , for  $i = 1, \dots, N$ , with corresponding likelihood  $q_{\theta,i}$ . Using a misspecified model naturally results in information loss and requires care with uncertainty quantification. We investigate the accuracy of inference on  $\theta_0$  using this Gaussian, misspecified model, and quantify the amount of information loss and the uncertainty of the procedure.

In our analysis we consider a Bayesian approach and investigate its frequentist properties. We endow the parameter  $\theta \in \mathbb{R}^d$  in the misspecified likelihood  $q_{\theta,i}$  with a prior distribution  $\Pi$  with a density  $\pi$  satisfying standard regularity assumptions, to be specified later. We then investigate the asymptotic behaviour of the misspecified posterior distribution for  $\theta$ , i.e. the random measure

$$\Pi(\theta \in A | X_1, \dots, X_N) = \frac{\int_A \prod_{i=1}^N q_{\theta,i}(X_i) \pi(\theta) d\theta}{\int_{\Theta} \prod_{i=1}^N q_{\theta,i}(X_i) \pi(\theta) d\theta}, \quad (4.3)$$

for Borel sets  $A \subset \Theta$ . We derive a Bernstein-von Mises type result in this misspecified model, which shows that asymptotically this random measure

converges to a Gaussian distribution. Under mild assumptions this contracts around the parameter  $\theta_N^* \in \Theta$  minimizing the Kullback-Leibler divergence between the misspecified Gaussian class and the original model, i.e.

$$\theta_N^* = \arg \min_{\theta \in \Theta} P_{\theta_0}^{(N)} \left( \sum_{i=1}^N \log \frac{p_{\theta_0,i}(X_i)}{q_{\theta,i}(X_i)} \right).$$

We show that in our setup  $\theta_N^*$  coincides with the parameter of interest  $\theta_0$ , and derive the covariance matrix of the limiting Gaussian distribution. We investigate how this differs from the well-specified posterior distribution, both theoretically and numerically for various synthetic data sets.

### 4.3 Misspecified Bernstein-von Mises Theorem

The asymptotic behaviour of the posterior distribution in regular parametric models is characterized by the celebrated Bernstein-von Mises theorem. However, in practice often simplified models are considered, to speed up computations and improve the interpretation. In such cases, we talk about misspecified models. As an example, in the previous section we have introduced the complex hierarchical model (4.1), and we use a surrogate Gaussian model  $q_{\theta,i}$ . In this section, we first derive a Bernstein-von Mises theorem for misspecified models that goes beyond our hierarchical model (4.1) and next apply it to the latter special case.

#### 4.3.1 A general Bernstein-von Mises theorem

Misspecified Bernstein-von Mises results were derived in the literature before, see for instance (Kleijn and van der Vaart 2012), but under conditions too strong for our situation. In particular, to obtain a  $\sqrt{N}$ -contraction rate, the paper (Kleijn and van der Vaart 2012) assumes that the log misspecified likelihood is locally Lipschitz with a majorizing function that has exponential moments, which is violated for examples of our model (4.1), as demonstrated in Remark 4.1 in Section 4.4 below. In this section, we present a theorem that is appropriate for the model (4.1), and applies to hierarchical observational models with a selection of non-linear operators  $T_{\theta,i}$ .

The conditions of the theorem are essentially the classical ones given in Theorem 7.1 in (Lehmann 1983) (given there for the well-specified case, i.i.d. data and a one-dimensional parameter), but are at a high enough level of abstraction that they apply to general models, also beyond independent observations.

Let  $l^{(N)}(\theta)$  and  $\nabla l^{(N)}(\theta)$  denote a possibly misspecified log-likelihood and its gradient, of a model for an observation  $X_N$ . Given  $\theta_N^* \in \mathbb{R}^d$  and symmetric nonnegative-definite  $d \times d$  matrices  $V_{\theta_N^*, N}$ , define, for  $\theta \in \Theta$ , symmetric matrices  $R_N(\theta)$  through

$$\begin{aligned} l^{(N)}(\theta) &= l^{(N)}(\theta_N^*) + \nabla l^{(N)}(\theta_N^*)^\top (\theta - \theta_N^*) \\ &\quad - \frac{1}{2} N (\theta - \theta_N^*)^\top [V_{\theta_N^*, N} + N^{-1} R_N(\theta)] (\theta - \theta_N^*). \end{aligned} \quad (4.4)$$

We may think of  $\theta_N^*$  as the vector minimizing  $\theta \mapsto P_0^{(N)} l^{(N)}(\theta)$  for a given true distribution  $P_0^{(N)}$  of the observation, and  $V_{\theta_N^*, N}$  as the Hessian of the map  $\theta \mapsto -N^{-1} P_0^{(N)} l^{(N)}(\theta)$  at the point  $\theta = \theta_N^*$ , but this is not assumed in the following. The following assumptions are imposed.

**Assumption 4.1.**

**4.1.1** The sequence  $\theta_N^*$  tends to an interior point  $\theta^*$  of  $\Theta \subset \mathbb{R}^d$ .

**4.1.2**  $\|V_{\theta_N^*, N} - V_{\theta^*}\| \rightarrow 0$  for a positive-definite  $d \times d$  matrix  $V_{\theta^*}$ .

**4.1.3**  $\frac{1}{\sqrt{N}} \nabla l^{(N)}(\theta_N^*) = O_P(1)$  in  $P_0^{(N)}$ -probability as  $n \rightarrow \infty$ .

**4.1.4** For any  $\delta > 0$ , there exists an  $\epsilon > 0$  such that

$$\lim_{N \rightarrow \infty} P_0^{(N)} \left( \sup_{\theta: \|\theta - \theta_N^*\| \geq \delta} \frac{1}{N} (l^{(N)}(\theta) - l^{(N)}(\theta_N^*)) \leq -\epsilon \right) = 1.$$

**4.1.5** Given any  $\epsilon > 0$ , there exists a  $\delta > 0$  such that

$$\lim_{N \rightarrow \infty} P_0^{(N)} \left( \sup_{\theta: \|\theta - \theta_N^*\| \leq \delta} \frac{1}{N} \|R_N(\theta)\| \geq \epsilon \right) = 0.$$

**4.1.6** The prior has a density  $\pi$  that is continuous and positive at  $\theta = \theta^*$ .

**4.1.7** The prior has a finite first moment:  $\int_{\Theta} \|\theta\| \pi(\theta) d\theta < \infty$ .

**Proposition 4.3.1.** Let  $t \mapsto \pi_N(t|X_N)$  be the density of  $\sqrt{N}(\vartheta - T_N)$  if  $\vartheta$  follows the (misspecified) posterior distribution with density proportional to  $\theta \mapsto \exp(l^{(N)}(\theta)) \pi(\theta)$  and let

$$T_N = \theta_N^* + \frac{1}{N} V_{\theta_N^*, N}^{-1} \nabla l^{(N)}(\theta_N^*). \quad (4.5)$$

Under Assumptions 4.1.1–4.1.6,

$$\int_{\sqrt{N}(\Theta - T_N)} |\pi_N(t | X_N) - \varphi_{0, V_{\theta^*}^{-1}}(t)| dt \xrightarrow{P_0^{(N)}} 0, \quad (4.6)$$

where  $\varphi_{\mu, \Sigma}$  denotes the density of the  $d$ -dimensional normal distribution with mean vector  $\mu \in \mathbb{R}^d$  and covariance matrix  $\Sigma \in \mathbb{R}^{d \times d}$ . If in addition Assumption 4.1.7 holds, then also

$$\int_{\sqrt{N}(\Theta - T_N)} (1 + \|t\|) |\pi_N(t | X_N) - \varphi_{0, V_{\theta^*}^{-1}}(t)| dt \xrightarrow{P_0^{(N)}} 0. \quad (4.7)$$

The proof of the proposition is deferred to the appendix (see Section 4.9). In contrast to the well-specified case, the covariance matrix  $V_{\theta^*}^{-1}$  of the limiting distribution does not necessarily match the covariance matrix of the true posterior distribution or the covariance matrix of the misspecified MLE. In fact, the marginals of  $V_{\theta^*}^{-1}$  can be larger, equal, or smaller than the corresponding marginal variances of the MLE, see (Kleijn and van der Vaart 2012). Therefore, the credible sets resulting from the misspecified posterior distribution can either be overconfident or conservative, depending on the true distribution  $P_0^{(N)}$ , and can deviate from the credible sets resulting from the well-specified model.

### 4.3.2 Bernstein-von Mises in hierarchical models

In this section, we apply the above misspecified Bernstein-von Mises theorem in the context of the hierarchical data-generating model (4.1) and a misspecified surrogate Gaussian likelihood. Given observations  $X_1, \dots, X_N$ , we form the posterior distribution in Equation (4.3) with the  $q_{\theta, i}$  equal to the density of the normal distribution with the correctly specified means and covariance matrices, as in Equation (4.2). The true distribution is given by the hierarchical model specified by (4.1)

The precision matrix  $V_{\theta^*} \in \mathbb{R}^{d \times d}$  of the limiting Gaussian distribution takes a particular form. The proof of the following lemma is deferred to Section 4.6.1.

**Lemma 4.3.2.** *Consider the hierarchical model (4.1) and suppose that  $(\mu_{\theta}^{(i)}, \Sigma_{\theta}^{(i)}) \neq (\mu_{\theta_0}^{(i)}, \Sigma_{\theta_0}^{(i)})$ , for all  $\theta \neq \theta_0$ . Then the map  $\theta \mapsto P_{\theta_0, i} \log(p_{\theta_0, i}/q_{\theta, i})$  has a unique minimum at  $\theta = \theta_0$ , i.e.  $\theta_N^* = \theta_0$ . Assume furthermore that the maps  $\theta \mapsto \mu_{\theta}^{(i)}$  and  $\theta \mapsto \Sigma_{\theta}^{(i)}$  are twice continuously differentiable in a neighborhood around  $\theta_0$ ,*

and that  $\Sigma_{\theta_0}^{(i)}$  is invertible. Then, with  $v_l^{(i)} := (\Sigma_{\theta_0}^{(i)})^{-1/2} \frac{d}{d\theta_l} \mu_{\theta}^{(i)}|_{\theta=\theta_0}$  and  $A_l^{(i)} := (\Sigma_{\theta_0}^{(i)})^{-1} \frac{d}{d\theta_l} \Sigma_{\theta}^{(i)}|_{\theta=\theta_0}$ ,  $l = 1, \dots, d$ , the Hessian of the Kullback-Leibler divergence at  $\theta = \theta_0$  is given by the positive semi-definite matrix

$$V_{\theta_0}^{(i)} = \frac{1}{2} \begin{pmatrix} \text{tr}(A_1^{(i)} A_1^{(i)}) & \cdots & \text{tr}(A_1^{(i)} A_d^{(i)}) \\ \vdots & \ddots & \vdots \\ \text{tr}(A_d^{(i)} A_1^{(i)}) & \cdots & \text{tr}(A_d^{(i)} A_d^{(i)}) \end{pmatrix} + \begin{pmatrix} (v_1^{(i)})^\top v_1^{(i)} & \cdots & (v_1^{(i)})^\top v_d^{(i)} \\ \vdots & \ddots & \vdots \\ (v_d^{(i)})^\top v_1^{(i)} & \cdots & (v_d^{(i)})^\top v_d^{(i)} \end{pmatrix}. \quad (4.8)$$

The matrix  $V_{\theta_0}^{(i)}$  given in Equation (4.8) is the sum of two nonnegative-definite matrices and is strictly positive-definite as soon as one of the matrices on the right side of the display is invertible. The second matrix is invertible if and only if the vectors  $\frac{d}{d\theta_l} \mu_{\theta}^{(i)}|_{\theta=\theta_0}$ ,  $l = 1, \dots, d$ , are linearly independent. The matrix  $V_{\theta_0}^{(i)}$  in Equation (4.8) typically does not have a closed form analytic expression and hence numerical methods have to be used to evaluate it; see Section 4.5 for examples.

The following theorem specializes Proposition 4.3.1 to the Gaussian misspecification of model (4.1). The proof of the theorem consists of verifying the conditions of the proposition, and is given in Section 4.6.2.

**Theorem 4.3.3.** (i) Consider the hierarchical model (4.1) with compact parameter set  $\Theta \subset \mathbb{R}^d$ , and mean and variance functions  $\theta \mapsto \mu_{\theta}^{(i)}$  and  $\theta \mapsto \Sigma_{\theta}^{(i)}$  that are independent of  $i$ . Assume that  $(\mu_{\theta}^{(i)}, \Sigma_{\theta}^{(i)}) \neq (\mu_{\theta_0}^{(i)}, \Sigma_{\theta_0}^{(i)})$ , for every  $\theta \neq \theta_0$ . Assume that  $\theta_0 \in \text{int}(\Theta)$  and that the maps  $\theta \mapsto \mu_{\theta}^{(i)}$  and  $\theta \mapsto \Sigma_{\theta}^{(i)}$  are twice continuously differentiable in a neighborhood of  $\theta_0$ , and that the matrices  $\Sigma_{\theta_0}^{(i)}$  and  $V_{\theta_0}^{(i)}$  as given by Equation (4.8) are invertible. Assume that  $\sup_{i \in \mathbb{N}} P_{\theta_0}^{(i)} \|X_i\|^4 < \infty$  and that the prior density  $\pi$  is continuous at  $\theta = \theta_0$ , with  $\pi(\theta_0) > 0$ . Then Equation (4.6) of Proposition 4.3.1 holds with  $\theta^* = \theta_0$  and  $V_{\theta^*} = V_{\theta_0}$ .

(ii) The condition that the mean and variance functions  $\mu_{\theta}^{(i)}$  and  $\Sigma_{\theta}^{(i)}$  do not depend on  $i$  can be dropped if the second derivatives of these functions are uniformly equicontinuous ( $i \in \mathbb{N}$ ), the limit  $V_{\theta_0} := \lim_{N \rightarrow \infty} N^{-1} \sum_{i=1}^N V_{\theta_0}^{(i)}$  exists and is invertible, and, for every  $\delta > 0$ ,

$$\limsup_{N \rightarrow \infty} \sup_{\theta: \|\theta - \theta_0\| \geq \delta} \frac{1}{N} \sum_{i=1}^N \left( (\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)})^{\top} (\Sigma_{\theta}^{(i)})^{-1} (\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)}) + \text{tr}(\Sigma_{\theta_0}^{(i)} (\Sigma_{\theta}^{(i)})^{-1} - I) - \log \det(\Sigma_{\theta_0}^{(i)} (\Sigma_{\theta}^{(i)})^{-1}) \right) > 0. \quad (4.9)$$

The theorem ensures that the misspecified posterior distribution accumulates its mass within balls with a radius of the order  $1/\sqrt{N}$  around the true parameter  $\theta_0$  and hence the posterior distribution is consistent at the optimal rate. However, the posterior covariance matrix might not match the inverse Fisher information, and hence credible sets do not necessarily coincide with confidence sets, not even asymptotically. See our numerical analysis in Section 4.5 for examples.

The left-hand side of Equation (4.9) involves the average of Kullback-Leibler divergence between the surrogate Gaussian distributions as defined in Equation (4.2). This condition can be regarded as ensuring that the model is, asymptotically, identifiable.

## 4.4 Applications

In this section, we investigate the behaviour of the misspecified posterior distribution resulting from the surrogate Gaussian likelihoods more closely in a number of instantiations of the model (4.1). First, we consider the integral of a shifted and squared Brownian motion, where the scalar shift is the parameter of interest. This involves perhaps one of the simplest possible non-linear transformation  $T_{\theta,i}$  of the nuisance functional parameter  $f_i$ . Taking different limits of the integral gives a non-i.i.d. model for the observations. Next, we investigate two more practically relevant examples, in which the operator  $T_{\theta,i}$  is the forward map of a partial differential equation (PDE). These models provide a more complex non-linear operator  $T_{\theta,i}$  and can serve as a platform for investigating more complicated PDE-driven operators. First, we consider the time-independent Schrödinger equation, with the boundary condition set equal to a function parametrized by the finite-dimensional parameter  $\theta$ . In this case the operators  $T_{\theta,i}$  are the same for each observation hence the data are i.i.d. As a second example, we consider general parabolic PDEs and evaluate the time variable at different time points  $t_i$ . We let the boundary condition and the potential in the parabolic PDE depend on the finite-dimensional parameter  $\theta$  of interest. The different time points render this in a non-i.i.d. setting. This resembles situations in astronomy,



where snapshots of the universe at different times of its evolvement are observed today.

#### 4.4.1 Square integral operator

Take  $\Theta$  to be a compact set in  $\mathbb{R}$ , and let  $z \geq 0$ . Define the map  $T_{z,\theta} : L^2[0, z] \rightarrow L^2[0, z]$  to be the map given by

$$T_{z,\theta}(h)(t) = \int_0^t (h(s) - \theta)^2 ds, \quad t \in [0, z]. \quad (4.10)$$

The operator  $T_{z,\theta}$  is non-linear and parametrized by the one-dimensional shift parameter  $\theta$ . We assume that the functional parameter  $f$  in (4.1) is a Brownian motion on  $[0, z]$ .

We consider observing the projection of the function  $t \mapsto T_{z,\theta}f(t)$  on the Legendre polynomials. The normalized Legendre polynomials on the interval  $[0, z]$  take the form, for  $x \in [0, z]$  and  $j \in \mathbb{N}_0$ ,

$$e_j^z(x) = \sum_{k=0}^j a_{j,k} z^{-j-1/2} x^k \quad \text{where} \quad a_{j,k} = \sqrt{2j+1} (-1)^{j+k} \binom{j}{k} \binom{j+k}{k}. \quad (4.11)$$

By elementary computations, it can be seen that the expected values of the coefficients in the series decomposition take the form

$$\mathbb{E}\langle T_{z,\theta}(f), e_j^z \rangle = \begin{cases} \frac{1}{6} z^{5/2} + \frac{1}{2} z^{3/2} \theta^2 & j = 0, \\ \frac{1}{12} \sqrt{3} z^{5/2} + \frac{1}{6} \sqrt{3} z^{3/2} \theta^2 & j = 1, \\ \frac{1}{60} \sqrt{5} z^{5/2} & j = 2, \\ 0 & j \geq 3. \end{cases} \quad (4.12)$$

Since the expectations disappear after the third coordinate, we omit the rest of the coefficients from our model.

Take  $\eta_{z,\theta}(f) = (\langle T_{z,\theta}(f_i), e_j \rangle)_{j=0,1,2} \in \mathbb{R}^3$  and consider the observational model

$$X_i = \eta_{z_i, \theta_0}(f_i) + \gamma_i, \quad \text{where } \gamma_i \sim N(0, \Lambda), \quad i = 1, \dots, N, \quad (4.13)$$

with  $\Lambda \in \mathbb{R}^{3 \times 3}$  a known positive definite covariance matrix. The variables  $f_i$  and  $\gamma_i$  are independent, and so are the vectors  $(f_i, \gamma_i, X_i)$  across  $i = 1, \dots, N$ . The goal is to make inference about the parameter  $\theta_0$ .

The hierarchical structure renders the distribution of  $X_i$  non-Gaussian. To simplify the estimation procedure, we fit a Gaussian surrogate model to  $X_i$ , with mean  $\mu_\theta^{(i)} = \mathbb{E}_{\theta, z_i}(X_i) = (\mathbb{E}\langle T_{z_i, \theta}(f), e_j^{z_i} \rangle)_{j=1,2,3} \in \mathbb{R}^3$  and covariance matrix  $\Sigma_\theta^{(i)} = \text{Cov}_{\theta, z_i}(X_i) \in \mathbb{R}^{3 \times 3}$ . We endow  $\theta$  with a prior  $\Pi$ , and compute the misspecified posterior Equation (4.3) using the given normal distribution as the surrogate likelihood  $q_{\theta, i}$ , for  $i = 1, \dots, N$ . The following corollary gives the limiting Gaussian law of the posterior distribution under mild assumption on the prior.

**Corollary 4.4.1.** *Let  $\Theta \subset (0, \infty)$  be compact and endowed with a prior  $\Pi$  with density  $\pi$  that is continuous and strictly positive at  $\theta = \theta_0$ . Let  $(z_i)_{i \in \mathbb{N}}$  be a given positive sequence and assume that the observations  $X_1, \dots, X_N$  follow Equation (4.13). Then*

1. *If the sequence  $(z_i)_{i \in \mathbb{N}}$  is bounded, satisfies  $\limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N z_i^3 > 0$  and  $\lim_{N \rightarrow \infty} V_{\theta^*, N} = V_{\theta^*}$  holds, then the misspecified posterior in Equation (4.3) converges to the Gaussian limit as in Equation (4.6).*
2. *If  $z_i \rightarrow 0$ , then  $\lim_{N \rightarrow \infty} V_{\theta^*, N} = 0$ .*

The proof of the corollary is based on Theorem 4.3.3 and is deferred to Section 4.7.1. The limiting misspecified posterior variance  $V_{\theta^*}^{-1}$  does not necessarily match the variance of the true posterior distribution, but in fact, depends on the value of the true shift parameter  $\theta_0$ . We compare the behaviour of the true and misspecified posterior numerically for synthetic data sets in Section 4.5.

**Remark 4.1.** *This model does not satisfy the condition imposed in (Kleijn and van der Vaart 2012) for obtaining contraction at  $\sqrt{N}$ -rate. Specifically, in that paper it is assumed that  $P_0(e^{sm^*(X)}) < \infty$ , where  $P_0$  is the true data-generating distribution and  $m^*$  satisfies*

$$\log \frac{q_{\theta_1}}{q_{\theta_2}}(X) \leq m^*(X) \|\theta_1 - \theta_2\|_2 \quad P_0 - a.s., \quad (4.14)$$

for  $\theta_1, \theta_2$  in an open neighborhood of  $\theta^*$ . To see that this fails, for simplicity take  $\theta_0 = 0$  (hence  $\theta^* = 0$  as well) and  $d = 1$ , i.e. consider only the first coordinate of the three-dimensional observation. Recalling that the 0th Legendre-polynomial is equal to 1 on the whole unit interval, we get that  $X_i = \int_0^1 \int_0^t f_i(s)^2 ds dt + \gamma_i$ , for  $f_i$  i.i.d. Brownian motion and  $\gamma_i \stackrel{i.i.d.}{\sim} N(0, \Lambda)$ . Then due to Equations (4.12) and

(4.24), we have

$$\begin{aligned} \left| \log \frac{q_{\theta_1, i}}{q_{\theta_2, i}}(X_i) \right| &= \frac{1}{2} \left| \log \left( \frac{\sigma_{\theta_2}^{(i)}}{\sigma_{\theta_1}^{(i)}} \right)^2 + \left( \frac{X_i - \mu_{\theta_2}^{(i)}}{\sigma_{\theta_2}^{(i)}} \right)^2 - \left( \frac{X_i - \mu_{\theta_1}^{(i)}}{\sigma_{\theta_1}^{(i)}} \right)^2 \right| \\ &\leq m_0(X_i) \|\theta_1 - \theta_2\| \end{aligned}$$

for some  $m_0(X_i)$  that is quadratic in  $X_i$ . Then by Jensen's inequality

$$\int_0^1 \int_0^t f_i(s)^2 ds dt \geq \left( \int_0^1 \int_0^t f_i(s) ds dt \right)^2$$

and noting that the term on the right-hand side is a multiple of a  $\chi_1^2$  distributed random variable we conclude that for any  $s > 0$ ,  $P_{0, i} e^{sm_0(X_i)}$  is not finite.

#### 4.4.2 Schrödinger equation

The time-independent Schrödinger equation is a simple PDE-constrained non-linear inverse problem, which can serve as a benchmark for testing methodology. The main focus of the literature so far has been the recovery of the underlying potential function  $f$ . Minimax posterior contraction rates using multi-scale analysis were derived for Gaussian priors in (Nickl, van de Geer and Wang 2020; Monard, Nickl and Paternain 2021a) and uniform sequence priors in (Nickl 2018). Furthermore, semi-parametric Bernstein-von Mises results on linear functionals were derived in (Nickl 2018; Monard, Nickl and Paternain 2021a). An alternative method based on the underlying equation was shown to achieve adaptive minimax contraction rate and reliable uncertainty quantification in Chapter 2. However, none of these results consider a hierarchical setting, where in the data-generating process the function  $f$  is taken to be random and the focus is on recovering certain parametric aspects of the forward operator, which is the focus here. This study is motivated by applications from physics, astronomy, and medicine, as discussed in the introduction.

Consider a bounded domain  $\mathcal{O} \subset \mathbb{R}^p$  with  $C^\infty$ -boundary and define the elliptic operator  $T_\theta : L^2(\mathcal{O}) \rightarrow L^2(\mathcal{O})$  to be the solution operator of the time-independent Schrödinger equation with parametric boundary condition  $g_\theta$ ,  $\theta \in \Theta \subset \mathbb{R}^d$ , i.e.  $u = T_\theta(f)$  solves

$$\begin{cases} \Delta u - 2fu = 0, & \text{on } \mathcal{O}, \\ u = g_\theta, & \text{on } \partial\mathcal{O}. \end{cases} \quad (4.15)$$

The existence of a unique solution  $T_\theta(f) \in C^2(\overline{\mathcal{O}})$  is guaranteed whenever  $f > 0$ ,  $\|f\|_{C^s(\mathcal{O})} < \infty$  for some  $s > 0$  and  $g_\theta \in C^{s+2}(\mathcal{O})$  (see Proposition 25 in (Nickl 2018)). Take the first  $p \geq d$  coordinates of  $T_\theta(f)$  relative to an appropriate orthonormal basis  $(e_j)_{j \in \mathbb{N}}$  of  $L^2(\mathcal{O})$ , and define  $\eta_\theta(f) = (\langle T_\theta(f), e_j \rangle)_{j=1, \dots, p} \in \mathbb{R}^p$ .

Our observational model is given by (4.1), with  $G$  a distribution that is supported on the set of positive functions  $f \in L^2(\mathcal{O})$ , and given  $f_i \stackrel{\text{i.i.d.}}{\sim} G$ ,

$$X_i = \eta_{\theta_0}(f_i) + \gamma_i, \quad \gamma_i \stackrel{\text{i.i.d.}}{\sim} N(0, \Lambda), \quad i = 1, \dots, N. \quad (4.16)$$

The observations  $X_1, \dots, X_N$  are i.i.d., but the functions  $u_{f_i} = T_{\theta_0}(f_i)$  and hence the mean functions of the  $X_i$  are random.

As a surrogate likelihood  $q_\theta$ , we take the Gaussian distribution with mean and covariance set equal to the mean  $\mu_\theta$  and covariance matrix  $\Sigma_\theta$  of the  $X_i$  under the true data-generating distribution. We endow  $\theta \in \Theta \subset \mathbb{R}^d$  with a prior distribution  $\Pi$ . The following corollary shows that the corresponding misspecified posterior is asymptotically Gaussian with a covariance matrix given in Equation (4.8). The proof is based on Theorem 4.3.3 and is deferred to Section 4.7.2.

**Corollary 4.4.2.** *Consider the observational model (4.16) and assume that the function  $\theta \mapsto g_\theta(x)$  in Equation (4.15) is twice differentiable in  $\theta$  for every  $x \in \partial\mathcal{O}$ , and the functions  $g_\theta$ ,  $\nabla g_\theta$  and  $\nabla^2 g_\theta$  are uniformly bounded in  $\theta$  on  $\partial\mathcal{O}$ . Furthermore, assume that the vectors  $\frac{d}{d\theta_j} \mu_\theta|_{\theta=\theta_0} \in \mathbb{R}^p$ ,  $j = 1, \dots, d$  are linearly independent. Finally, assume that the prior  $\Pi$  possesses a density  $\pi$  that is continuous and strictly positive at  $\theta = \theta_0$ . Then the misspecified posterior distribution admits the Gaussian approximation given in Equation (4.7) with  $\theta^* = \theta_0$  and  $V_{\theta^*}$  given by Equation (4.8).*

We investigate the non-asymptotic behaviour of the misspecified posterior in a numerical study with simulated data in Section 4.5. We show that the size of the deviation between the true and misspecified posterior depends on the boundary condition. In all cases, the misspecified posterior distribution achieves the parametric recovery rate  $1/\sqrt{N}$  for the underlying true parameter  $\theta_0$  of interest.

#### 4.4.3 Parabolic PDE-constrained inverse problem

In this section, we consider parabolic PDEs with different time horizons for the observations. This model is a step towards the more complex PDE-constrained

inverse problems considered in practice. It illustrates the case of data that is observed at different time points, which influences the operator and hence the distribution of the data.

The existing frequentist Bayesian literature on versions of this model considers the case in which the driving function ( $f$  in the following) is fixed, and focuses on inference on this function. Specifically (Kekkonen 2022) derives a minimax posterior contraction rate, using multi-scale analysis in the context of the heat-equation with additional absorption term. In the same model adaptation and frequentist uncertainty quantification was obtained using a linearisation approach in Chapter 2. Here we consider the different setups of a hierarchical model with a random nuisance parameter, and aim to recover parametric aspects of the non-linear partial differential equation.

Let  $\mathcal{O}$  be an open, bounded domain in  $\mathbb{R}^d$  with boundary  $\partial\mathcal{O}$ . For a non-negative continuous function  $f : \mathcal{O} \rightarrow \mathbb{R}$ , let  $u = T_\theta f$  be the solution to the following elliptic PDE, for given  $t_{\max} > 0$ ,

$$\begin{cases} -\frac{\partial u(t, x)}{\partial t} + \mathcal{A}u(t, x) = f(x), & \text{on } (0, t_{\max}) \times \mathcal{O}, \\ u(t, x) = g_\theta(t, x), & \text{on } [(0, t_{\max}) \times \partial\mathcal{O}] \cup [\{t_{\max}\} \times \overline{\mathcal{O}}], \end{cases} \quad (4.17)$$

where, for  $v \in C^2(\mathcal{O})$ ,

$$\mathcal{A}v(t, x) := -\frac{1}{2} \text{tr}(a(x)D^2v(t, x)) - \langle b(x), Dv(t, x) \rangle + c_\theta(x)v(t, x).$$

The function  $b : \mathcal{O} \rightarrow \mathbb{R}^d$ ,  $c_\theta : \mathcal{O} \rightarrow [0, \infty)$  is the transport vector and the positive-definite matrix  $a : \mathcal{O} \rightarrow \mathbb{R}^{d \times d}$  the diffusion of the equation. We assume that the functions satisfy the regularity conditions given in (Feehan, Gong and Song 2015). In particular, the functions  $a, b$  and  $c_\theta$  are continuous, and  $a$  is positive definite with  $a(x) = x_d^\beta \tilde{a}(x)$  for some positive-definite matrix-valued function  $\tilde{a} : \mathcal{O} \rightarrow \mathbb{R}^{d \times d}$ , and  $\beta \in (0, 2]$ , where  $x_d$  denotes the  $d$ th coordinate of the vector  $x$ . Furthermore,  $g_\theta : (0, t_{\max}) \times \partial\mathcal{O} \mapsto \mathbb{R}$  is continuous, indexed by a  $d$ -dimensional parameter  $\theta \in \Theta \subset \mathbb{R}^d$  of interest.

For a fixed sequence  $(t_i)_{i \in \mathbb{N}}$  of positive numbers, define  $T_{\theta, i} f := u_\theta(t_i, \cdot)$ . Let  $(e_j)_{j \in \mathbb{N}}$  be an orthonormal basis for  $L^2(\mathcal{O})$ , and take the first  $p$  coordinates of the corresponding series decomposition of the image  $T_{\theta, i} f$ , i.e. let  $\eta_{\theta, i}(f) = (\langle T_{\theta, i} f, e_j \rangle)_{j=1, \dots, p} \in \mathbb{R}^p$ . We assume that  $\sup_{j=1, \dots, p} \|e_j\|_\infty < \infty$ . Our  $p$ -dimensional observations are generated as in model (4.1) with  $G$  a distribution  $G$  on the space of nonnegative continuous functions such that  $\mathbb{E}_G \|f\|_\infty^4 < \infty$

and

$$X_i = \eta_{\theta,i}(f_i) + \gamma_i \quad \gamma_i \stackrel{\text{i.i.d.}}{\sim} N_p(0, \Lambda), \quad i = 1, \dots, N,$$

for some given positive-definite covariance matrix  $\Lambda \in \mathbb{R}^{p \times p}$ . Denote the mean and covariance functions of the observation by  $\mu_\theta^{(i)} = (\mathbb{E}_G \langle u_\theta(t_i, \cdot), e_j \rangle)_{j=1, \dots, p}$  and  $\Sigma_\theta^{(i)} = \text{Cov}_{\theta,i}(X_i)$ , respectively. Then we form the misspecified posterior distribution from Equation (4.3) with the misspecified likelihood  $X_i \stackrel{\text{ind}}{\sim} q_{\theta,i} = N_p(\mu_\theta^{(i)}, \Sigma_\theta^{(i)})$ . The next corollary describes the limiting law of this posterior. The proof is deferred to Section 4.7.3.

**Corollary 4.4.3.** *Let  $\theta_0 \in \text{int}(\Theta)$  and assume that  $(t_i)_{i \in \mathbb{N}}$  is a sequence of positive numbers such that  $V_{\theta^*} = \lim_{N \rightarrow \infty} N^{-1} \sum_{i=1}^N V_{\theta^*}^{(i)} \succ 0$  exists, where  $V_{\theta^*}^{(i)}$  is defined in Equation (4.8), and the boundary condition  $g_\theta(t, x)$  and the potential functions  $c_\theta(x)$  are twice differentiable with respect to  $\theta$  with equicontinuous second derivatives. Assume that  $(\mu_\theta^{(i)}, \Sigma_\theta^{(i)}) \neq (\mu_{\theta_0}^{(i)}, \Sigma_{\theta_0}^{(i)})$  holds for all  $\theta \neq \theta_0$ . Assume furthermore that for all  $\delta > 0$ ,*

$$\limsup_{N \rightarrow \infty} \sup_{\theta: \|\theta - \theta_0\| > \delta} \frac{1}{N} \sum_{i=1}^N \|\mu_\theta^{(i)} - \mu_{\theta_0}^{(i)}\|^2 > 0.$$

*Finally, let  $\Pi$  be a prior on the compact set  $\Theta$  with a density  $\pi$  that is positive and continuous at  $\theta_0$ . Then the misspecified posterior admits the Gaussian approximation given in Equation (4.7) with  $\theta^* = \theta_0$  and precision matrix  $V_{\theta^*}$ .*

## 4.5 Numerical analysis

In this section, we study the numerical accuracy of the misspecified Gaussian approximation relative to the misspecified and true posterior distributions, both in the toy model with the integral of the shifted squared Brownian motion and in the model involving the time-independent Schrödinger equation.

### 4.5.1 Square integral operator

We consider the shifted, square integral operator  $T_{z_i, \theta}(f) = \int_0^{z_i} (f(s) - \theta)^2 ds$  given in Equation (4.10), where  $f$  follows a Brownian motion. The true observational model for the data  $X_1, \dots, X_N \in \mathbb{R}^p$  is given in Equation (4.13), where for simplicity we take  $z_i = 1$ , for every  $i = 1, \dots, N$ . We consider the misspecified model using the surrogate Gaussian, as discussed in Section 4.4.1.

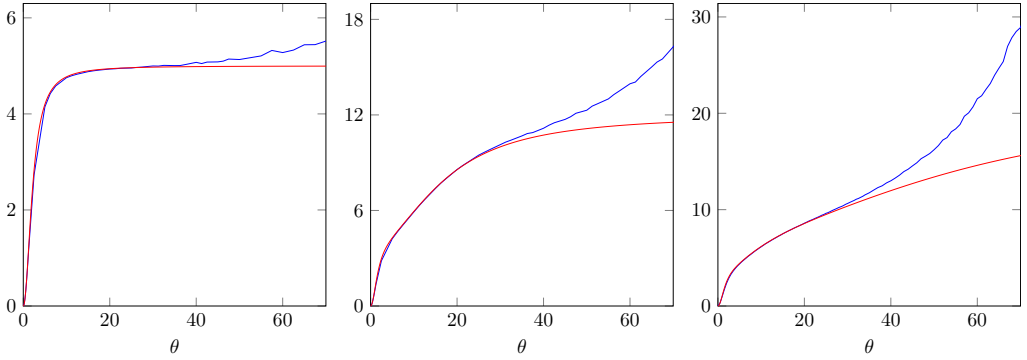


Figure 4.5.1: Square integral observational model Equation (4.13). Numerical Fisher information  $I_{\theta_0}$  (blue) and the misspecified inverse variance  $V_{\theta^*}$  (red) for dimensions  $p = 1, 2$  and  $3$ .

We start our analysis by comparing the true and misspecified Fisher informations  $I_{\theta_0}$  and  $V_{\theta^*}$ . The latter is derived in Lemma 4.3.2 (where all  $V_{\theta^*}^{(i)}$  are equal as the data are i.i.d.), with entries to the formula given in Equations (4.12) and (4.24). The computation of the true Fisher information is numerically more involved. The Fisher information is given by

$$I_{\theta} = \mathbb{E} \left( \frac{\int \frac{d}{d\theta} p_{\theta,f}(X) dP(f)}{\int p_{\theta,f}(X) dP(f)} \right)^2, \quad (4.18)$$

where  $P$  denotes the law of a Brownian motion, the outer expectation is with respect to the marginal distribution of  $X$ , and  $p_{\theta,f}$  denotes the conditional density of  $X$  given  $\theta$  and  $f$ , i.e.

$$p_{\theta,f}(x) = (2\pi)^{-d/2} (\det \Lambda)^{-1/2} e^{-\frac{1}{2} \sum (x_i - \langle T_{z_i, \theta} f, e_i \rangle) (\Lambda)_{ij}^{-1} (x_j - \langle T_{z_j, \theta} f, e_j \rangle)}.$$

We estimated the expectation in Equation (4.18) numerically using the Monte Carlo method, based on  $10^6$  draws from the Brownian motion  $f$ . For each draw, we approximated the forward map  $T_{1,\theta} f$  by computing the integral on a grid of size 100 on  $[0, 1]$ . We approximated the inner integrals in Equation (4.18) at each data point  $X$  by averaging the likelihoods  $p_{\theta,f}(x)$  and their derivatives over the  $10^6$  draws from the Brownian motion, and next the outer expectation by averaging the resulting quotients over  $10^6$  draws from  $X$ .

We considered the observational model with  $p = 1, 2$  and  $3$  coordinates and evaluated the true and misspecified Fisher informations for a wide collection of true shift parameters  $\theta_0$ , ranging from  $\theta_0 = 0$  to  $\theta_0 = 70$ . The corresponding

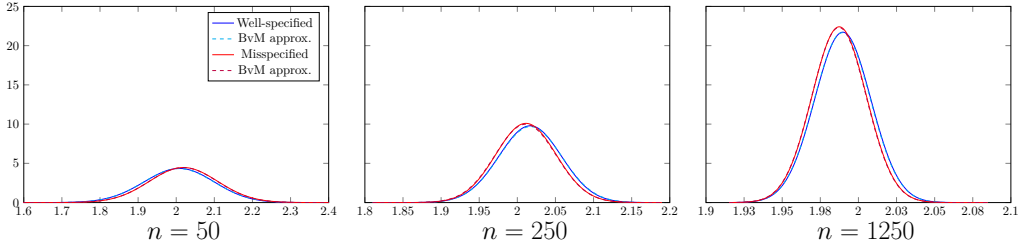


Figure 4.5.2: Square integral observational model Equation (4.13). True (blue) and misspecified (red) posterior densities and corresponding limiting Gaussian approximations in the case that the true shift parameter is  $\theta_0 = 2$ . The true curve is solid and the dashed curve is the Gaussian approximation. The sample size increases from left to right  $N = 50, 250$ , and  $1250$ .

Fisher informations are plotted in Figure 4.5.1, using blue for the true and red for the misspecified information. For small values of  $\theta_0$  they are aligned, and the misspecified model provides similar uncertainty quantification as the well-specified one. However, for large values of  $\theta_0$  the scissor opens and there is a significant and increasing discrepancy between the two informations. We observe that the misspecified Fisher information gets substantially smaller than the true one, resulting in overly conservative credible sets in the misspecified model, for all choices of  $p = 1, 2, 3$ .

We also compared the true and misspecified posterior distributions and their Gaussian approximations, where we varied the sample sizes over  $N = 50, 250$ , and  $1250$  and took the shift parameter equal to  $\theta_0 = 2$  or  $\theta_0 = 40$ , examples for which the true and misspecified Fisher informations match or not. In the former case, we considered a uniform prior for  $\theta$  supported on the interval  $[1, 3]$ . Figure 4.5.2 shows in blue and red the true and the misspecified posterior distributions in the case  $\theta_0 = 2$ , where dashed curves are their Gaussian approximations. One can see that both the true and misspecified posteriors concentrate around the true parameter  $\theta_0$ , their spread decreases at the same rate, and in both cases the Gaussian approximations are accurate. Figure 4.5.3 shows the same curves in the case  $\theta_0 = 40$ . By comparing the true density to the misspecified density, it can be seen that the centring points of the densities differ slightly. However, the interval where the centred 95% credible intervals intersect, remains large.



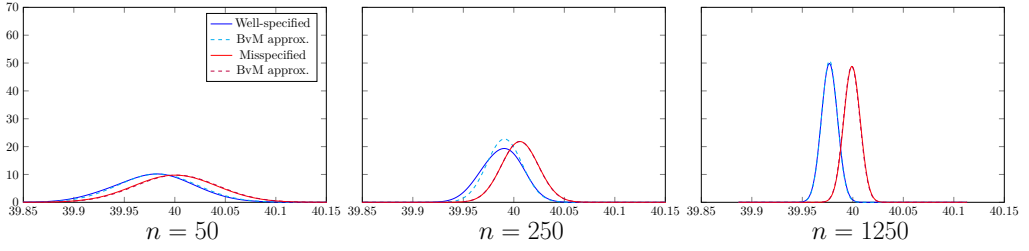


Figure 4.5.3: Square integral observational model Equation (4.13). True (blue) and misspecified (red) posterior densities and corresponding limiting Gaussian approximations in the case that the true shift parameter is  $\theta_0 = 40$ . The true curve is solid and the dashed curve is the Gaussian approximation. The sample size increases from left to right  $N = 50, 250$ , and  $1250$ .

## 4.5.2 Schrödinger equation

We consider the observational model Equation (4.16) resulting from the time-independent Schrödinger equation (4.15), investigated in Section 4.4.2. We choose  $\mathcal{O}$  equal to the unit square, and in the first layer of the hierarchical model we take  $f(x, y) = e^{2B_1(x)+3B_2(y)}$ , for  $(x, y) \in [0, 1]^2$ , where  $B_1, B_2$  are i.i.d. standard Brownian motions. Furthermore, we set the boundary condition to  $g_\theta(x, y) = (x - \frac{1}{2})^2 + \theta^2 y$ , for  $(x, y) \in \partial[0, 1]^2$ .

We focus on comparing the true and misspecified Fisher informations  $I_\theta$  and  $V_{\theta^*}$ , for values of the true parameter ranging from  $\theta_0 = 0$  to  $\theta_0 = 10$ . We considered the  $p = 1, 2$  and 3-dimensional observational models. The true Fisher information was computed numerically from Equation (4.26). We simulated a Brownian motion  $(\tilde{B}_t)_{t \geq 0}$  and function  $f$  on an equidistantly spaced  $100 \times 100$  grid of  $[0, 1]^2$ . For each  $X_i$ , we drew 100 i.i.d. observations of  $f$ , and used these to estimate the Fisher information through Equation (4.18). We used the basis functions  $e_i(x)e_j(y)$  for  $0 \leq i, j \leq 3$ , where  $e_i$  denotes the  $i$ th normalized Legendre basis on  $[0, 1]$ .

Figure 4.5.4 shows the true (blue) and misspecified (red) Fisher informations as a function of  $\theta_0$ . One observes that there is a slight deviation between the two informations. For  $p = 1$ , the misspecified Fisher information dominates the true Fisher information, but for  $p = 3$  the true Fisher dominates the misspecified Fisher information.

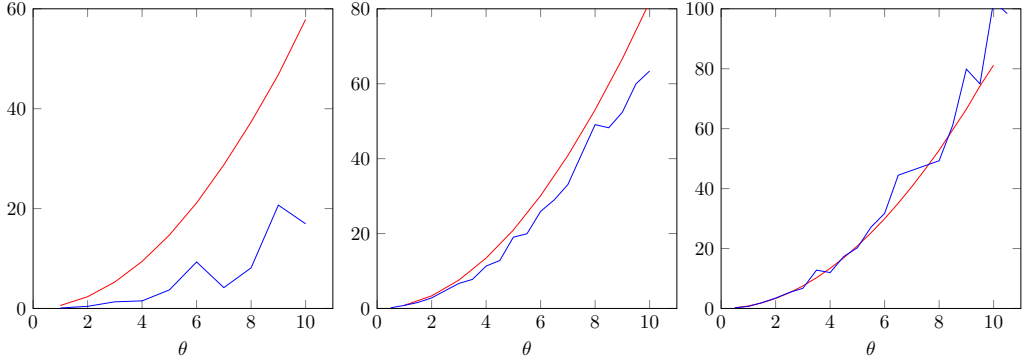


Figure 4.5.4: Schrödinger model (4.16). True Fisher information  $I_{\theta_0}$  (blue) and misspecified Fisher information  $V_{\theta_0}$  (red) as function of  $\theta_0$  for observations of dimensions  $p \times p$  with  $p = 1, 3$  and  $4$  (left to right).

## 4.6 Proof of the hierarchical Bernstein von-Mises theorem

In this section, we give the proofs of Lemma 4.3.2 and Theorem 4.3.3.

### 4.6.1 Proof of Lemma 4.3.2

For simplicity, we omit the index  $i$  from the notation.

We show the more general statement that the Kullback-Leibler divergence is uniquely minimized at  $(\mu, \Sigma) = (\mu_{\theta_0}, \Sigma_{\theta_0})$ , from which it follows that  $\theta^* = \theta_0$  is the minimizer of  $\theta \mapsto P_0 \log(p_0/p_\theta)$ .

Up to an additive constant minus twice the log-likelihood of the normal distribution with mean  $\mu$  and covariance matrix  $\Sigma$  is equal to  $\log \det \Sigma + (X - \mu)^\top \Sigma^{-1} (X - \mu)$ . The expectation of this relative to an arbitrary distribution with  $\mathbb{E}X = \mu_{\theta_0}$  and covariance matrix  $\text{Cov } X = \Sigma_{\theta_0}$  is equal to

$$\begin{aligned} \log \det \Sigma + \mathbb{E}((X - \mu_{\theta_0})^\top \Sigma^{-1} (X - \mu_{\theta_0}) + (\mu - \mu_{\theta_0})^\top \Sigma^{-1} (\mu - \mu_{\theta_0})) \\ = \log \det \Sigma + \text{tr}(\Sigma_{\theta_0} \Sigma^{-1}) + (\mu - \mu_{\theta_0})^\top \Sigma^{-1} (\mu - \mu_{\theta_0}). \end{aligned}$$

This is clearly minimized at  $\mu = \mu_{\theta_0}$ . It remains to find the minimizer of  $\Sigma \mapsto \log \det(\Sigma) + \text{tr}(\Sigma^{-1} \Sigma_{\theta_0})$ . Taking  $B := \Sigma_{\theta_0}^{-1/2} \Sigma \Sigma_{\theta_0}^{-1/2} \succ 0$ , equivalently we aim to minimize

$$B \mapsto \log \det(\Sigma_{\theta_0}^{1/2} B \Sigma_{\theta_0}^{1/2}) + \text{tr}(B^{-1}) = \log \det \Sigma_{\theta_0} + \log \det B + \text{tr}(B^{-1}).$$

Let  $(\lambda_j)_{j=1,\dots,p}$  denote the eigenvalues of  $B$ . Then the above expression simplifies to  $c + \sum_{j=1}^p \log \lambda_j + \lambda_j^{-1}$ , which is minimized at  $\lambda_j = 1$ , for all  $j = 1, \dots, p$ . Hence the minimizer is  $\Sigma = \Sigma_{\theta_0}$ .

We now verify the formula for the Hessian at  $\theta = \theta_0$ . The second derivative of the Kullback-Leibler divergence with respect to  $\theta_l$  and  $\theta_k$  is given by

$$\begin{aligned} \frac{d^2}{d\theta_l d\theta_k} P_0 \log(p^\circ/p_\theta)|_{\theta=\theta_0} &= -\frac{1}{2} \det(\Sigma_{\theta_0})^{-2} \left( \frac{d}{d\theta_l} \det(\Sigma_{\theta_0}) \right) \left( \frac{d}{d\theta_k} \det(\Sigma_{\theta_0}) \right) \\ &+ \frac{1}{2} \det(\Sigma_{\theta_0})^{-1} \frac{d^2}{d\theta_l d\theta_k} \det(\Sigma_{\theta_0}) + \left( \frac{d}{d\theta_l} \mu_{\theta_0} \right)^\top \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_k} \mu_{\theta_0} \right) \\ &+ \frac{1}{2} \sum_{i,j} \Sigma_{\theta_0,i,j} \frac{d^2}{d\theta_l d\theta_k} \Sigma_{\theta_0,i,j}^{-1}. \end{aligned}$$

Combining this with Lemmas 4.8.1 and 4.8.2, we see that

$$\begin{aligned} \frac{d^2}{d\theta_l d\theta_k} P_0 \log(p_0/p_\theta)|_{\theta=\theta_0} &- \left( \frac{d}{d\theta_l} \mu_{\theta_0} \right)^\top \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_k} \mu_{\theta_0} \right) \\ &= \frac{1}{2} \text{tr} \left( \Sigma_{\theta_0}^{-1} \frac{d^2}{d\theta_l d\theta_k} \Sigma_{\theta_0} \right) - \frac{1}{2} \text{tr} \left( \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_l} \Sigma_{\theta_0} \right) \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_k} \Sigma_{\theta_0} \right) \right) \\ &+ \frac{1}{2} \text{tr} \left( \Sigma_{\theta_0} \frac{d^2}{d\theta_l d\theta_k} \Sigma_{\theta_0}^{-1} \right). \end{aligned}$$

By applying Equation (4.30) twice, we obtain

$$\begin{aligned} \frac{d^2}{d\theta_l d\theta_k} \Sigma_\theta^{-1} &= \Sigma_\theta^{-1} \left( \frac{d}{d\theta_l} \Sigma_\theta \right) \Sigma_\theta^{-1} \left( \frac{d}{d\theta_k} \Sigma_\theta \right) \Sigma_\theta^{-1} - \Sigma_\theta^{-1} \left( \frac{d^2}{d\theta_l d\theta_k} \Sigma_\theta \right) \Sigma_\theta^{-1} \\ &+ \Sigma_\theta^{-1} \left( \frac{d}{d\theta_k} \Sigma_\theta \right) \Sigma_\theta^{-1} \left( \frac{d}{d\theta_l} \Sigma_\theta \right) \Sigma_\theta^{-1}. \end{aligned}$$

Plugging this expression into the formula one before the previous display results in

$$\begin{aligned} \frac{d^2}{d\theta_l d\theta_k} P_0 \log(p_0/p_\theta)|_{\theta=\theta_0} & \tag{4.19} \\ &= \frac{1}{2} \text{tr} \left( \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_l} \Sigma_{\theta_0} \right) \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_k} \Sigma_{\theta_0} \right) \right) + \left( \frac{d}{d\theta_l} \mu_{\theta_0} \right)^\top \Sigma_{\theta_0}^{-1} \left( \frac{d}{d\theta_k} \mu_{\theta_0} \right). \end{aligned}$$

This in turn implies the formula in Equation (4.8) for  $V_{\theta^*}$ .

Now note that Equation (4.8) can be rewritten in the form

$$\frac{1}{2} \sum_{1 \leq k, l \leq p} (A_{1,kl}, \dots, A_{p,kl})(A_{1,kl}, \dots, A_{p,kl})^\top + [v_1^\top, \dots, v_p^\top][v_1^\top, \dots, v_p^\top]^\top,$$

where  $A_l = (A_{l,kl})_{k,l=1,\dots,p}$ . This expression is the sum of Gram matrices and hence is positive semi-definite. If the vectors  $v_j$ ,  $j = 1, \dots, d$  are linearly independent, then the second Gram matrix is positive definite, which implies the positive definiteness of  $V_{\theta^*}$ .

### 4.6.2 Proof of Theorem 4.3.3

Below we verify Assumptions 4.1, with  $\theta_N^* = \theta_0$  and  $V_{\theta_N^*, N}$  equal to the Hessian of the Kullback-Leibler divergence. Then the theorem is implied by Proposition 4.3.1 and Lemma 4.3.2.

*Assumption 4.1.1.* The minimizing point  $\theta_N^* = \theta_0$  is independent of  $N$  and is interior to  $\Theta$  by assumption.

*Assumption 4.1.2.* In the case that the mean and covariance functions do not depend on  $i$ , the Hessian is given by Equation (4.8). In the case that they do depend on  $i$ , the average Hessians converge by assumption. In both cases, the (limiting) Hessian  $V_{\theta_0}$  is positive definite by assumption.

*Assumption 4.1.3.* Because the Kullback-Leibler divergence is minimized at  $\theta = \theta_0$ , its gradient vanishes at this point. This implies that  $\nabla l^{(N)}(\theta_0)$  has mean zero. The random vector  $\nabla l^{(N)}(\theta_0)$  is a sum of independent variables, each a polynomial of degree 2 in the observations  $X_{i,k}$ ,  $k = 1, \dots, p$ . Therefore  $\text{Var}_{P_0}(N^{-1/2} \nabla l^{(N)}(\theta_0)) \lesssim N^{-1} \sum_{i=1}^N P_{0,i} \|X_i\|_2^4$ . Because  $\sup_i P_{0,i} \|X_i\|_2^4 < \infty$  by assumption, the variances are bounded in  $N$ . Combined with the vanishing means, this implies that the sequence  $N^{-1/2} \nabla l^{(N)}(\theta_0)$  is bounded in probability.

*Assumption 4.1.4.* By the formula for the expected log Gaussian likelihood given in the proof of Lemma 4.3.2, we have that  $2N^{-1} P_0(l^{(N)}(\theta) - l^{(N)}(\theta_0))$  is equal to minus the average given in Equation (4.9). In the case that the mean and variance functions  $\mu_\theta^{(i)}$  and  $\Sigma_\theta^{(i)}$  do not depend on  $i$ , these averages are fixed in  $N$ . In view of Lemma 4.3.2 and the assumptions, these averages are positive for every  $\theta \neq \theta_0$ . Furthermore, as a function of  $\theta$  they are continuous. Since  $\Theta$  is compact, it follows that the averages are bounded away from zero on the set  $\{\theta : \|\theta - \theta_0\| \geq \delta\}$ , for every  $\delta > 0$ . In the case that the mean and covariance functions do depend on  $i$ , this is true by assumption. Therefore, for the verification of Assumption 4.1.4, it suffices to show that  $\sup_\theta N^{-1} |l^{(N)}(\theta) - l^{(N)}(\theta_0)| - P_0(l^{(N)}(\theta) - l^{(N)}(\theta_0))$  tends to zero in probability.

We prove this using a bracketing argument, as given in Lemma 4.6.1. The Gaussian log-likelihood can be written in exponential family form as

$$l^{(N)}(\theta) = c_N(\theta) - \sum_{i=1}^N \sum_{k=1}^K T_{i,k}(X_i) g_{i,k}(\theta) \quad (4.20)$$

with sufficient statistics and natural parameter vectors given by

$$T_i(X_i) = \begin{pmatrix} X_i \\ \text{vec}(X_i X_i^\top) \end{pmatrix} \quad \text{and} \quad g_i(\theta) = \begin{pmatrix} (\Sigma_\theta^{(i)})^{-1} \mu_\theta^{(i)} \\ -\frac{1}{2} \text{vec}((\Sigma_\theta^{(i)})^{-1}) \end{pmatrix}. \quad (4.21)$$

Here  $\text{vec}(A)$  denotes a column vector constructed from the elements of a matrix  $A$  (in an arbitrary, but fixed way). By assumption the functions  $\theta \mapsto (\Sigma_\theta^{(i)})^{-1}$  and  $\theta \mapsto \mu_\theta^{(i)}$  are twice differentiable with equicontinuous second derivatives. The same holds for the functions  $\theta \mapsto g_{i,k}(\theta)$ ,  $i = 1, \dots, N$  and  $k = 1, \dots, K$ .

It suffices to prove that  $\sup_\theta \frac{1}{N} \left| \sum_{i=1}^N (T_{i,k}(X_i) - P_0 T_{i,k}) (g_{i,k}(\theta) - g_{i,k}(\theta_0)) \right|$  tends to zero in probability, for every fixed  $k$ . For a given covering  $\Theta = \cup_j B_j$  by balls of radius  $\delta$ , consider the brackets

$$\begin{aligned} \ell_{i,j} &= T_{i,k}^+ \inf_{\theta \in B_j} (g_{i,k}(\theta) - g_{i,k}(\theta_0)) - T_{i,k}^- \sup_{\theta \in B_j} (g_{i,k}(\theta) - g_{i,k}(\theta_0)) \\ \varepsilon_{i,j} &= T_{i,k}^+ \sup_{\theta \in B_j} (g_{i,k}(\theta) - g_{i,k}(\theta_0)) - T_{i,k}^- \inf_{\theta \in B_j} (g_{i,k}(\theta) - g_{i,k}(\theta_0)). \end{aligned}$$

We have  $\ell_{i,j} \leq T_{i,k}(g_{i,k}(\theta) - g_{i,k}(\theta_0)) \leq \varepsilon_{i,j}$ , for all  $\theta \in B_j$  and  $i \in \mathbb{N}$ , and

$$|\varepsilon_{i,j} - \ell_{i,j}| \leq |T_{i,k}| \left| \sup_{\theta \in B_j} g_{i,k}(\theta) - \inf_{\theta \in B_j} g_{i,k}(\theta) \right|.$$

Because the functions  $g_{i,k}$  are uniformly equicontinuous, for every given  $\epsilon > 0$  there exists  $\delta > 0$  such that the variation of  $g_{i,k}$  over an arbitrary ball of radius  $\delta$  is at most  $\epsilon$ . Because  $\Theta$  is compact, for given  $\delta > 0$ , there is a finite cover  $\Theta = \cup_j B_j$  by balls of radius  $\delta$ . For such a cover the right side of the display is bounded above by  $|T_{i,k}| \epsilon$ , for every  $j$  and  $i$ . It follows that  $N^{-1} \sum_{i=1}^N P_0 |\varepsilon_{i,j} - \ell_{i,j}| \leq N^{-1} \sum_{i=1}^N P_0 |T_{i,k}| \epsilon \lesssim \sup_i P_0 \|X_i\|^4 \epsilon$ . Furthermore, since  $\text{Var} T^- \vee \text{Var} T^+ \leq \text{Var} T$  and  $\text{Var}(S + T) \leq 2 \text{Var} S + 2 \text{Var} T$ , for any random variables  $S, T$ , we have

$$\begin{aligned} & \frac{1}{N^2} \sum_{i=1}^N (\text{Var}_{P_0} \ell_{i,j}(X_i) + \text{Var}_{P_0} \varepsilon_{i,j}(X_i)) \\ & \leq \frac{4}{N^2} \sum_{i=1}^N \text{Var}_{P_0} T_{i,k}(X_i) \max_j \sup_{\theta \in B_j} |g_{i,k}(\theta) - g_{i,k}(\theta_0)|^2, \end{aligned}$$

which tends to zero as  $N \rightarrow \infty$ .

*Assumption 4.1.5.* By using the Taylor expansion of  $l^{(N)}(\theta)$  around  $\theta^*$ , we obtain that

$$\begin{aligned} l^{(N)}(\theta) - l^{(N)}(\theta^*) &= \nabla l^{(N)}(\theta^*)^\top (\theta - \theta^*) + \frac{1}{2} (\theta - \theta^*)^\top \sum_{i=1}^N \sum_{k=1}^K H_{\theta^*}^{g_{i,k}}(\theta) T_{i,k}(X_i) (\theta - \theta^*), \end{aligned}$$

where  $H_{\theta^*}^{g_{i,k}}(\theta)$ , are symmetric matrices with value at index  $(r, s)$  given by

$$H_{\theta^*, r, s}^{g_{i,k}}(\theta) = 2 \int_0^1 (1-t) \frac{\partial^2}{\partial \theta_r \partial \theta_s} g_{i,k}(\theta^* + t(\theta - \theta^*)) dt. \quad (4.22)$$

Since  $V_{\theta^*, N}$  is defined as the Hessian of  $N^{-1}l^{(N)}(\theta)$  at  $\theta^*$ , we have

$$-V_{\theta^*, N} = \frac{1}{N} \sum_{i=1}^N \sum_{k=1}^K H_{\theta^*}^{g_{i,k}}(\theta^*) P_{0,i}(T_{i,k}(X_i)),$$

and arrive at

$$\frac{1}{N} R_N(\theta) = \frac{1}{N} \sum_{i=1}^N \left( \sum_{k=1}^K (H_{\theta^*}^{g_{i,k}}(\theta^*) P_{0,i} T_{i,k}(X_i) - H_{\theta^*}^{g_{i,k}}(\theta) T_{i,k}(X_i)) \right). \quad (4.23)$$

By the triangle inequality this is bounded above by

$$\begin{aligned} &\frac{1}{N} \left\| \sum_{i=1}^N \sum_{k=1}^K H_{\theta^*}^{g_{i,k}}(\theta^*) (T_{i,k}(X_i) - P_{0,i} T_{i,k}(X_i)) \right\| \\ &+ \max_{i,k} \sup_{\theta: \|\theta - \theta^*\| \leq \delta} (H_{\theta^*}^{g_{i,k}}(\theta^*) - H_{\theta^*}^{g_{i,k}}(\theta)) \frac{1}{N} \sum_{i=1}^N \sum_{k=1}^K P_{0,i} |T_{i,k}(X_i)|. \end{aligned}$$

The first term is an average of independent random variables with mean zero and bounded variances, and hence this term tends to zero in probability. The second term can be made arbitrarily small by choice of  $\delta$ , by the equicontinuity of the functions  $H_{\theta^*}^{g_{i,k}}$ .

*Assumption 4.1.6.* This is true by assumption.

**Lemma 4.6.1.** *For  $i \in \mathbb{N}$  and  $\theta \in \Theta$ , let  $f_{i,\theta} : \mathbb{R}^d \rightarrow \mathbb{R}$  be a measurable function. Assume that for every  $\epsilon > 0$ , there exist functions  $\ell_{i,j}, \omega_{i,j} : \mathbb{R}^d \rightarrow \mathbb{R}$ , for  $j = 1, \dots, J_\epsilon$  and  $i \in \mathbb{N}$  such that for every  $\theta \in \Theta$  there exists a  $j \in \{1, \dots, J_\epsilon\}$  such that  $\ell_{i,j} \leq f_{i,\theta} \leq \omega_{i,j}$  for every  $i \in \mathbb{N}$ , and such that*

$$\limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E}(\omega_{i,j}(X_i) - \ell_{i,j}(X_i)) \leq \epsilon$$

$$\lim_{N \rightarrow \infty} N^{-2} \sum_{i=1}^N \frac{1}{N} \sum_{i=1}^N (\text{Var } \ell_{i,j}(X_i) + \text{Var } \omega_{i,j}(X_i)) = 0,$$

for every  $j \in \{1, \dots, J_\epsilon\}$ . Then  $\frac{1}{N} \sup_{\theta \in \Theta} |\sum_{i=1}^N (f_{i,\theta}(X_i) - \mathbb{E} f_{i,\theta}(X_i))| \rightarrow 0$  in probability.

*Proof.* Take arbitrary  $\theta \in \Theta$ . By assumption there exists a  $j \in \{1, \dots, J_\epsilon\}$  such that, for  $N$  large enough,

$$\begin{aligned} \frac{1}{N} \sum_{i=1}^N [f_{i,\theta}(X_i) - P_i f_{i,\theta}] &\leq \frac{1}{N} \sum_{i=1}^N (\omega_{i,j}(X_i) - P_i \omega_{i,j}) + \frac{1}{N} \sum_{i=1}^N P_i (\omega_{i,j} - \ell_{i,j}) \\ &\leq \frac{1}{N} \sum_{i=1}^N (\omega_{i,j}(X_i) - P_i \omega_{i,j}) + \epsilon. \end{aligned}$$

Hence the supremum over  $\theta$  of the left side is bounded above by

$$\max_{j=1, \dots, J_\epsilon} \frac{1}{N} \sum_{i=1}^N (\omega_{i,j}(X_i) - P_i \omega_{i,j}) + \epsilon.$$

By the assumption  $\lim_{N \rightarrow \infty} N^{-2} \sum_{i=1}^N \text{Var } \omega_{i,j}(X_i) = 0$ , this tends in probability to  $\epsilon$ . In combination with a similar argument using the lower brackets  $\ell_{i,j}$ , we find that the sequence  $N^{-1} \sup_{\theta \in \Theta} |\sum_{i=1}^N (f_{i,\theta}(X_i) - \mathbb{E} f_{i,\theta}(X_i))| \rightarrow 0$  is asymptotically bounded in probability by  $\epsilon$ . This being true for every  $\epsilon > 0$ , gives the result.  $\square$

## 4.7 Proofs for the applications

In this section, we collect the proofs for the three examples discussed in Section 4.4. In each case we verify that the conditions of Theorem 4.3.3 hold.

### 4.7.1 Proof of Corollary 4.4.1

The mean function  $\theta \mapsto \mu_\theta^{(i)}$  is given in Equation (4.12) with  $z = z_i$ . These functions are twice differentiable, and for bounded set  $\{z_i : i \in \mathbb{N}\}$  the second derivative is bounded and equicontinuous.

Next, we deal with the function  $\theta \mapsto \Sigma_\theta^{(i)}$ . For  $k, l \in \mathbb{N}_0$ , define the coefficients

$$b_{k,l} = \frac{2}{3}((4+k)^{-1}(k+l+6)^{-1} + (k+2)^{-1}((l+4)^{-1} - (k+l+6)^{-1})) \\ - \frac{1}{3}((5+k)^{-1}(k+l+6)^{-1} + (k+1)^{-1}((l+5)^{-1} - (k+l+6)^{-1}))$$

and

$$c_{k,l} = 2((k+3)^{-1}(k+l+5)^{-1} + (k+2)^{-1}((l+3)^{-1} - (k+l+5)^{-1})) \\ - \frac{2}{3}((k+4)^{-1}(k+l+5)^{-1} + (k+1)^{-1}((l+4)^{-1} - (k+l+5)^{-1})).$$

Then for all  $j_1, j_2 \in \mathbb{N}_0$ , we have

$$\Sigma_{\theta, j_1, j_2}^{(i)} = \Lambda_{j_1, j_2} + \sum_{k=0}^{j_1} \sum_{l=0}^{j_2} a_{j_1, l} a_{j_2, k} (z_i^5 b_{k, l} + \theta^2 z_i^4 c_{k, l}). \quad (4.24)$$

See Section 4.8 for the details of the cumbersome, but elementary computations to derive this result. The limiting covariance matrix of the misspecified posterior, given in Equation (4.8), is completely determined by Equations (4.12) and (4.24).

The function  $\theta \mapsto \Sigma_{\theta, j_1, j_2}^{(i)}$  is twice differentiable and for bounded set  $\{z_i : i \in \mathbb{N}\}$  the second derivative is bounded and equicontinuous. Also note that the matrix can be written in the form  $\Sigma_\theta^{(i)} = \Lambda + Az_i^5 + \theta^2 z_i^4 B$  for positive semi-definite matrices  $\Lambda, B$  and  $C$  not depending on  $z_i \geq 0$  or  $\theta \in \mathbb{R}$  (the positive semi-definiteness of  $B$  and  $C$  follows from the positive semi-definiteness of  $\Sigma_\theta$  for all  $z_i > 0$  and  $\theta \in \mathbb{R}$ ). Using the above formula one can compute the misspecified posterior covariance matrix given in Equation (4.8). In view of  $\frac{d}{d\theta} \mu_\theta = \theta z^{3/2}(1, \frac{1}{3}, 0, \dots, 0)$ , we get that

$$V_{\theta^*}^{(i)} = \frac{1}{2} \text{tr} \left( \left[ (\Lambda + Az_i^5 + B\theta_0^2 z_i^4)^{-1} 2\theta_0 B z_i^4 \right]^2 \right) \\ + \theta_0 z_i^{3/2} (1, \frac{1}{3}, 0, \dots, 0) (\Lambda + Az_i^5 + B\theta_0^2 z_i^4)^{-1} \theta_0 z_i^{3/2} (1, \frac{1}{3}, 0, \dots, 0)^\top \\ = z_i^8 G_{\theta_0, 1}(z_i) + z_i^3 G_{\theta_0, 2}(z_i)$$



with functions  $G_{\theta_0,1}(z)$  and  $G_{\theta_0,2}(z)$  such that  $\sup_{0 \leq z \leq C} |G_{\theta_0,1}(z)| + |G_{\theta_0,2}(z)| < \infty$  for arbitrary  $C > 0$ . One can immediately obtain that  $V_{\theta^*}^{(i)} > 0$  if and only if  $\theta^* = \theta_0 \neq 0$ .

Since  $\Theta$  and  $(z_i)_{i \in \mathbb{N}}$  are bounded, we know that  $\mu_{\theta,j}^{(i)}$  are uniformly bounded. Furthermore, in view of

$$\Sigma_{\theta}^{-1} \leq \Lambda^{-1},$$

the elements of  $\Sigma_{\theta}^{-1}$  are bounded in absolute value.

The map  $(\theta, z) \mapsto \mu_{\theta,z}$  is continuous. Since  $\Theta \times \overline{\{z_i : i \in \mathbb{N}\}}$  is compact, it is uniformly continuous. Similarly,  $(\theta, z) \mapsto \Sigma_{\theta,z}$  is uniformly continuous. Note also that the function  $\theta \mapsto \Sigma_{\theta,z}$  is twice differentiable and the second derivative is also equicontinuous. Furthermore, since  $\det(\Sigma_{\theta,z}) \geq \det(\text{Cov}(\eta_{z,\theta}(f_i))) + \det(\Lambda) \geq \det(\Lambda)$  is lower bounded for all  $(\theta, z) \in \Theta \times \{z_i : i \in \mathbb{N}\}$ , the function  $(\theta, z) \mapsto \Sigma_{\theta,z}^{-1}$  is uniformly continuous as well. For the case where the sequence  $(z_i)_{i \in \mathbb{N}}$  is constant we have i.i.d. data and therefore  $\lim_{N \rightarrow \infty} V_{\theta^*,N} = V_{\theta^*} > 0$ . If  $z_i \rightarrow 0$ , then  $\lim_{z_i \rightarrow 0} V_{\theta^*}^{(i)} = 0$ , and thus  $\lim_{N \rightarrow \infty} N^{-1} \sum_{i=1}^N V_{\theta^*}^{(i)} = 0$ .

In order to show that Equation (4.9) holds, we first note that the Kullback-Leibler divergence between a  $N(0, \Sigma_{\theta})$  distribution, and a  $N(0, \Sigma_{\theta_0})$  distribution is equal to the difference  $\text{tr}(\Sigma_{\theta_0}^{(i)} (\Sigma_{\theta}^{(i)})^{-1} - I) - \log \det(\Sigma_{\theta_0}^{(i)} (\Sigma_{\theta}^{(i)})^{-1})$ . Hence this term must be non-negative. Therefore it suffices to show that

$$\limsup_{N \rightarrow \infty} \sup_{\theta: |\theta - \theta_0| \geq \delta} \frac{1}{N} \sum_{i=1}^N (\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)})^{\top} (\Sigma_{\theta}^{(i)})^{-1} (\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)}) > 0.$$

With  $\lambda_1^{(i)}, \dots, \lambda_d^{(i)}$  the eigenvalues of  $\Sigma_{\theta}^{(i)}$ , the left-hand side of the previous display is lower-bounded by

$$\limsup_{N \rightarrow \infty} \sup_{\theta: |\theta - \theta_0| \geq \delta} \frac{1}{N} \sum_{i=1}^N \inf_{j=1, \dots, d} (\lambda_j^{(i)})^{-1} \|\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)}\|^2. \quad (4.25)$$

We first lower bound  $\inf_{j=1, \dots, d} (\lambda_j^{(i)})^{-1}$  uniformly in  $i$ , which is equivalent to giving an upper bound for the eigenvalues of  $\Sigma_{\theta}^{(i)} = \text{Cov}(T_{z_i, \theta}(f)) + \Lambda$  that is uniform in  $i$ . The uniform boundedness of the coefficients  $a_{k,l}$ ,  $b_{k,l}$  and  $c_{k,l}$  combined with Equation (4.24) shows that the elements of  $\Sigma_{\theta_0}^{(i)}$  are uniformly bounded as long as  $(z_i)_{i \in \mathbb{N}}$  is bounded. This in turn shows that the eigenvalues of  $\Sigma_{\theta_0}^{(i)}$  are uniformly upper bounded, if the sequence  $(z_i)_{i \in \mathbb{N}}$  is bounded. Therefore we see that  $\inf_{j=1, \dots, d} (\lambda_j^{(i)})^{-1}$  is lower bounded by some fixed constant. This

shows that the eigenvalues of  $\Sigma_{\theta_0}^{(i)}$  are uniformly bounded as well. In order to lower bound  $\|\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)}\|^2$ , we Equation (4.12) to see that

$$\|\mu_{\theta}^{(i)} - \mu_{\theta_0}^{(i)}\|^2 = \frac{1}{3} z_i^3 (\theta + \theta_0)^2 (\theta - \theta_0)^2.$$

This is lower bounded by  $\frac{1}{3} z_i^3 |\theta + \theta_0|^2 \delta^2$  for all  $\theta \in \Theta$  such that  $|\theta - \theta_0| \geq \delta$ . Therefore Equation (4.25) is lower bounded, up to some fixed constant, by

$$\limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N z_i^3,$$

which is positive by assumption.

We finish the proof by providing upper bounds for the fourth moment. By Jensen's inequality and Fubini's theorem, we have

$$P_0^{(N)} |\langle T_{\theta, i} f, e_{i, j} \rangle|^4 \leq \int_0^{z_i} \int_0^t P_0^{(N)} (f(s) - \theta)^8 ds |e_{i, j}(t)|^4 dt < \infty,$$

where in the last inequality we used that  $f(s) \sim \mathcal{N}(0, s)$ , implying a finite eighth moment.

### 4.7.2 Proof of Corollary 4.4.2

In view of the maximum principle (see Theorem 4 in Chapter 6.4 of (Evans 2010)), for all  $x \in \mathcal{O}$ ,

$$0 \wedge \inf_{y \in \partial \mathcal{O}} g_{\theta_0}(y) \leq T_{\theta_0} f(x) \leq 0 \vee \sup_{y \in \partial \mathcal{O}} g_{\theta_0}(y),$$

implying  $P_0^{(N)} \|X_i\|^4 < \infty$ .

Proposition 25 in (Nickl 2018) implies that the solution  $T_{\theta} f$  can be represented by the Feynman-Kac formula

$$T_{\theta} f(x) = \mathbb{E}^x \left( g_{\theta}(\tilde{B}_{\tau_{\mathcal{O}}}) e^{-\int_0^{\tau_{\mathcal{O}}} f(\tilde{B}_s) ds} \right). \quad (4.26)$$

Here  $(\tilde{B}_s : s \geq 0)$  denotes a  $d$ -dimensional Brownian motion, where the superscript  $x$  denotes the fact that it has started at  $x \in \mathcal{O}$ . Here  $\tau_{\mathcal{O}}$  is the exit time of  $(\tilde{B}_s)_{s \geq 0}$  from  $\mathcal{O}$ . Recall that the multivariate Gaussian model belongs to the exponential family of distributions Equation (4.20), with natural parameter

vector  $g_\theta$  and sufficient statistics vector  $T_i(X_i)$  given in Equation (4.21). In view of the regularity assumptions on  $\theta \mapsto \mu_\theta$  and  $\theta \mapsto \Sigma_\theta$ , one can apply Fubini's theorem twice in the Feynman-Kac formula, and conclude

$$\frac{d}{d\theta_j} P_0^{(N)} \langle T_\theta f, e_l \rangle = \int_{\mathcal{O}} \mathbb{E}^x \left( \frac{d}{d\theta_j} g_\theta(\tilde{B}_{\tau_{\mathcal{O}}}) \mathbb{E}_f e^{-\int_0^{\tau_{\mathcal{O}}} f(\tilde{B}_s) ds} \right) e_l(x) dx, \quad (4.27)$$

where  $\mathbb{E}^x$  denotes the expectation with respect to Brownian motion starting at  $x \in \mathcal{O}$ , and  $\mathbb{E}_f$  denotes the expectation with respect to the random variable  $f$ . Therefore, the map  $\theta \mapsto \mu_\theta$  is smooth. We can extend this argument to show that these maps are differentiable to the same order as  $\theta \mapsto g_\theta$  is differentiable and bounded. It follows that the map  $\theta \mapsto P_0^{(N)} \log p_\theta$  is a composition of twice differentiable maps, and therefore twice differentiable. Reasoning along similar lines shows that  $\theta \mapsto \text{Cov}_\theta(X_i, X_j)$  is twice differentiable.

Finally, to prove the positive definiteness of  $V_{\theta^*}$ , note that the linear independence assumption of the vectors  $\frac{d}{d\theta_j} \mu_\theta \in \mathbb{R}^p$ , for  $j = 1, \dots, d$ , implies that the vectors  $v_j$ , for  $j = 1, \dots, d$ , are linearly independent in Equation (4.8). Therefore, the second Gram matrix in Equation (4.8) is positive definite, which implies the positive definiteness of  $V_{\theta^*}$ .

### 4.7.3 Proof of Corollary 4.4.3

Define  $X = (X_t)_{t \in [0, t_{\max}]}$  to be the solution to the stochastic differential equation defined in terms of  $a, b$  and  $c_\theta$  and initial conditions  $X_0 = x \in \mathcal{O}$  (see (Feehan, Gong and Song 2015)). By the Feynman-Kac formula the solution  $u_\theta : (0, t_{\max}) \times \mathcal{O}$  of Equation (4.17) can be expressed as

$$\begin{aligned} u_\theta(t, x) = & \mathbb{E}^x \left( g_\theta(\tau_{\mathcal{O}} \wedge t_{\max}, X_{\tau_{\mathcal{O}} \wedge t_{\max}}) \exp \left( - \int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} c_\theta(X_s) ds \right) \right) \\ & + \mathbb{E}^x \left( \int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} f(X_s) \exp \left( - \int_t^s c_\theta(X_v) dv \right) ds \right). \end{aligned} \quad (4.28)$$

Here  $\tau_{\mathcal{O}} = \tau_{\mathcal{O}, x}$  is defined as the exit time of  $X$  out of  $\mathcal{O}$ , i.e.

$$\tau_{\mathcal{O}, x} = \inf\{s \geq 0 : X_s \notin \mathcal{O}\}, \quad X_0 = x.$$

Using this formula, we proceed by showing that the collections of functions  $\{\theta \mapsto \mu_\theta^{(i)} : i \in \mathbb{N}\}$  and  $\{\theta \mapsto \Sigma_\theta^{(i)} : i \in \mathbb{N}\}$  are equicontinuous. Let  $\theta_1, \theta_2 \in \Theta$  be given. Note that

$$|\mu_{\theta_1, l}^{(i)} - \mu_{\theta_2, l}^{(i)}| \leq \|e_l\|_\infty \mathbb{E}_g \|u_{\theta_1}(t_i, \cdot) - u_{\theta_2}(t_i, \cdot)\|_\infty.$$

The second factor on the right-hand side of the preceding display can be bounded from above by triangle inequality as

$$\begin{aligned}
 & \|u_{\theta_1}(t_i, \cdot) - u_{\theta_2}(t_i, \cdot)\|_\infty \\
 & \leq \mathbb{E}^x |g_{\theta_1}(\tau_{\mathcal{O}} \wedge t_{\max}, X_{\tau_{\mathcal{O}} \wedge t_{\max}})| \\
 & \quad \times \underbrace{\left| \exp\left(-\int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} c_{\theta_1}(X_s) ds\right) - \exp\left(-\int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} c_{\theta_2}(X_s) ds\right) \right|}_{=:\Delta_1} \\
 & + \mathbb{E}^x \left| \exp\left(-\int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} c_{\theta_2}(X_s) ds\right) \right| \\
 & \quad \times \underbrace{|g_{\theta_2}(\tau_{\mathcal{O}} \wedge t_{\max}, X_{\tau_{\mathcal{O}} \wedge t_{\max}}) - g_{\theta_1}(\tau_{\mathcal{O}} \wedge t_{\max}, X_{\tau_{\mathcal{O}} \wedge t_{\max}})|}_{=:\Delta_2} \\
 & + \mathbb{E}^x \int_t^{\tau_{\mathcal{O}} \wedge t_{\max}} |f(X_s)| \underbrace{\left| \exp\left(-\int_t^s c_{\theta_1}(X_v) dv\right) - \exp\left(-\int_t^s c_{\theta_2}(X_v) dv\right) \right|}_{=:\Delta_3} ds.
 \end{aligned}$$

In view of the bounds

$$\begin{aligned}
 \Delta_1 & \leq t_{\max} \sup_{\theta \in \Theta} \|\nabla_{\theta} c_{\theta}\|_{\infty} \|\theta_1 - \theta_2\|, \\
 \Delta_2 & \leq \sup_{\theta \in \Theta} \|\nabla_{\theta} g_{\theta}\|_{\infty} \|\theta_1 - \theta_2\|, \\
 \Delta_3 & \leq t_{\max} \sup_{\theta \in \Theta} \|\nabla_{\theta} c_{\theta}\|_{\infty} \|\theta_1 - \theta_2\|,
 \end{aligned}$$

combined with the assumed bounds on  $g_{\theta}$  and  $c_{\theta}$ , we find that

$$\begin{aligned}
 \|u_{\theta_1}(t_i, \cdot) - u_{\theta_2}(t_i, \cdot)\|_{\infty} & \leq t_{\max} \sup_{\theta \in \Theta} \|g_{\theta}\|_{\infty} \sup_{\theta \in \Theta} \|\nabla_{\theta} c_{\theta}\|_{\infty} \|\theta_1 - \theta_2\| \\
 & + \sup_{\theta \in \Theta} \|\nabla_{\theta} g_{\theta}\|_{\infty} \|\theta_1 - \theta_2\| + t_{\max}^2 \sup_{\theta \in \Theta} \|f\|_{\infty} \sup_{\theta \in \Theta} \|\nabla_{\theta} c_{\theta}\|_{\infty} \|\theta_1 - \theta_2\|.
 \end{aligned} \tag{4.29}$$

We conclude that  $\{\theta \mapsto \mu_{\theta, l}^{(i)} : i \in \mathbb{N}\}$  is equicontinuous. Since  $\Sigma_{\theta, l, k}^{(i)} = P_{0, i} X_{\theta, i, l} X_{\theta, i, k} - P_{0, i} X_{\theta, i, l} P_{0, i} X_{\theta, i, k}$ , and the second term is equicontinuous, we only need to show that  $\theta \mapsto P_{0, i} X_{\theta, i, k} X_{\theta, i, l}$  is equicontinuous. We have

$$\begin{aligned}
 & |P_{0, i} X_{\theta_1, i, k} X_{\theta_1, i, l} - P_{0, i} X_{\theta_2, i, k} X_{\theta_2, i, l}| \\
 & \leq \|e_l\|_{\infty} \|e_k\|_{\infty} P_{0, i} \int_{\mathcal{O}} \int_{\mathcal{O}} \left( |u_{\theta_1}(t_i, x) - u_{\theta_2}(t_i, x)| |u_{\theta_1}(t_i, y)| \right. \\
 & \quad \left. + |u_{\theta_1}(t_i, y) - u_{\theta_2}(t_i, y)| |u_{\theta_2}(t_i, x)| \right) dx dy.
 \end{aligned}$$

If we show that the first term of the integrand on the right-hand side is equicontinuous, the claim follows. By Hölder's inequality, we bound the integral by

$$\begin{aligned} \mathbb{E}_G \int_{\mathcal{O}} \int_{\mathcal{O}} |u_{\theta_1}(t_i, x) - u_{\theta_2}(t_i, x)| |u_{\theta_1}(t_i, y)| dx dy \\ \leq \sqrt{\mathbb{E}_G \int_{\mathcal{O}} |u_{\theta_1}(t_i, x) - u_{\theta_2}(t_i, x)|^2 dx} \sqrt{\mathbb{E}_G \int_{\mathcal{O}} |u_{\theta_1}(t_i, y)|^2 dy}. \end{aligned}$$

The second factor on the right-hand side is bounded by  $\text{Vol}(\mathcal{O})^2 (\mathbb{E}_G \|g_{\theta_1}\|_{\infty} + t_{\max} \|f\|_{\infty})^2$ . The integrand in the first factor is bounded as

$$\mathbb{E}_G |u_{\theta_1}(t_i, x) - u_{\theta_2}(t_i, x)| \leq C \|\theta_1 - \theta_2\|$$

by Equation (4.29), for a constant  $C < \infty$  that does not depend on  $i$ . We conclude that  $\theta \mapsto \Sigma_{\theta, k, l}^{(i)}$  is also equicontinuous. Similar reasoning shows that the map  $(\theta, t) \mapsto (\mu_{\theta, t}, \Sigma_{\theta, t})$  is continuous on  $\Theta \times [0, t_{\max}]$  if  $g_{\theta}$ ,  $f$  and  $c_{\theta}$  are continuous and bounded in  $\theta$ . This can be extended to differentiability in  $\theta$  up to the minimum order of differentiability of  $g_{\theta}$  and  $c_{\theta}$ .

Similar to the proof in Section 4.7.1, by the assumption on  $\mu_{\theta}$  it suffices to show that the eigenvalues of  $\Sigma_{\theta_0}^{(i)}$  are uniformly bounded in  $i$ . We do this by showing that  $|\eta_{\theta_0}(f_i)|$  is uniformly bounded in  $i$ . Due to Equation (4.28), we can bound  $|\eta_{\theta_0}(f_i)|$  from above by

$$|\eta_{\theta_0}(f_i)| \leq \sup_{j=1, \dots, p} \|e_j\|_{\infty} \left( \sup_{\substack{t \in (0, t_{\max}), \\ x \in \mathcal{O}}} g_{\theta_0}(t, x) + t_{\max} \mathbb{E}^x \|f\|_{\infty} \right) < \infty.$$

Since this is uniform in  $i$ , we conclude that the entries of  $\Sigma_{\theta_0}^{(i)}$  are uniformly bounded in  $i$ , and therefore its eigenvalues  $\lambda_1^{(i)}, \dots, \lambda_d^{(i)}$  are also uniformly bounded in  $i$ .

Finally, we note that  $P_{0,i} \|X_i\|^4 \lesssim \mathbb{E}_G \|u(t_i, \cdot)\|_{\infty}^4 + P_{0,i} \|\gamma_i\|^4$ . The first term is bounded by a multiple of  $t_{\max}^4 (\|g_{\theta}\|_{\infty}^4 + \mathbb{E}_G \|f\|_{\infty}^4) < \infty$ , while the boundedness of the second term follows from the Gaussianity of  $\gamma_i$ .

## 4.8 Technical lemmas

In this section, we collect technical lemmas used to prove our main results. We start by recalling Jacobi's formula (see for instance (Magnus and Neudecker 2019)), and extending this to the second derivative of the determinant. Next,

we present the computations for deriving assertion Equation (4.24) used in the application on the square integral operator.

**Lemma 4.8.1** (Jacobi's formula). *Let  $\Sigma : \Theta \rightarrow \mathbb{R}^{n \times n}$  with  $\Theta \subset \mathbb{R}^d$  be a coordinate-wise differentiable map, such that  $\Sigma_\theta$  is invertible for each  $\theta \in \Theta$ . Then the partial derivative of the determinant for every  $j = 1, \dots, d$  is given by*

$$\frac{d}{d\theta_j} \det \Sigma_\theta = \det(\Sigma_\theta) \operatorname{tr} \left( \Sigma_\theta^{-1} \frac{d}{d\theta_j} \Sigma_\theta \right).$$

**Lemma 4.8.2.** *Let  $\Sigma : \Theta \rightarrow \mathbb{R}^{n \times n}$  be as in Lemma 4.8.1, where each coordinate is twice differentiable. Then the second derivative of  $\det(\Sigma_\theta)$  with respect to  $\theta_l$  and  $\theta_k$  is given by*

$$\begin{aligned} \frac{d^2}{d\theta_l d\theta_k} \det(\Sigma_\theta) &= \det(\Sigma_\theta) \left( \operatorname{tr} \left( \Sigma_\theta^{-1} \frac{d}{d\theta_l} \Sigma_\theta \right) \operatorname{tr} \left( \Sigma_\theta^{-1} \frac{d}{d\theta_k} \Sigma_\theta \right) \right. \\ &\quad \left. - \operatorname{tr} \left( \Sigma_\theta^{-1} \left( \frac{d}{d\theta_l} \Sigma_\theta \right) \Sigma_\theta^{-1} \frac{d}{d\theta_k} \Sigma_\theta \right) + \operatorname{tr} \left( \Sigma_\theta^{-1} \frac{d^2}{d\theta_l d\theta_k} \Sigma_\theta \right) \right). \end{aligned}$$

*Proof.* Writing the first derivative with the help of Jacobi's formula and interchanging the order of trace and differentiation, we see

$$\begin{aligned} \frac{d}{d\theta_l d\theta_k} \det(\Sigma_\theta) &= \frac{d}{d\theta_l} \operatorname{tr} \left( \det(\Sigma_\theta) \Sigma_\theta^{-1} \frac{d}{d\theta_k} \Sigma_\theta \right) \\ &= \operatorname{tr} \left( \left( \frac{d}{d\theta_l} \det(\Sigma_\theta) \right) \Sigma_\theta^{-1} \frac{d}{d\theta_k} \Sigma_\theta + \det(\Sigma_\theta) \left( \frac{d}{d\theta_l} \Sigma_\theta^{-1} \right) \frac{d}{d\theta_k} \Sigma_\theta \right. \\ &\quad \left. + \det(\Sigma_\theta) \Sigma_\theta^{-1} \frac{d^2}{d\theta_l d\theta_k} \Sigma_\theta \right). \end{aligned}$$

Next, we apply again Jacobi's formula to the first term on the right-hand side and the formula

$$\frac{d}{d\theta_l} \Sigma_\theta^{-1} = -\Sigma_\theta^{-1} \left( \frac{d}{d\theta_l} \Sigma_\theta \right) \Sigma_\theta^{-1}, \quad (4.30)$$

for the derivative of the inverse of  $\Sigma_\theta$ , see for instance (Magnus and Neudecker 2019), in the second term on the right-hand side. This results in the stated formula, concluding the proof.  $\square$

**Proof of assertion (4.24)**

Note that for the Brownian motion  $f$

$$\text{Cov}(X_{z,i}, X_{z,j}) = \mathbb{E}\langle T_{z,\theta}f, e_{z,i} \rangle \langle T_{z,\theta}f, e_{z,j} \rangle - \mathbb{E}\langle T_{z,\theta}f, e_{z,i} \rangle \mathbb{E}\langle T_{z,\theta}f, e_{z,j} \rangle.$$

The first term is equal to

$$\begin{aligned} & \mathbb{E}\langle T_{z,\theta}f, e_{z,i} \rangle \langle T_{z,\theta}f, e_{z,j} \rangle \\ &= \int_0^z \int_0^z \int_0^t \int_0^s \mathbb{E}(f_v - \theta)^2 (f_r - \theta)^2 dr dv e_{z,i}(t) e_{z,j}(s) ds dt. \end{aligned}$$

Assuming  $v \geq r$ , by elementary computations using the definition of the Brownian motion, one can obtain that  $\mathbb{E}f_v^2 f_r^2 = 3r^2 + (v - r)r$ , and  $\mathbb{E}f_r f_v^2 = 0 = \mathbb{E}f_r f_v^2$ . By elementary algebra this results in

$$\mathbb{E}(f_v - \theta)^2 (f_r - \theta)^2 = 3r^2 + (v - r)r + \theta^2 v + 4\theta^2 r + \theta^2 r + \theta^4.$$

This in turn implies the following formula for the covariance

$$\mathbb{E}\langle T_{z,\theta}f, e_{z,i} \rangle \langle T_{z,\theta}f, e_{z,j} \rangle \tag{4.31}$$

$$\begin{aligned} &= \int_0^z \int_0^z \int_0^t \int_0^s 3(r \wedge v)^2 + ((r \vee v) - (r \wedge v))(r \wedge v) \\ &+ (5(r \wedge v) + (r \vee v))\theta^2 + \theta^4 dr dv e_{z,i}(t) e_{z,j}(s) ds dt. \end{aligned} \tag{4.32}$$

The double inner integral can be explicitly calculated to be

$$\left( \frac{2}{3}(s \vee t)(s \wedge t)^3 - \frac{1}{3}(s \wedge t)^4 + \frac{1}{4}(s \wedge t)^2(s \vee t)^2 \right) \tag{4.33}$$

$$+ \theta^2 \left( \frac{1}{2}(s \wedge t)(s \vee t)^2 + \frac{5}{2}(s \vee t)(s \wedge t)^2 - \frac{2}{3}(s \wedge t)^3 \right) + \theta^4 st. \tag{4.34}$$

Next, for general  $k, \ell, m, n \in \mathbb{N}$ ,

$$\begin{aligned} & \int_0^z \int_0^z (s \wedge t)^k (s \vee t)^\ell s^m t^n ds dt \\ &= \int_0^z \int_0^t s^{k+m} t^{\ell+n} ds dt + \int_0^z \int_t^z t^{k+n} s^{\ell+m} ds dt \\ &= \frac{(2 + 2k + m + n)z^{2+k+l+m+n}}{(1 + k + m)(1 + k + n)(2 + k + l + m + n)}. \end{aligned}$$

By recalling the definitions of Legendre polynomials, see Equation (4.11), and using the formula above one can compute all integrals in Equation (4.31), where the double inner integral can be rewritten in the form Equation (4.33). For instance, the coefficient in the  $\theta^4$  term is

$$\int_0^z \int_0^z s t e_{z,i}(t) e_{z,j}(s) ds dt = \sum_{\ell=0}^i \sum_{k=0}^j a_{z,i,\ell} a_{z,j,k} \frac{z^{\ell+k+4}}{(\ell+2)(k+2)}.$$

Similarly, the coefficient of the  $\theta^2$ -term can be seen to be equal to the sum

$$\sum_{\ell=0}^i \sum_{k=0}^j a_{z,i,\ell} a_{z,j,k} \left\{ \frac{1}{2} \frac{(4+\ell+k)z^{5+k+\ell}}{(2+\ell)(2+k)(5+\ell+k)} + \frac{5}{2} \frac{(6+\ell+k)z^{5+\ell+k}}{(3+\ell)(3+k)(5+\ell+k)} - \frac{2}{3} \frac{(8+\ell+k)z^{5+\ell+k}}{(4+\ell)(4+k)(5+\ell+k)} \right\},$$

while the constant term takes the form

$$\sum_{\ell=0}^i \sum_{k=0}^j a_{z,i,\ell} a_{z,j,k} \left\{ \frac{2}{3} \frac{(8+\ell+k)z^{6+\ell+k}}{(4+\ell)(4+k)(6+\ell+k)} - \frac{1}{3} \frac{(10+\ell+k)z^{6+\ell+k}}{(5+\ell)(5+k)(6+\ell+k)} + \frac{1}{4} \frac{z^{6+\ell+k}}{(3+\ell)(3+k)} \right\}.$$

Furthermore, by similar, elementary computations we can compute the expectation of  $\langle T_{z,\theta} f, e_{z,i} \rangle$ , showing that it is equal to

$$\mathbb{E} \langle T_{z,\theta} f, e_{z,i} \rangle = \sum_{\ell=0}^i a_{z,i,\ell} \left\{ \frac{1}{2} \frac{1}{\ell+3} z^{\ell+3} + \theta^2 \frac{1}{\ell+2} z^{\ell+2} \right\}.$$

Combining the above displays and simplifying the expression, we see that the covariance  $\text{Cov}(X_{z,i}, X_{z,j})$  equals

$$\sum_{\ell=0}^i \sum_{k=0}^j a_{1,i,\ell} a_{1,j,k} \left( z^5 \left\{ \frac{2(8+\ell+k)}{3(4+\ell)(4+k)(6+\ell+k)} - \frac{10+\ell+k}{3(5+\ell)(5+k)(6+\ell+k)} \right\} + \theta^2 z^4 \left\{ \frac{2}{(2+k)(3+\ell)} - \frac{2}{3(1+k)(4+\ell)} + \frac{4}{(1+k)(2+k)(3+k)(4+k)(5+k+\ell)} \right\} \right).$$

This concludes the proof of the claim.



## 4.9 Proof of Proposition 4.3.1

The main component of the proof is Lemma 4.9.1 below, which is modelled on Lemma 7.1 in (Lehmann 1983).

For  $\tilde{\Theta}_N = \sqrt{N}(\Theta - T_N)$  the rescaled parameter set, we write the posterior density as

$$\begin{aligned}\pi_N(t|X_N) &= \frac{\pi\left(T_N + \frac{t}{\sqrt{N}}\right) \exp\left(l^{(N)}\left(T_N + \frac{t}{\sqrt{N}}\right)\right)}{\int_{\tilde{\Theta}_N} \pi\left(T_N + \frac{u}{\sqrt{N}}\right) \exp\left(l^{(N)}\left(T_N + \frac{u}{\sqrt{N}}\right)\right) du} \\ &= \pi\left(T_N + \frac{t}{\sqrt{N}}\right) e^{\omega_N(t)} \frac{1}{C_N},\end{aligned}$$

where

$$\omega_N(t) = l^{(N)}\left(T_N + \frac{t}{\sqrt{N}} - l^{(N)}(\theta_N^*)\right) - \frac{1}{2N} \nabla l^{(N)}(\theta_N^*)^\top V_{\theta_N^*, N}^{-1} \nabla l^{(N)}(\theta_N^*), \quad (4.35)$$

$$C_N = \int_{\tilde{\Theta}_N} \pi\left(T_N + \frac{u}{\sqrt{N}}\right) e^{\omega_N(u)} du. \quad (4.36)$$

Since  $\int_{\mathbb{R}^d} e^{-\frac{1}{2}t^\top V_{\theta^*} t} dt = (2\pi)^{d/2} \det(V_{\theta^*})^{-1/2}$  and the sets  $\tilde{\Theta}_N$  grow to the full space by Assumption 4.1.1, Lemma 4.9.1 implies

$$C_N \xrightarrow{P_0^{(N)}} \pi(\theta^*) (2\pi)^{d/2} \det(V_{\theta^*})^{-1/2} > 0. \quad (4.37)$$

It follows that the sequence  $C_N^{-1}$  is bounded in probability.

We can rewrite Equation (4.6) as  $C_N^{-1}$  times

$$\begin{aligned}& \int_{\tilde{\Theta}_N} \left| \pi\left(T_N + \frac{t}{\sqrt{N}}\right) e^{\omega_N(t)} - C_N (2\pi)^{-d/2} \det(V_{\theta^*})^{1/2} e^{-\frac{1}{2}t^\top V_{\theta^*} t} \right| dt \\ & \leq \int_{\tilde{\Theta}_N} \left| \pi\left(T_N + \frac{t}{\sqrt{N}}\right) e^{\omega_N(t)} - e^{-\frac{1}{2}t^\top V_{\theta^*} t} \pi(\theta^*) \right| dt \\ & \quad + \left| C_N (2\pi)^{-d/2} \det(V_{\theta^*})^{1/2} - \pi(\theta^*) \right| \int_{\tilde{\Theta}_N} e^{-\frac{1}{2}t^\top V_{\theta^*} t} dt,\end{aligned}$$

by the triangle inequality. Both terms on the right converge to zero in probability, as follows from Lemma 4.9.1 and Equation (4.37), respectively. This concludes the proof of Equation (4.6).

For the proof of Equation (4.7), it is enough to show that the preceding line of argument goes through with an added factor  $(1 + |t|)$  in the integrands. This follows by an appropriately adapted version of Lemma 4.9.1.

**Lemma 4.9.1.** *Under the assumptions of Proposition 4.3.1, as  $N \rightarrow \infty$ ,*

$$\int_{\tilde{\Theta}_N} \left| e^{\omega_N(t)} \pi\left(T_N + \frac{t}{\sqrt{N}}\right) - e^{-\frac{1}{2}t^\top V_{\theta^*} t} \pi(\theta^*) \right| dt \xrightarrow{P_0^{(N)}} 0.$$

*Proof.* In view of the definitions of  $R_N(\theta)$  and  $\omega_N(t)$ , given in Equations (4.4) and (4.35),

$$\omega_N(t) = -\frac{1}{2}t^\top V_{\theta_N^*, N} t - \frac{1}{2N}(t + Z_N)^\top R_N\left(T_N + \frac{t}{\sqrt{N}}\right)(t + Z_N), \quad (4.38)$$

where  $Z_N := N^{-1/2}V_{\theta_N^*, N}^{-1}\nabla l^{(N)}(\theta_N^*)$ . By Assumptions 4.1.2 and 4.1.3, the sequence  $Z_N$  is bounded in probability. Definition (4.5) of  $T_N$  gives that  $T_N = \theta_N^* + N^{-1/2}Z_N$  and hence  $T_N \rightarrow \theta^*$  in probability, by Assumption 4.1.1. We now partition  $\tilde{\Theta}_N$  into three subsets: for a fixed  $M$  and  $\eta$ , we consider the sets of  $t \in \tilde{\Theta}_N$  with  $\|t\| \leq M$ ,  $M < \|t\| < \eta\sqrt{N}$  and  $\|t\| \geq \eta\sqrt{N}$ .

Because the volume of the ball with radius  $M$  is finite, for the range  $\|t\| \leq M$ , it suffices to show that

$$\sup_{\|t\| \leq M} \left| e^{\omega_N(t)} \pi\left(T_N + \frac{t}{\sqrt{N}}\right) - e^{-\frac{1}{2}t^\top V_{\theta^*} t} \pi(\theta_N^*) \right| = o_P(1). \quad (4.39)$$

In view of Assumption 4.1.5, we have that  $\sup_{\|t\| \leq M} \|N^{-1}R_N(T_N + t/\sqrt{N})\| \rightarrow 0$  in probability. It follows that the supremum over  $\|t\| \leq M$  of the second term on the right of Equation (4.38) tends to zero in probability. Combined with Assumption 4.1.2 this gives that  $\sup_{\|t\| \leq M} |\omega_N(t) + \frac{1}{2}t^\top V_{\theta^*} t| \rightarrow 0$ . Since also  $\sup_{\|t\| \leq M} |\pi(T_N + \frac{t}{\sqrt{N}}) - \pi(\theta^*)| \rightarrow 0$  by Assumption 4.1.6, it follows that Equation (4.39) holds. This is true for every fixed  $M$ .

Next, we deal with the range  $M < \|t\| < \eta\sqrt{N}$ . The function defined by  $t \mapsto \exp(-\frac{1}{2}t^\top V_{\theta^*} t)$  is integrable and can be made arbitrarily small by choosing large enough  $M$ . We shall show that

$$\int_{M < \|t\| < \eta\sqrt{N}} \pi(T_N + t/\sqrt{N}) e^{\omega_N(t)} dt \quad (4.40)$$

can be made arbitrarily small by choosing sufficiently large  $M$  and sufficiently small  $\eta > 0$ . On the range of the integral, we have  $\|T_N + t/\sqrt{N} - \theta^*\| < 2\eta$ , with probability tending to one. Therefore, by Assumption 4.1.6, we have for sufficiently small  $\eta > 0$  that  $\pi(T_N + \frac{t}{\sqrt{N}})$  is bounded in probability. By Assumption 4.1.5 we can choose  $\eta > 0$  still smaller, if necessary, so that  $\sup_{\|t\| < \eta\sqrt{N}} \left\| \frac{1}{N} R_N(T_N + \frac{t}{\sqrt{N}}) \right\| < \frac{1}{4} \lambda_{\min, V_{\theta^*}}$ , with probability tending to one, where  $\lambda_{\min, V_{\theta^*}} > 0$  denotes the smallest eigenvalue of  $V_{\theta^*}$ . Assumption 4.1.2 gives that  $\|V_{\theta_N^*, N} - V_{\theta^*}\| < \frac{1}{4} \lambda_{\min, V_{\theta^*}}$ , for large enough  $N$ . Then, in view of Equation (4.38),

$$\begin{aligned} \omega_N(t) &\leq -\frac{1}{2} \lambda_{\min, V_{\theta^*}} \|t\|^2 + \frac{1}{2} \|V_{\theta_N^*, N} - V_{\theta^*}\| \|t\|^2 + \frac{1}{4} \lambda_{\min, V_{\theta^*}} (\|t\|^2 + \|Z_N\|^2) \\ &\leq -\frac{1}{8} \lambda_{\min, V_{\theta^*}} \|t\|^2 + \frac{1}{4} \lambda_{\min, V_{\theta^*}} \|Z_N\|^2, \end{aligned}$$

with probability tending to one. The second term is bounded in probability, while the first term is quadratically decreasing, and hence dominates the expression. We conclude that Equation (4.40) can be made arbitrarily small by choosing sufficiently large  $M$ , for  $\eta$  sufficiently small so that the preceding estimates hold.

Finally, we consider the range  $\|t\| > \eta\sqrt{N}$ . As before, the contribution of the term involving  $\exp(-\frac{1}{2} t^\top V_{\theta^*} t)$  tends to zero, and hence we need only deal with the term  $\pi(T_N + t/\sqrt{N}) e^{\omega_N(t)}$ . In view of Assumption 4.1.4, there exists an  $\epsilon > 0$  such that  $\sup_{\|\theta - \theta_N^*\| > \eta} l^{(N)}(\theta) - l^{(N)}(\theta_N^*) \leq -\epsilon N$  with probability going to one. On this event,

$$\begin{aligned} &\int_{\|t\| > \eta\sqrt{N}} \left| e^{\omega_N(t)} \pi\left(T_N + \frac{t}{\sqrt{N}}\right) \right| dt \\ &\leq N^{d/2} e^{-n\epsilon - \frac{1}{2N} \nabla l^{(N)}(\theta_N^*)^\top V_{\theta_N^*, N}^{-1} \nabla l^{(N)}(\theta_N^*)} \int_{\|\theta - T_N\| \geq \eta} \pi(\theta) d\theta. \end{aligned}$$

This tends to zero in probability. This finishes the proof of the lemma.  $\square$

# Chapter 5

## Conclusion

To this day, statistical inverse problems pose a challenging problem for mathematicians. Although they frequently appear in many different scientific fields, there is no unifying theory that encompasses the many different types that exist. This thesis attempts to expand the applicability of Bayesian methods to a larger class of problems.

An exception to this is the set of linear inverse problems, for which the theory is substantially more developed. Although they are only a small subset of all the relevant inverse problems, they do provide a good starting point for discovering new results. Due to the conjugacy of Gaussian priors in white noise models, the involved analysis is more tractable.

The Gaussian structure that appears in linear inverse problems with Gaussian priors immediately breaks down when one considers non-linear inverse problems. Therefore the analysis of the statistical model is significantly more involved. Recently advancements in general methods for non-linear inverse problems in PDE models have been made, but the techniques still require a detailed analysis of the PDE in question.

Many existing Bayesian methods for solving inverse problems require an extensive computational load. In Gaussian models, this often appears due to the need for inverting a large matrix. Another source for the computational complexity can be the forward map in a PDE model, which is often computationally difficult to compute.

In this thesis, we proposed a method for applying techniques for linear inverse problems to solve non-linear problems. It has been proven that this method, which utilises Gaussian priors on the singular value decomposition of a linear operator, is consistent at minimax rates and provides good uncertainty quantification. This technique proves to be particularly useful for PDE models, as

it can provide benefits in the mathematical analysis, do exact posterior inference, and speed up the generation of draws from the posterior distribution.

With this method, the mathematical analysis is now focussed on showing the existence of appropriate observational models and corresponding inverse maps. The original observations may need to be written as the observations coming from an equivalent linear inverse problem for a new parameter. Furthermore, this new parameter needs to permit an inverse map that recovers the parameter of interest in a suitable way. Proving the existence of these models and inverse maps is not trivial and might even be impossible in certain cases.

Exact inference on the posterior is often impossible to perform in non-linear PDE models as the likelihood does not admit to a tractable form. Therefore, methods such as MCMC are often used to generate approximate posterior draws. The need for the calibration of the parameters for the MCMC amplifies the difficulties that arise in implementing these techniques. The method proposed in this thesis allows to generate exact draws from the posterior in the provided examples, due to the explicit expression for the Gaussian posterior and the inverse map.

The improvement in computational speed can be achieved through the means of distributed methods. Although these have not yet been proven directly for non-linear models, they can be attained by applying the results in this thesis for linear models to the non-linear problem. Theoretically, when there are  $N$  observations and  $o(N^{\frac{1}{3}})$  machines are working in parallel, the computational time can be reduced by an order of  $o(N)$ . For inverse problems, which naturally require a high number of observations, this optimisation substantially reduces the computational time and can be proven to be highly useful in practice.

Whether it is possible to apply the proposed method, depends on the structure of the PDE and which parameter within the PDE needs to be recovered. It remains an open question whether this can be applied to other PDE models such as the multi-dimensional Darcy's problem. Furthermore, the model is restricted to recovering functions that are an element of the Sobolev space generated by eigenfunctions defined by the linear operator. More analysis is needed to see if these methods work for more general priors that can recover functions from a larger space.

In this thesis, new conditions for the Bernstein-von Mises theorem were proven for hierarchical, finite-dimensional models that are potentially misspecified. These results expand the models in which it is known that the posterior asymptotically converges to a normal distribution. More concrete analysis has been done for certain Gaussian misspecifications, for which it is known that the misspecified posterior is consistent. However, one does not know whether the

---

uncertainty quantification asymptotically behaves well, or whether it is over or under-confident. Therefore it is paramount that one is careful with the inference coming from such misspecified posterior distributions. Misspecification of models with a likelihood that is intractable, or computationally expensive to evaluate, can be of great use. This provides a way of doing inference on complex models for which it would be impossible to do so otherwise. The resulting statistical analysis can be useful in these models, but a careful interpretation of the conclusions remains necessary.



# Bibliography

- Agapiou, S., Larsson, S. and Stuart, A. M. (2013). ‘Posterior contraction rates for the Bayesian approach to linear ill-posed inverse problems’. In: *Stochastic Process. Appl.* 123.10, pp. 3828–3860. DOI: 10.1016/j.spa.2013.05.001.
- Agapiou, S., Stuart, A. M. and Zhang, Y.-X. (2014). ‘Bayesian Posterior Contraction Rates for Linear Severely Ill-posed Inverse Problems’. In: *Journal of Inverse and Ill-posed Problems* 22.3, pp. 297–321. DOI: 10.1515/jip-2012-0071.
- Alsing, J. et al. (2016). ‘Hierarchical cosmic shear power spectrum inference’. In: *Monthly Notices of the Royal Astronomical Society* 455.4, pp. 4452–4466. DOI: 10.1093/mnras/stv2501.
- Bayarri, M. J. and Berger, J. O. (2004). ‘The Interplay of Bayesian and Frequentist Analysis’. In: *Statist. Sci.* 19.1, pp. 58–80. DOI: 10.1214/088342304000000116.
- Bayes, T. (1763). ‘LII. An essay towards solving a problem in the doctrine of chances. By the late Rev. Mr. Bayes, F. R. S. communicated by Mr. Price, in a letter to John Canton, A. M. F. R. S’. In: *Philos. Trans. R. Soc.* 53, pp. 370–418. DOI: 10.1098/rstl.1763.0053.
- Belitser, E. and Ghosal, S. (2003). ‘Adaptive Bayesian Inference on the Mean of an Infinite-Dimensional Normal Distribution’. In: *Ann. Statist.* 31.2, pp. 536–559. DOI: 10.1214/aos/1051027880.
- Berger, J. O. (1980). *Statistical decision theory and Bayesian analysis*. Springer-Verlag New York Inc. ISBN: 0-387-96098-8. DOI: 10.1007/978-1-4757-4286-2.
- Bernstein, S. (1917). *Theory of Probability*. Russian.
- Bochkina, N. (2022). *Bernstein - von Mises theorem and misspecified models: a review*. arXiv: 2204.13614 [math.ST].



- Bowman, F. D. et al. (2008). 'A Bayesian hierarchical framework for spatial modeling of fMRI data'. In: *NeuroImage* 39.1, pp. 146–156. DOI: 10.1016/j.neuroimage.2007.08.012.
- Brooks, S. et al. (2011). *Handbook of Markov Chain Monte Carlo*. Chapman and Hall/CRC Handbooks of Modern Statistical Methods Ser. CRC Press LLC. ISBN: 9781420079425. DOI: 10.1201/b10905.
- Cai, T. T. and Low, M. G. (2004). 'An adaptation theory for nonparametric confidence intervals'. In: *Ann. Statist.* 32.5, pp. 1805–1840. DOI: 10.1214/009053604000000049.
- Le Cam, L. (2012). *Asymptotic Methods in Statistical Decision Theory*. Springer Science & Business Media. ISBN: 978-1-4612-4946-7. DOI: 10.1007/978-1-4612-4946-7.
- Le Cam, L. and Yang, G. (1990). *Asymptotics in statistics: some basic concepts*. Springer-Verlag. ISBN: 978-1-4684-0377-0. DOI: 10.1007/978-1-4684-0377-0.
- Castillo, I. (2008). 'Lower bounds for posterior rates with Gaussian process priors'. In: *Electron. J. Statist.* 2, pp. 1281–1299. DOI: 10.1214/08-EJS273.
- Castillo, I. (2012). 'A semiparametric Bernstein–von Mises theorem for Gaussian process priors'. In: *Probab. Theory Related Fields* 152.1-2, pp. 53–99. ISSN: 0178-8051. DOI: 10.1007/s00440-010-0316-5.
- Castillo, I. and Nickl, R. (2014). 'On the Bernstein–von Mises phenomenon for nonparametric Bayes procedures'. In: *Ann. Statist.* 42.5, pp. 1941–1969. ISSN: 0090-5364. DOI: 10.1214/14-AOS1246.
- Castillo, I. and Rousseau, J. (2015). 'A Bernstein–von Mises theorem for smooth functionals in semiparametric models'. In: *Ann. Statist.* 43.6, pp. 2353–2383. ISSN: 0090-5364. DOI: 10.1214/15-AOS1336.
- Cavalier, L. (2008). 'Nonparametric statistical inverse problems'. In: *Inverse Problems* 24.3, pp. 379–452. DOI: 10.1088/0266-5611/24/3/034004.
- Cornish, N. J. and Littenberg, T. B. (2015). 'Bayeswave: Bayesian inference for gravitational wave bursts and instrument glitches'. In: *Classical and Quantum Gravity* 32.13, p. 135012. DOI: 10.1088/0264-9381/32/13/135012.
- Cox, D. D. (1993). 'An Analysis of Bayesian Inference for Nonparametric Regression'. In: *Ann. Statist.* 21.2, pp. 903–923. DOI: 10.1214/aos/1176349157.

- Donnet, S. et al. (2018). ‘Posterior concentration rates for empirical Bayes procedures with applications to Dirichlet process mixtures’. In: *Bernoulli* 24.1, pp. 231–256. DOI: 10.3150/16-BEJ872.
- Doob, J. L. (1955). ‘A probability approach to the heat equation’. In: *Trans. Amer. Math. Soc.* 80.1, pp. 216–280. DOI: 10.2307/1993013.
- Doob, J. (1949). ‘Application of the theory of martingales’. In: *Le Calcul des probabilités et ses applications* 13, pp. 23–27.
- Durante, D., Pozza, F. and Szabó, B. (2023). *Skewed Bernstein-von Mises theorem and skew-modal approximations*. arXiv: 2301.03038 [math.ST].
- Engl, H. W., Hanke, M. and Neubauer, A. (1996). *Regularization of inverse problems*. Vol. 375. Mathematics and its Applications. Kluwer Academic Publishers Group, Dordrecht, pp. viii+321. ISBN: 0-7923-4157-0.
- Eriksen, H. K. et al. (2004). ‘Power spectrum estimation from high-resolution maps by Gibbs sampling’. In: *The Astrophysical Journal Supplement Series* 155.2, p. 227. DOI: 10.1086/425219.
- Evans, L. C. (2010). *Partial differential equations*. Second. Vol. 19. Graduate Studies in Mathematics. American Mathematical Society, Providence, RI, pp. xxii+749. ISBN: 978-0-8218-4974-3. DOI: 10.1090/gsm/019.
- Evans, S. N. and Stark, P. B. (2002). ‘Inverse problems as statistics’. In: *Inverse Problems* 18.4, R55–R97. DOI: 10.1088/0266-5611/18/4/201.
- Ezquiaga, J. and Zumalacárregui, M. (2018). ‘Dark energy in light of multimessenger gravitational-wave astronomy’. In: *Frontiers in Astronomy and Space Sciences* 5, p. 44. DOI: 10.3389/fspas.2018.00044.
- Fearnhead, P. et al. (2016). *Piecewise Deterministic Markov Processes for Continuous-Time Monte Carlo*. arXiv: 1611.07873 [stat.CO].
- Feehan, P. M. N., Gong, R. and Song, J. (2015). *Feynman-Kac Formulas for Solutions to Degenerate Elliptic and Parabolic Boundary-Value and Obstacle Problems with Dirichlet Boundary Conditions*. arXiv: 1509.03864 [math.PR].
- de Finetti, B. (1931). ‘Funzione caratteristica di un fenomeno aleatorio’. In: *Matematiche e Naturale* 4, pp. 251–299.
- Fisher, R. (1925). ‘Theory of Statistical Estimation’. In: *Math. Proc. Cambridge Philos. Soc.* 22, pp. 700–725. DOI: 10.1017/S0305004100009580.

- Florens, J.-P. and Simoni, A. (2012). ‘Regularized Posteriors in Linear Ill-Posed Inverse Problems’. In: *Scand. J. Stat.* 39.2, pp. 214–235. DOI: 10.1111/j.1467-9469.2011.00784.x.
- Freedman, D. A. (1963). ‘On the Asymptotic Behavior of Bayes’ Estimates in the Discrete Case’. In: *Ann. Math. Statist.* 34.4, pp. 1386–1403. DOI: 10.1214/aoms/1177703871.
- Freedman, D. A. (1965). ‘On the Asymptotic Behavior of Bayes’ Estimates in the Discrete Case II’. In: *Ann. Math. Statist.* 36.2, pp. 454–456. DOI: 10.1214/aoms/1177700155.
- Gelman, A. et al. (2013). *Bayesian Data Analysis*. Third. CRC Press. ISBN: 978-1-4398-9822-2. DOI: 10.1201/b16018.
- Ghosal, S., Ghosh, J. K. and van der Vaart, A. W. (2000). ‘Convergence rates of posterior distributions’. In: *Ann. Statist.* 28.2, pp. 500–531. ISSN: 0090-5364. DOI: 10.1214/aos/1016218228.
- Ghosal, S. and van der Vaart, A. W. (2017). *Fundamentals of Nonparametric Bayesian Inference*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, Cambridge. DOI: 10.1017/9781139029834.
- Gilbarg, D. and Trudinger, N. S. (2015). *Elliptic Partial Differential Equations of Second Order*. eng. Vol. 224. Grundlehren der mathematischen Wissenschaften. Berlin, Heidelberg: Springer Berlin / Heidelberg. ISBN: 354013025X. DOI: 10.1007/978-3-642-61798-0.
- Giné, E. and Nickl, R. (2016). *Mathematical Foundations of Infinite-Dimensional Statistical Models*. Cambridge series in statistical and probabilistic mathematics. Cambridge University Press, Cambridge. ISBN: 978-1-107-04316-9. DOI: 10.1017/CBO9781107337862.
- Giordano, M. and Kekkonen, H. (2020). ‘Bernstein-von Mises Theorems and Uncertainty Quantification for Linear Inverse Problems’. In: *SIAM/ASA J. Uncertain. Quantification* 8.1, pp. 342–373. DOI: 10.1137/18M1226269.
- Giordano, M. and Nickl, R. (2020). ‘Consistency of Bayesian inference with Gaussian process priors in an elliptic inverse problem’. In: *Inverse Problems* 36.8, pp. 085001, 35. ISSN: 0266-5611. DOI: 10.1088/1361-6420/ab7d2a.
- Gugushvili, S., van der Vaart, A. W. and Yan, D. (2020). ‘Bayesian linear inverse problems in regularity scales’. In: *Ann. Inst. Henri Poincaré Probab. Statist.* 56.3, pp. 2081–2107. DOI: 10.1214/19-AIHP1029.

- Hadji, A. (2023). ‘Scalability and uncertainty of Gaussian processes’. PhD thesis. Leiden University. URL: <https://hdl.handle.net/1887/3513272>.
- Hadji, A., Hesselink, T. and Szabó, B. (2022). *Optimal recovery and uncertainty quantification for distributed Gaussian process regression*. arXiv: 2205.03150 [math.ST].
- Hjort, N. L. et al. (2010). *Bayesian Nonparametrics*. Cambridge University Press. ISBN: 9780511672118. DOI: 10.1017/CBO9780511802478.
- Ibragimov, I. A. and Khas’minskii, R. Z. (1985). ‘On Nonparametric Estimation of the Value of a Linear Functional in Gaussian White Noise’. In: *Theory Probab. Appl.* 29.1, pp. 1281–1299. DOI: 10.1137/1129002.
- Isakov, V. (2017). *Inverse Problems for Partial Differential Equations*. Springer Cham. ISBN: 978-3-319-51658-5. DOI: 10.1007/978-3-319-51658-5.
- Ivrii, V. (2016). ‘100 years of Weyl’s law’. In: *Bull. Math. Sci.* 6, pp. 379–452. DOI: 10.1007/s13373-016-0089-y.
- Kekkonen, H. (2022). ‘Consistency of Bayesian inference with Gaussian process priors for a parabolic inverse problem’. In: *Inverse Problems* 38.3, p. 035002. DOI: 10.1088/1361-6420/ac4839.
- Kimeldorf, G. S. and Wahba, G. (1970). ‘A Correspondence Between Bayesian Estimation on Stochastic Processes and Smoothing by Splines’. In: *Ann. Math. Statist.* 41.2, pp. 495–502. DOI: 10.1214/aoms/1177697089.
- Kinderlehrer, D. and Stampacchia, G. (2000). *An Introduction to Variational Inequalities and Their Applications*. First. Vol. 31. Classics in Applied Mathematics. Society for Industrial and Applied Mathematics, pp. xxii+306. ISBN: 0-89871-466-4. DOI: 10.1137/1.9780898719451.
- Kleijn, B. J. K. and van der Vaart, A. W. (2012). ‘The Bernstein-Von-Mises theorem under misspecification’. In: *Electron. J. Stat.* 6, pp. 354–381. DOI: 10.1214/12-EJS675.
- Knapik, B. T., van der Vaart, A. W. and van Zanten, J. H. (2013). ‘Bayesian recovery of the initial condition for the heat equation’. In: *Comm. Statist. Theory Methods* 42.7, pp. 1294–1313. ISSN: 0007-4497. DOI: 10.1080/03610926.2012.681417.

- Knapik, B. T., Szabó, B. T. et al. (2016). ‘Bayes procedures for adaptive inference in inverse problems for the white noise model’. In: *Probab. Theory Related Fields* 164.3, pp. 771–813. DOI: 10.1007/s00440-015-0619-7.
- Knapik, B. T., van der Vaart, A. W. and van Zanten, J. H. (2011). ‘Bayesian Inverse Problems with Gaussian Priors’. In: *Ann. Statist.* 39.5, pp. 2626–2657. DOI: 10.1214/11-AOS920.
- Laplace, P.-S. (1810). ‘Mémoire sur les formules qui sont fonctions de très grands nombres et sur leurs applications aux probabilités’. French. In: *Oeuvres de Laplace* 12, pp. 301–345.
- Leahu, H. (2011). ‘On the Bernstein-von Mises phenomenon in the Gaussian white noise model’. In: *Electronic Journal of Statistics* 5.none, pp. 373–404. DOI: 10.1214/11-EJS611.
- Ledoux, M. and Talagrand, M. (1991). *Probability in Banach spaces*. Vol. 23. Berlin: Springer-Verlag, pp. xii+480. ISBN: 3-540-52013-9. DOI: 10.1007/978-3-642-20212-4.
- Lehmann, E. L. (1983). *Theory of Point Estimation*. First. Probability and Mathematical Statistics. John Wiley & Sons, pp. xii+506. ISBN: 0-471-05849-1.
- Lions, J.-L. and Magenes, E. (1972a). *Non-Homogeneous Boundary Value Problems and Applications : Volume I*. Die Grundlehren der mathematischen Wissenschaften. Springer Berlin, Heidelberg. ISBN: 978-3-642-65161-8. DOI: 10.1007/978-3-642-65161-8.
- Lions, J.-L. and Magenes, E. (1972b). *Non-Homogeneous Boundary Value Problems and Applications : Volume II*. Die Grundlehren der mathematischen Wissenschaften. Berlin, Heidelberg: Springer Berlin, Heidelberg. ISBN: 978-3-642-65217-2. DOI: 10.1007/978-3-642-65217-2.
- Magnus, J. R. and Neudecker, H. (2019). *Matrix Differential Calculus with Applications in Statistics and Econometrics*. Third. Probability and Statistics. John Wiley & Sons, pp. xviii+479. ISBN: 9781119541202. DOI: 10.1002/9781119541219.
- Marinucci, D. (2004). ‘Testing for non-Gaussianity on cosmic microwave background radiation: A review’. In: *Statist. Sci.* 19.2, pp. 294–307. DOI: 10.1214/088342304000000783.
- Maz’ya, V. (2011). *Sobolev Spaces: with Applications to Elliptic Partial Differential Equations*. Springer Berlin, Heidelberg. ISBN: 978-3-642-15564-2. DOI: 10.1007/978-3-642-15564-2.

- von Mises, R. (1931). *Wahrscheinlichkeitsrechnung*. German. Springer-Verlag.
- Moler, C. B. and Stewart, G. W. (1973). ‘An Algorithm for Generalized Matrix Eigenvalue Problems’. In: *SIAM J. Numer. Anal.* 10.2, pp. 241–256. DOI: 10.1137/0710024.
- Monard, F., Nickl, R. and Paternain, G. P. (2021a). ‘Statistical guarantees for Bayesian uncertainty quantification in nonlinear inverse problems with Gaussian process priors’. In: *Ann. Statist.* 49.6, pp. 3255–3298. ISSN: 0090-5364. DOI: 10.1214/21-aos2082.
- Monard, F., Nickl, R. and Paternain, G. P. (2021b). ‘Consistent inversion of noisy non-Abelian X-ray transforms’. In: *Comm. Pure Appl. Math.* 74.5, pp. 1045–1099. ISSN: 0010-3640. DOI: 10.1002/cpa.21942.
- Monthus, C. and Garel, T. (2011). ‘A critical Dyson hierarchical model for the Anderson localization transition’. In: *Journal of Statistical Mechanics: Theory and Experiment* 2011.05, P05005. DOI: 10.1088/1742-5468/2011/05/P05005.
- Di Nezza, E., Palatucci, G. and Valdinoci, E. (2012). ‘Hitchhiker’s guide to the fractional Sobolev spaces’. In: *Bulletin des Sciences Mathématiques* 136.5, pp. 521–573. ISSN: 0007-4497. DOI: 10.1016/j.bulsci.2011.12.004.
- Nickl, R. (2018). ‘Bernstein-von Mises Theorems for statistical inverse problems I: Schrödinger Equation’. In: *J. Eur. Math. Soc* 22.8, pp. 2697–2750. DOI: 10.4171/JEMS/975.
- Nickl, R. (2023). *Bayesian Non-Linear Statistical Inverse Problems*.
- Nickl, R., van de Geer, S. and Wang, S. (2020). ‘Convergence rates for penalized least squares estimators in PDE constrained regression problems’. In: *SIAM/ASA J. Uncertain. Quantification* 8.1, pp. 374–413. DOI: 10.1137/18M1236137.
- Nickl, R. and Wang, S. (2020). *On polynomial-time computation of high-dimensional posterior measures by Langevin-type algorithms*. arXiv: 2009.05298 [math.ST].
- Oppen, M. and Vivaerlli, F. (1999). ‘General bounds on bayes errors for regression with gaussian processes’. In: *Advances in Neural Information Processing Systems* 11, pp. 302–308.
- Petrone, S. et al. (2014). ‘Empirical Bayes methods in classical and Bayesian inference’. In: *METRON* 72, pp. 201–215. DOI: 10.1007/s40300-014-0044-1.

- Pinsker, M. S. (1980). 'Optimal Filtering of Square-Integrable Signals in Gaussian Noise'. In: *Probl. Peredachi Inf.* 16.2, pp. 52–68.
- Ray, K. (2013). 'Bayesian inverse problems with non-conjugate priors'. In: *Electron. J. Statist.* 7, pp. 2516–2549. DOI: 10.1214/13-EJS851.
- Ray, K. (2017). 'Adaptive Bernstein–von Mises theorems in Gaussian white noise'. In: *Ann. Statist.* 45.6, pp. 2511–2536. ISSN: 0090-5364. DOI: 10.1214/16-AOS1533.
- Ray, K. and Schmidt-Hieber, J. (2016). 'Minimax theory for a class of nonlinear statistical inverse problems'. In: *Inverse Problems* 32.6, p. 065003. DOI: 10.1088/0266-5611/32/6/065003.
- Reiß, M. (2008). 'Asymptotic Equivalence for Nonparametric Regression with Multivariate and Random Design'. In: *Ann. Statist.* 36.4, pp. 1957–1982. ISSN: 0090-5364. DOI: 10.1214/07-AOS525.
- Robert, C. and Casella, G. (1999). *Monte Carlo Statistical Methods*. Springer Texts in Statistics. Springer New York, NY. ISBN: 9781475730715. DOI: 10.1007/978-1-4757-4145-2.
- Robins, J. and van der Vaart, A. (2006). 'Adaptive nonparametric confidence sets'. In: *Ann. Statist.* 34.1, pp. 229–253. DOI: 10.1214/009053605000000877.
- Rousseau, J. and Szabo, B. (2017). 'Asymptotic behaviour of the empirical Bayes posteriors associated to maximum marginal likelihood estimator'. In: *Ann. Statist.* 45.2, pp. 833–865. DOI: 10.1214/16-AOS1469.
- Schwartz, L. (1965). 'On Bayes Procedure'. In: *Z. Wahrscheinlichkeitstheorie* 4.1, pp. 10–26. DOI: 10.1007/BF00535479.
- Scott, S. L. (2017). 'Comparing consensus Monte Carlo strategies for distributed Bayesian computation'. In: *Braz. J. Probab. Stat.* 31.4, pp. 668–685. DOI: 10.1214/17-BJPS365.
- Scott, S. L. et al. (2016). 'Bayes and big data: the consensus Monte Carlo algorithm'. In: *International Journal of Management Science and Engineering Management* 11, pp. 78–88. DOI: 10.1080/17509653.2016.1142191.
- Seeley, R. T. (1965). 'Integro-Differential Operators on Vector Bundles'. In: *Trans. Amer. Math. Soc.* 117, pp. 167–204. DOI: 10.2307/1994203.
- Sellentin, E., Heymans, C. and Harnoïs-Déraps, J. (2018). 'The skewed weak lensing likelihood: why biases arise, despite data and theory being sound'.

- In: *Monthly Notices of the Royal Astronomical Society* 477.4, pp. 4879–4895. DOI: 10.1093/mnras/sty988.
- Shang, Z. and Cheng, G. (2015). *A Bayesian Splitotic Theory For Nonparametric Models*. arXiv: 1508.04175 [math.ST].
- Sniekers, S. and van der Vaart, A. (2015). ‘Adaptive Bayesian credible sets in regression with a Gaussian process prior’. In: *Electron. J. Stat.* 9.2, pp. 2475–2527. DOI: 10.1214/15-EJS1078.
- Sniekers, S. and van der Vaart, A. (2020). ‘Adaptive Bayesian credible bands in regression with a Gaussian process prior’. In: *Sankhya A* 82.2, pp. 386–425. ISSN: 0976-836X. DOI: 10.1007/s13171-019-00185-0.
- Stigler, S. (1999). *Statistics on the table: the history of statistical concepts and methods*. Harvard University Press. ISBN: 0-674-83601-4. DOI: 10.2307/j.ctv1pdrpsj.
- Stigler, S. (2007). ‘The Epic Story of Maximum Likelihood’. In: *Statist. Sci.* 22.4, pp. 598–620. DOI: 10.1214/07-STS249.
- Stuart, A. M. (2010). ‘Inverse problems: A Bayesian perspective’. In: *Acta Numerica* 19, pp. 451–559. DOI: 10.1017/S0962492910000061.
- Szabó, B. T., van der Vaart, A. W. and van Zanten, J. H. (2013). ‘Empirical Bayes scaling of Gaussian priors in the white noise model’. In: *Electron. J. Statist.* 7, pp. 991–1018. DOI: 10.1214/13-EJS798.
- Szabó, B., van der Vaart, A. W. and van Zanten, J. H. (2015). ‘Frequentist coverage of adaptive nonparametric Bayesian credible sets’. In: *Ann. Statist.* 43.4, pp. 1391–1428. ISSN: 0090-5364. DOI: 10.1214/14-AOS1270.
- Szabó, B. and van Zanten, H. (2019). ‘An asymptotic analysis of distributed nonparametric methods’. In: *Journal of Machine Learning Research* 20, pp. 1–30.
- Taylor, M. (2011). *Partial Differential Equations I: Basic Theory*. Applied Mathematical Sciences. Springer-Verlag, New York, pp. xxii+654. ISBN: 978-1-4419-7054-1. DOI: 10.1007/978-1-4419-7055-8.
- Triebel, H. (2008). *Function Spaces and Wavelets on Domains*. Second. EMS Tracts in Mathematics. European Mathematical Society Publishing House, pp. X+256. ISBN: 978-3-03719-019-7. DOI: 10.4171/019.



- Tröltzsch, F. (2010). *Optimal Control of Partial Differential Equations: Theory, Methods and Applications*. Vol. 112. Graduate Studies in Mathematics. American Mathematical Society. ISBN: 978-0-8218-4904-0. DOI: 10.1090/gsm/112.
- Tsybakov, A. B. (2008). *Introduction to Nonparametric Estimation*. Springer New York, NY. ISBN: 978-0-387-79052-7. DOI: 10.1007/b13794.
- van der Vaart, A. W. and van Zanten, J. H. (2008). ‘Rates of contraction of posterior distributions based on Gaussian process priors’. In: *Ann. Statist.* 36.3, pp. 1435–1463. DOI: 10.1214/009053607000000613.
- van der Vaart, A. W. (1998). *Asymptotic Statistics*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, Cambridge. ISBN: 978-0-521-49603-2. DOI: 10.1017/CBO9780511802256.
- van der Vaart, A. W. and Wellner, J. A. (2023). *Weak convergence and empirical processes*. Second Edition. Springer Series in Statistics. With applications to statistics. Springer-Verlag, New York.
- Wasserman, L. (2006). *All of Nonparametric Statistics*. Springer Texts in Statistics. Springer New York, NY. ISBN: 978-0-387-30623-0. DOI: 10.1007/0-387-30623-4.
- Wikle, C. K. (2003). ‘Hierarchical Bayesian models for predicting the spread of ecological processes’. In: *Ecology* 84.6, pp. 1382–1394. DOI: 10.1890/0012-9658(2003)084[1382:HBMFPT]2.0.CO;2.
- Woodhouse, J. H. and Dziewonski, A. M. (1984). ‘Mapping the upper mantle: Three-dimensional modeling of earth structure by inversion of seismic waveforms’. In: *J. Geophys. Res.* 89, pp. 5953–5986. DOI: 10.1029/JB089iB07p05953.
- Yan, D. and van der Vaart, A. (2023). ‘Bayesian Linear Inverse Problems in Regularity Scales with Discrete Observation’.
- Yan, D. (2020). ‘Bayesian Inference for Gaussian Models’. PhD thesis. Leiden University. URL: <https://hdl.handle.net/1887/86070>.
- Yang, Y., Bhattacharya, A. and Pati, D. (2017). *Frequentist coverage and sup-norm convergence rate in Gaussian process regression*. arXiv: 1708.04753 [math.ST].
- Zhao, Z. and Chung, K. L. (2012). *From Brownian Motion to Schrodinger’s Equation*. eng. Vol. 312. Grundlehren der mathematischen Wissenschaften. Springer. ISBN: 978-3-642-63381-2. DOI: 10.1007/978-3-642-57856-4.

- Zou, Y. and Guo, Z. (2003). 'A review of electrical impedance techniques for breast cancer detection'. In: *Medical Engineering & Physics* 25.2, pp. 79–90. doi: 10.1016/S1350-4533(02)00194-7.



# Acknowledgements

Before I thank all the people who have helped me write this manuscript, I would like to say that even though the front of the thesis states my name, this was far from a solo effort. All the results in this book are an accumulation of many hours of work, which would not have been possible if not for the support of many people.

Years of research would have been a fruitless effort if not for my supervisors Aad van der Vaart and Botond Szabó. Every time the progress halted, which happened rather frequently, the guidance and insight of both of you were incredibly helpful in making progress.

Seeing a path out once I was hopelessly stuck in the middle of a proof is a great skill Aad, and it was crucial in many occasions. In many theorems, the corresponding proofs followed the exact path that you outlined in the beginning, most of which I only started beginning to comprehend in the end. After more than four years of working in Bayesian statistics, it seems that I only explored the smallest part of an enormous ocean. Nonetheless, you always seemed to be able to navigate through the rough waters that sometimes arose.

Since meeting you Botond, I have always admired your optimism and dedication to mathematics. Time and time again you have helped me in so many different areas of academia, and the journey would not have been the same without your positive outlook. As you always made time to look into my results, it has given all my work the polishing that it so desperately required. The new ideas you always proposed on top of that, have made the projects much more diverse, extensive, and interesting than they would have been otherwise.

In both Leiden and in Delft I have met many wonderful people, even if I was not always around. Such an environment was the result of a great group of colleagues who made work, conferences, and reading groups much more pleasant. To the second group of fellow PhD's and friends, with whom I share all of my life in mathematics, I wish to extend similar gratitude. I can confidently say that sharing the experience and sorrow of this life with you made it somewhat

bearable at times.

Crediting my parents for everything they have done for me would simply be impossible, as you have made everything that I could have ever wished for possible. Sincere thanks seem to be only a negligible return of gratitude for creating the wonderful life that you have given me.

I would like to say the same to my brother, who at times I might have punished by being interested in maths, but who supported me nonetheless.

Since meeting you Nhi, you have helped me feel more positive and passionate about my thesis even when I thought I felt indifferent. For all the harsh words you so often spoke to me I want to thank you, as they were very helpful in finding the motivation to continue in difficult times. Until the very end of writing this thesis, you have pushed me to continue, even when it came at the cost of your own sanity. Now is the time to finally take a break.

# Curriculum Vitae

Geerten Jens Koers was born on the 30th of March, 1995 in Groningen, the Netherlands. After obtaining a high school degree at Vrije School Zutphen VO in 2013, he started his bachelors degree at Universiteit Leiden (LEI), graduating in 2016. The same year he continued his masters degree in Mathematics at LEI, and obtained his M.Sc. degree (cum laude) in 2018 with a thesis titled *"Invariant Distributions for a Generalization of Wright's delay equation"* under the supervision of dr.ir. Onno van Gaans. On the 1st of May 2019 he began as a PhD candidate at LEI, with Prof. dr. Aad van der Vaart as promoter and dr. Botond Szabó as supervisor. On the 1st of June 2021 he started working at Delft University of Technology, where he continued his work as a PhD candidate under the same supervision.