

CONSISTENCY OF EMPIRICAL BAYES AND KERNEL FLOW FOR HIERARCHICAL PARAMETER ESTIMATION

YIFAN CHEN, HOUMAN OWHADI, AND ANDREW M. STUART

Abstract. Gaussian process regression has proven very powerful in statistics, machine learning and inverse problems. A crucial aspect of the success of this methodology, in a wide range of applications to complex and real-world problems, is hierarchical modeling and learning of hyperparameters. The purpose of this paper is to study two paradigms of learning hierarchical parameters: one is from the probabilistic Bayesian perspective, in particular, the empirical Bayes approach that has been largely used in Bayesian statistics; the other is from the deterministic and approximation theoretic view, and in particular the kernel flow algorithm that was proposed recently in the machine learning literature. Analysis of their consistency in the large data limit, as well as explicit identification of their implicit bias in parameter learning, are established in this paper for a Mat'ern-like model on the torus. A particular technical challenge we overcome is the learning of the regularity parameter in the Mat'ern-like field, for which consistency results have been very scarce in the spatial statistics literature. Moreover, we conduct extensive numerical experiments beyond the Mat'ern-like model, comparing the two algorithms further. These experiments demonstrate learning of other hierarchical parameters, such as amplitude and lengthscale; they also illustrate the setting of model misspecification in which the kernel flow approach could show superior performance to the more traditional empirical Bayes approach.

1. Introduction

1.1. Background and context. Gaussian process regression (GPR) is important in its own right, and as a prototype for more complex inverse problems in which there is a possibly indirect, nonlinear set of observations. An important reason for the success of GPR in applications is its ability to learn hyperparameters, entering through a hierarchical prior, from data. Learning of these hyperparameters is typically achieved through fully Bayesian (sampling) or empirical Bayesian (optimization) methods. However, new approaches suggested in the machine learning literature, particularly the kernel flow method [25], rely on approximation theoretic criteria that can be traced back to the classical idea of cross-validation for model selection. The primary goal of this paper is to study and compare these two approaches. Special attention will be paid to their large data consistency, implicit bias, and robustness to model misspecification.

Received by the editor May 22, 2020, and, in revised form, February 1, 2021.
2020 *Mathematics Subject Classification.* Primary 62C10, 41A05, 35Q62.

The first author gratefully acknowledged the support of the Caltech Kortchack Scholar Program. The second author gratefully acknowledged support from AFOSR (grant FA9550-18-1-0271) and ONR (grant N00014-18-1-2363). The third author was grateful to AFOSR (grant FA9550-17-1-0185) and NSF (grant DMS 18189770) for financial support. The first, second, and third authors gratefully acknowledged support from AFOSR MURI (FA9550-20-1-0358).

1.2. Gaussian process regression. We start with a brief introduction to GPR; for simplicity, we focus on the noise-free scenario. The target is to recover a function $u^\dagger : D \rightarrow \mathbb{R}$ from pointwise data $y_i = u^\dagger(x_i)$ for $1 \leq i \leq N$, where $x_i \in D \subset \mathbb{R}^d$ and D is a compact domain. This problem often appears in fields such as supervised learning in machine learning, non-parametric regression in statistics, and interpolation in numerical analysis.

The GPR solution to this problem is as follows. Given a family of positive definite covariance/kernel functions $K_\theta : D \times D \rightarrow \mathbb{R}$ where $\theta \in \Theta$ is a hyperparameter, GPR approximates $u^\dagger$ with the conditional expectation

$$(1.1) \quad u(\cdot, \theta, X) := \mathbb{E}[\xi(\cdot, \theta) \mid \xi(X, \theta) = u^\dagger(X)] = K_\theta(\cdot, X)[K_\theta(X, X)]^{-1}u^\dagger(X),$$

where $\xi(\cdot, \theta) \sim \text{GP}(0, K_\theta)$ is a centered Gaussian process¹ (GP) with covariance function K_θ . We have used the following compressed notation:

$$X := (x_1, \dots, x_N)^T \quad \text{and} \quad u^\dagger(X) := (u^\dagger(x_1), \dots, u^\dagger(x_N))^T.$$

Moreover, $K_\theta(X, X)$ denotes the $N \times N$ dimensional Gram matrix with (i, j) -th entry $K_\theta(x_i, x_j)$, and $K_\theta(\cdot, X)$ is a function mapping D to $\mathbb{R}^N$ with i -th component $K_\theta(\cdot, x_i) : D \rightarrow \mathbb{R}$.

Normally, every $\theta \in \Theta$ produces a solution $u(\cdot, \theta, X)$ that agrees with $u^\dagger$ on X . Nevertheless, different choices may yield distinct out-of-sample errors, known as generalization errors in the machine learning context. Therefore, it is of paramount importance to learn a good hierarchical parameter θ adaptively from data.

1.3. Two approaches. In this paper, we study two approaches to the question posed above, both based on selecting θ as the optimizer of a variational problem.

1.3.1. Empirical Bayes approach. The empirical Bayes (EB) approach addresses the question by proposing a statistical model. It formulates a prior distribution on the pair (ξ, θ) by assuming that θ is sampled from a prior distribution and ξ is then sampled from the conditional distribution of $\xi \mid \theta$; then, it finds the posterior distribution of the pair (ξ, θ) conditioned on $\xi(X) = u^\dagger(X)$, and selects the parameter θ that maximizes the marginal probability of θ under this posterior. For simplicity, we work with uninformative priors, which lead to the following objective function:

$$(1.2) \quad L^{\text{EB}}(\theta, X, u^\dagger) = u^\dagger(X)^T [K_\theta(X, X)]^{-1} u^\dagger(X) + \log \det K_\theta(X, X).$$

This is also twice the negative marginal log likelihood of θ given the data $u^\dagger(X)$. Then, EB will choose θ by minimizing this objective function, namely

$$(1.3) \quad \theta^{\text{EB}}(X, u^\dagger) := \underset{\theta \in \Theta}{\operatorname{argmin}} L^{\text{EB}}(\theta, X, u^\dagger).$$

1.3.2. Approximation theoretic approach. Approximation theoretic considerations, on the other hand, provide a different answer without proposing statistical models. This methodology proceeds by asking for an ideal θ that minimizes $d(u^\dagger, u(\cdot, \theta, X))$

¹Recall that the covariance function K_θ of a Gaussian process $\text{GP}(0, K_\theta)$ is the kernel of the integral operator representation of C_θ in the covariance operator notation $N(0, C_\theta)$. Connections between these perspectives are reviewed in Subsection 2.1. We will use the covariance operator notation more frequently later in this paper.

for some cost function d . Though in practice $u^\dagger$ is not available, there are ideas in cross-validation that split χ into training data and validation data, and use the approximation error in validation data to estimate the exact error. Inspired by this idea, we could turn to optimize the following objective function:

$$(1.4) \quad d(u(\cdot, \theta, \chi), u(\cdot, \theta, \pi\chi)),$$

where we write $\pi\chi$ for a subset of χ obtained by subsampling a proportion, say one-half, of χ .

In this paper, we focus on a particular choice of d that originates from the Kernel Flow (KF) approach [25]. To describe it, we denote by $(H_\theta, \|\cdot\|_{K_\theta})$ the associated *Reproducing Kernel Hilbert Space* (RKHS) for the kernel K_θ ; note that $\|K_\theta(\cdot, x)\|_{K_\theta}^2 = K_\theta(x, x)$. The objective function in KF is chosen as

$$(1.5) \quad L^{\text{KF}}(\theta, \chi, \pi\chi, u^\dagger) := \frac{\|u(\cdot, \theta, \chi) - u(\cdot, \theta, \pi\chi)\|_{K_\theta}^2}{\|u(\cdot, \theta, \chi)\|_{K_\theta}^2}.$$

This measures the discrepancy in the RKHS norm between the GPR solution using the whole data χ and using a subset of the data $\pi\chi$, normalized by the RKHS norm of the former.

Remark 1.1. As explained above, we understand the numerator as an estimation of the error $\|u^\dagger - u(\cdot, \theta, \chi)\|_{K_\theta}^2$. Such error estimate, based on comparing solutions obtained via different data resolutions, is a widely used idea in numerical analysis.

Based on Galerkin orthogonality (see [25]), the objective function admits a finite dimensional representation formula that is convenient for numerical computation:

$$(1.6) \quad L^{\text{KF}}(\theta, \chi, \pi\chi, u^\dagger) = 1 - \frac{u^\dagger(\pi\chi)^\top [K_\theta(\pi\chi, \pi\chi)]^{-1} u^\dagger(\pi\chi)}{u^\dagger(\chi)^\top [K_\theta(\chi, \chi)]^{-1} u^\dagger(\chi)}.$$

Then, the KF estimator is defined as

$$(1.7) \quad \theta^{\text{KF}}(\chi, \pi\chi, u^\dagger) := \underset{\theta \in \Theta}{\operatorname{argmin}} L^{\text{KF}}(\theta, \chi, \pi\chi, u^\dagger).$$

Remark 1.2. The existence of the finite-sample formula (1.6) is attributed to the choice of the RKHS norm in comparing solutions. It is essentially a consequence of the standard representer theorem. Additional motivations for using the RKHS norm will be reviewed in Subsection 1.5.2.

1.3.3. Guiding observations and goals. The EB and KF algorithms estimate the parameter θ from the observed data, the number of which can vary considerably. Thus, a basic question to ask is whether the estimators attain meaningful limits as data accumulate:

- (1) Consistency: how do θ^{EB} and θ^{KF} behave in the large data limit, i.e., as the number of data N goes to infinity?

Meanwhile, since we have two estimators, it is natural to compare their performance. Indeed, we observe that EB and KF have distinct objectives: EB seeks to estimate the most likely parameters of the distribution assumed to generate the data, while KF chooses parameters to minimize an estimate of the approximation error in a parameter-dependent RKHS norm, targeting at the approximation efficiency of the underlying function. Moreover, EB is always probabilistic, while KF need not be.

These differences motivate the implicit bias question that has been popular in the machine learning community, and the model misspecification question that is common in mathematical modeling:

- (2) Implicit bias: what is the selection bias of EB and KF, or, how should the obtained estimators θ^{EB} and θ^{KF} be interpreted in practice?
- (3) Model misspecification: how do θ^{EB} and θ^{KF} behave when there is a mismatch between the data-generating mechanism and the model used to regress the data?

The precise goal of this paper is to address these questions for certain concrete models, either theoretically or experimentally.

1.4. Our contributions. Our contributions in this paper are twofold and explained in the following two subsections.

1.4.1. Consistency and implicit bias. The first part of this work is devoted to the questions of consistency and implicit bias. We study a Mat'ern-like model on the torus, in which $u^\dagger$ is a sample drawn from the Mat'ern-like Gaussian process, with three parameters $\theta = (\sigma, \tau, s)$ that quantify the amplitude, inverse lengthscale and regularity of the process. The detailed definition is in Subsection 2.1.

Our main analysis concerns learning the regularity parameter s using EB and KF. When the sampled points χ are equidistributed, we achieve the following contributions:

- Consistency: we prove that the EB estimator converges to s in the large data limit, while the KF estimator converges to $\frac{s-d/2}{2}$, so that s is also determined. Their variances are also computed and compared.
- Implicit bias: we characterize the selection bias of EB and KF algorithms, in terms of the L^2 error between $u^\dagger$ and the GPR solution using learned parameters — this is the so-called generalization error. It is found that EB selects the parameter that achieves the minimal L^2 error in expectation, while KF selects the minimal parameter that suffices for the fastest rate of convergence of the L^2 error to 0 as the data density increases.

We can interpret these contributions from two perspectives. From the machine learning side, we are able to show that KF, as a machine learning method, has a well-defined large data limit for the Mat'ern-like model. Furthermore we can characterize clearly its implicit bias in terms of L^2 generalization errors. Thus, this paper leads to a *first* theory for the KF learning algorithm.

From the spatial statistics side, our analysis contributes to a novel consistency theory for estimating the regularity parameter of Mat'ern-like fields in general dimensions. Such results are scarce in the spatial statistics literature; the techniques we use to prove consistency may be of independent interest and applicable beyond the setting considered here.

We also include numerical studies concerning the learning of the amplitude parameter σ and the inverse lengthscale parameter τ ; these experiments contribute to a more complete picture of GPR using the Mat'ern-like field with hierarchical parameters. Moreover, we provide numerical experiments for several other well-specified models beyond the Mat'ern-like model, thus further extending the scope of discussions.

1.4.2. *Model misspecification.* The second part of this work considers model misspecification: the data generating model for $u^\dagger$ and the model K_θ used for regression do not match. We adopt the following setting:

- We model the truth $u^\dagger$ either as a GP, using a variety of covariance functions, or as a deterministic function which solves a PDE.
- The kernel K_θ is chosen to be Green's function of various differential operators, where θ encodes information beyond the amplitude, lengthscale, and regularity of the field. For example we choose θ to be the location of a discontinuity within a conductivity field.

In this setting we observe distinct behavior distinguishing EB and KF. This raises the discussion of how to choose which algorithm to use when solving practical problems where misspecification is to be expected. Our numerical study explores several misspecification possibilities, showing that KF could be competitive with EB in certain scenarios.

1.5. **Literature review.** In this subsection, we review the related literature. Several fields are of relevance, so we label them to help organize the review.

1.5.1. *Regression and inverse problems.* Regression is a form of inverse problem [7], and if formulated in a Bayesian fashion, it falls within the scope of Bayesian nonparametric estimation [11, 15]. In the paper [19] a simple class of linear inverse problems was studied from the perspective of posterior consistency, and it was demonstrated that the rate of posterior convergence depends sensitively on the relationship between regularity of the true function being sought and the regularity of draws from the prior. This motivates the need for hierarchical procedures that adapt, on the basis of the data, the regularity of draws from the prior. In [18] the work in [19] was extended to cover the data-adapted learning of the regularity parameter in the prior; as the authors note: theoretical work “that supports the preference for empirical or hierarchical Bayes methods does not exist at the present time, however. It has until now been unknown whether these approaches can indeed robustify a procedure against prior mismatch. In this paper, we answer this question in the affirmative.” This analysis, however, requires simultaneous diagonalization of a self-adjoint operator formed from the forward model and the covariance operator, for all values of the hyper-parameter. Consistency is studied without this assumption in [42], and extended to the study of emulation within Bayesian inversion in [35] and to empirical Bayesian procedures in [36]. The papers [18] and [36] also use the EB loss function (1.2). In [9] estimation of hyper-parameters in Gaussian priors is discussed in the context of MAP estimators.

1.5.2. *Kernel flow and cross-validation.* The KF loss function in (1.6) was originally derived in [25] and motivated from the perspective of optimal recovery theory. It can be interpreted, from a numerical homogenization perspective [24], as the relative energy contained in the fine scales (in the unresolved part) of $u^\dagger$. In the paper [25], the proposed loss function to be optimized (via SGD) has the form

$$(1.8) \quad \mathbb{E}_{\pi_1} \mathbb{E}_{\pi_2} L^{\text{KF}}(\theta, \pi_1 X, \pi_2 \pi_1 X, u^\dagger),$$

where $\pi_1 X$ is a subsampling of X and $\pi_2 \pi_1 X$ is a further subsampling of $\pi_1 X$. This choice reduces the dimension of the kernel matrix and enables fast computation per iteration. Although the KF loss appears to be new, it can be seen as a variant of cross-validation (CV), which is a commonly used model selection/parameter

estimation criteria [1, 10, 20]. A theoretical understanding of the consistency of CV “is very much of interest” [45] since its convergence rate can be shown to be asymptotically minimax [34] or near minimax optimal [37, 39] while having a lower computational complexity [48] than MLE (maximum likelihood estimation). The consistency of parameter estimation for the Ornstein-Uhlenbeck process has been studied in [46] for MLE, and [4] for CV.

In the setting of hyperparameter estimation of GPs, comparing MLE with CV can be traced back to Wahba and Wendelberger [40] and Stein [32] who compared variants of these procedures² for choosing the smoothing parameter of a smoothing spline; they observed that while MLE is optimal when the model is well-specified, CV may perform better (than MLE) under misspecification (see also [3] for theoretical analysis and [41] for a practical example involving real data) and has a comparable rate of convergence when the model is correct (Stein [32] observed that “both estimates are asymptotically normal with the CV estimate having twice the asymptotic variance of the MLE estimate” and suggested that “The penalty for using CV instead of MLE when the stochastic model is correct is greater for higher-order smoothing splines, both in terms of the efficiency in estimating the smoothing parameter and the impact on subsequent predictions”). We also refer to [21] for a detailed numerical comparison between MLE and CV for estimating spline smoothing parameters. As observed in [30], these comparisons “are relevant for both numerical analysts and statisticians” since kernel interpolation can be interpreted as both approximating a deterministic unknown function from quadrature points or as estimating a sample from a Gaussian process from pointwise measurements.

1.5.3. Machine learning and kernel learning. Kernel methods and GPs have long been used in machine learning [16, 27]. Learning a good kernel for a given task is very important in practice. Many works have tried to learn a kernel from data based on different criteria; for example, in [2], the kernel is modified to make the model have a large margin in classification, and in [6], the kernel is selected to have a small local Rademacher complexity. EB and KF loss functions in this paper have also been used in [25, 27, 44].

The recent discovery of the neural tangent kernel regime for overparameterized models [17] and the identification [23] of warping kernels [25, 26, 29, 31] as the infinite depth limit of residual neural networks [14] also suggest that a theoretical understanding of kernel selections may lead to important insights for neural network based machine learning. This line of work suggests that it may be fruitful to consider machine learning directly as the problem of selecting an underlying kernel (by minimizing nonlinear functionals of the empirical distribution such as (1.2) or (1.6)) and learning based on this kernel; in this perspective one has hierarchical GPR with kernel itself as the hyperparameter. This may be more effective than simply fitting the data by minimizing a generalized moment, i.e., a linear functional, of the empirical distribution, which is popularly used in empirical risk minimization. Numerical experiments presented in [47] and [13], based on the KF methodology in [25], provide evidence that (1) this point of view could improve test errors, generalization gaps, and robustness to distribution shifts in the training of ANNs, and (2) kernel methods can be a simple and effective approach for learning dynamical

²Modified maximum likelihood estimation and generalized cross validation.

systems and surrogate models, with the underlying kernel also learned from data (using KF and its variants). This further motivates the desire to understand the KF-based estimation of θ .

1.6. Organization. The rest of this paper is organized as follows. Section 2 is devoted to learning the regularity parameter of the Matérn-like model, where the large data consistency is proved and implicit bias is characterized. Most of the detailed proofs are deferred to Section 6, and concise intuitive ideas are presented in Section 2 for the sake of readability. Section 3 considers other well-specified models, including the learning of the lengthscale and amplitude parameters in the Matérn-like model, or beyond the Matérn-like model. Experiments are provided concerning consistency and variance of these EB and KF estimators. Section 4 covers discussions on model misspecification through numerical studies. The purpose of the numerical experiments is twofold: (i) to demonstrate the extent to which the ideas learned through the analysis of consistency, which focuses primarily on the regularity parameter, extends to other parameters; (ii) to compare the performance of the EB and KF estimators quantitatively, since use of the latter is somewhat new in this area and its potential pros and cons need to be evaluated. Finally, we conclude this paper in Section 5.

2. Regularity parameter learning for the Matérn-like model

In this section, we study a Matérn-like model on the torus. We start with definitions of this model in Subsection 2.1, followed with definitions of EB and KF estimators in this context in Subsection 2.2. Then, in Subsection 2.3, we present our theory for the consistency of EB and KF estimators in learning the regularity parameter, with experiments included to demonstrate the correctness and implications of the theory. In particular, the implicit bias of these two estimators is explained. We outline the sketch of proofs for the theoretical result in Subsections 2.4, 2.5 and 2.6, and summarize several observations in Subsection 2.7. Subsection 2.8 provides additional experiments discussing the variance of these estimators.

2.1. The Matérn-like model. We follow the general set-up in Subsections 1.2 and 1.3, where we have mentioned all the abstract ingredients such as the physical domain D , the truth $u^\dagger$, the kernel K_θ , and the data location χ . In the current and next subsections, we will specify the exact meaning of these terms for a Matérn-like model on the torus. We will also make remarks to explain its connection to the standard Whittle-Matérn process in the whole domain; see Remark 2.2.

2.1.1. The physical domain. We set D to be $T^d = [0, 1]_{\text{per}}^d$, the d dimensional unit torus; this will be the domain that we use for all our analysis. We need to introduce some mathematical concepts related to functions defined on this torus T^d . First, the space of square integrable functions on T^d with mean 0 is denoted by

$$(2.1) \quad L^2(T^d) := \left\{ v : T^d \rightarrow \mathbb{R} : \int_{T^d} |v(x)|^2 dx < \infty, \int_{T^d} v(x) dx = 0 \right\}.$$

The L^2 inner product and norm are denoted by $[\cdot, \cdot]$ and $\|\cdot\|_0$ respectively.

In order both to define covariance operators and Sobolev spaces it is convenient to introduce the Laplacian operator. Let $-\Delta$ be the negative Laplacian equipped with periodic boundary conditions on T^d and restricted to functions with zero mean.

This operator has orthonormal eigenfunctions $\varphi_m(x) = e^{2\pi i \langle m, x \rangle}$ with corresponding eigenvalues $\lambda_m = 4\pi^2 |m|^2$, for every $m \in \mathbb{Z}^d \setminus \{0\}$, where $\mathbb{Z}^d$ denotes the d -fold tensor product of $\mathbb{Z}$, the set of non-negative integers. Here, i is the imaginary number, and $\langle m, x \rangle$ denotes the Euclidean inner product between $m, x \in \mathbb{R}^d$.

Now, we can write functions in $L^2(\mathbb{T}^d)$ as Fourier series:

$$(2.2) \quad v(x) = \sum_{m \in \mathbb{Z}^d} \hat{v}(m) e^{2\pi i \langle m, x \rangle},$$

where $\hat{v} : \mathbb{Z}^d \rightarrow \mathbb{R}$ is the Fourier coefficient that satisfies $\hat{v}(0) = 0$ and $\hat{v}(m) = [\hat{v}, \varphi_m]$ for $m \in \mathbb{Z}^d \setminus \{0\}$. This representation can be used to define useful Sobolev-like spaces. For every $t > 0$, the Sobolev-like space $H^t(\mathbb{T}^d) \subset L^2(\mathbb{T}^d)$ consists of functions with bounded $|\cdot|_t$ norm:

$$(2.3) \quad |\hat{v}|_t^2 := \sum_{m \in \mathbb{Z}^d} (4\pi^2 |m|^2)^t |\hat{v}(m)|^2 < \infty.$$

We note that $H^0(\mathbb{T}^d) = L^2(\mathbb{T}^d)$. For $t < 0$, the space $H^t(\mathbb{T}^d)$ is defined through duality. The Hilbert scale of function spaces defined through varying t serves as the basic ingredient to model the regularity of a function on $\mathbb{T}^d$.

2.1.2. The Mat'ern-like kernel and process. The Mat'ern-like covariance operator on the torus is defined by

$$(2.4) \quad C_\theta = \sigma^2 (-\Delta + \tau^2 I)^{-s},$$

where the parameter $\theta = (\sigma, \tau, s)$. The roles of the three parameters are reviewed in Remark 2.2. The orthonormal eigenfunctions of this operator are $\varphi_m(x) = e^{2\pi i \langle m, x \rangle}$ with corresponding eigenvalues $\sigma^2 (4\pi^2 |m|^2 + \tau^2)^{-s}$, for $m \in \mathbb{Z}^d \setminus \{0\}$.

The Mat'ern-like kernel function K_θ is related to the operator C_θ via

$$(2.5) \quad K_\theta(x, y) = [\delta(\cdot - x), C_\theta \delta(\cdot - y)],$$

where $\delta(\cdot - x)$ is the Dirac function centered at x . Equivalently, K_θ can be understood as the Green function of the differential operator C_θ^{-1} . Note that by Sobolev's embedding theorem, $s > d/2$ is required to make $K_\theta(x, y)$ pointwise well-defined (See Section 7.1.3 and Lemma 7.2 in [7]): $K_\theta(\cdot, y)$ then lies in the space of continuous functions for any $y \in \mathbb{T}^d$.

Remark 2.1. We also have the Mercer decomposition of the kernel function:

$$(2.6) \quad K_\theta(x, y) = \sum_{m \in \mathbb{Z}^d \setminus \{0\}} \sigma^2 (4\pi^2 |m|^2 + \tau^2)^{-s} \varphi_m(x) \varphi_m^*(y),$$

where φ_m^* is the complex conjugate of φ_m .

Given these function spaces and operators, we can define the Mat'ern-like process using the Gaussian measure notation:

$$(2.7) \quad \xi \sim \mathcal{N} \left(0, \sigma^2 (-\Delta + \tau^2 I)^{-s} \right).$$

This covariance operator viewpoint could be understood as follows: for any $f \in L^2(\mathbb{T}^d)$, the quantity $[f, \xi]$ is a Gaussian random variable with mean 0 and variance $[f, \sigma^2 (-\Delta + \tau^2 I)^{-s} f]$. We note that (2.7) is equivalent to the GP notation $\xi \sim \text{GP}(0, K_\theta)$. For more details on how to define Gaussian measures using operators

we refer to [5, 24]. A sample from this process can be realized by the Karhunen–Loève expansion

$$(2.8) \quad \xi(x) = \sum_{m \in \mathbb{Z}^d \setminus \{0\}} \sigma(4\pi^2|m|^2 + \tau^2)^{-s/2} \phi_m(x) \xi_m,$$

where ξ_m ($m \in \mathbb{Z}^d \setminus \{0\}$) are i.i.d. standard normal random variables; we have $E \xi(x) \xi(y) = K_\sigma(x, y)$. Numerically, we can draw a sample by truncating this series and restricting to a grid of values on the torus. Alternatively it is possible to discretize the differential operator C_σ^{-1} on a grid first, and then compute the discrete eigenfunctions to draw a sample. Such an idea is useful when the eigenvalues and eigenfunctions of C_σ^{-1} are not analytically known a priori. Indeed, when the operator is discretized into a matrix, the infinite dimensional Gaussian measure becomes a finite dimensional one with the covariance matrix being the discretization of C_σ . Drawing samples is then straightforward. In this section, however, we work on the torus and so the eigenvalues and eigenfunctions are known explicitly and the truncated Karhunen–Loève expansion could be employed.

Remark 2.2. The three parameters σ , τ and s quantify the amplitude, inverse lengthscale, and regularity of the process, respectively. This setting is similar to that of the standard Matérn process [12, 33], defined on the whole space $\mathbb{R}^d$, whose kernel function and associated covariance operator are both characterized by three parameters; see [22] for links to the solution of stochastic PDEs, an approach attributable to Whittle [12, 43]. The Matérn kernel function is

$$K_{\sigma,l,v}(x, y) = \sigma^2 \frac{2^{1-v}}{\Gamma(v)} \left(\frac{|x-y|}{l} \right)^v B_v \left(\frac{|x-y|}{l} \right),$$

for $x, y \in \mathbb{R}^d$, where B_v is the modified Bessel function of the second kind of order v . On $\mathbb{R}^d$, this kernel function corresponds to the covariance operator

$$C_{\sigma,l,v} = \frac{\sigma^2 l^d \Gamma(v + d/2) (4\pi)^{d/2}}{\Gamma(v)} (I - l^2 \Delta)^{-v-d/2}.$$

From this formula, the connection between the Matérn covariance operator in $\mathbb{R}^d$ and the Matérn-like kernel operator (2.4) on T^d becomes apparent. We restrict our analysis to the torus to exploit powerful Fourier series techniques. We will also comment on other boundary conditions in Subsection 2.7. For related results regarding the Matérn process in $\mathbb{R}^d$ or other bounded domains, we recommend the book [33]. We note that [33, Sec. 6.7] also considers a periodic version of the Matérn model and discusses (via the Fisher information matrix) the fixed domain asymptotics of the maximum likelihood estimate of the three parameters. By using the Mercer decomposition (2.6), the periodic case there is mathematically equivalent to the Matérn-like model on the torus that is considered in this paper. In the next subsection, we prove the consistency of estimators for the regularity parameter, providing a rigorous theory for this periodic model. It would be interesting, in future work, to combine this consistency with the properties of the Fisher information matrix established in [33, Sec. 6.7] to obtain Bernstein-von-Mises type theorems characterizing asymptotic normality of the estimator.

2.2. Regularity parameter learning. With the Matérn-like kernel and process defined, we move to discuss the parameter learning problem in this subsection. We fix $\sigma = 1$ and $\tau = 0$ in the Matérn-like model and focus on the regularity parameter

only. To proceed, we need to make precise the ground truth $u^\dagger$, the kernel, and the data location X , of the learning problem.

2.2.1. The ground truth. Our theoretical results regarding the consistency of EB and KF estimators will be based on the assumption that $u^\dagger$ is drawn from the GP $N(0, (-\Delta)^{-s})$ for some $s > d/2$.

Remark 2.3. We note some regularity properties of this GP here. The Cameron-Martin space for $\xi \sim (0, (-\Delta)^{-s})$ is $H^s(T^d)$ (for readers not familiar with the Cameron-Martin space, see Theorem 7.33 in [7]). However, ξ is not an element of this space, almost surely. Indeed, it holds that ξ belongs to $H^{s-d/2-\eta}(T^d)$ for any $\eta > 0$ almost surely (and to Hölder spaces with the same number of fractional derivatives; see Theorem 2.12 in [7]). Furthermore, since the Laplacian operator is homogeneous and thus the covariance operator is stationary in space, the regularity of the path is spatially homogeneous (the measure is space translation-invariant). Here, we refer, for this phenomenon, to ξ (as a function) having *homogeneous critical regularity* $s - d/2$ across T^d . If we drop the term “homogeneous”, we mean the property holds without the requirement of spatial homogeneity. Such behavior may occur for functions with spatial singularities.

Remark 2.4. We always require $s > d/2$, which ensures the continuity of the sample path of ξ almost surely and guarantees that $H^s(T)$ is a RKHS, according to discussions in Remark 2.3. Thus, the pointwise value of ξ makes sense.

2.2.2. The equidistributed data. We observe equidistributed pointwise values of $u^\dagger$ over the torus, i.e., the data lie on a lattice. To describe the data locations we introduce a level parameter $q \in \mathbb{N}$ such that, for a given q , we have the data locations $X_q := \{x_j : j \in J_q\}$, where $x_j = (j_1, j_2, \dots, j_d) \cdot 2^{-q}$ and $J_q := \{(j_1, j_2, \dots, j_d) \in \mathbb{N}^d : 0 \leq j_k \leq 2^q - 1 \forall k \leq d\}$. We also use the simplified notation $x_j = j2^{-q}$ throughout the paper.

2.2.3. The EB and KF estimators. We follow the definitions in Subsection 1.3. Here, the kernel function for the regularity learning problem will be

$$K_\theta(x, y) = [\delta(\cdot - x), (-\Delta)^{-\theta}(\cdot - y)],$$

where the parameter $\theta \in \mathbb{R}$. Similar to Remark 2.1, it has the following Mercer decomposition

$$(2.9) \quad K_\theta(x, y) = \sum_{m \in \mathbb{Z}^d \setminus \{0\}} (4\pi^2|m|^2)^{-\theta} \varphi_m(x) \varphi_m^*(y).$$

Numerically, we can compute it by truncating this infinite series. Fast Fourier Transform could be applied to speed up computation of the kernel matrix.

We adapt several notations from Subsection 1.3 to this specific problem, by writing t instead of θ , and q instead of α , and $K(t, q)$ instead of $K_\theta(x, y)$. These simplified notations make the analysis cleaner to present. Under such convention, the EB estimator for the regularity parameter is:

$$(2.10) \quad s^{\text{EB}}(q, u^\dagger) = \underset{t \in [d/2 + \delta, 1/\delta]}{\operatorname{argmin}} \mathcal{L}^{\text{EB}}(t, q, u^\dagger), \quad \mathcal{L}^{\text{EB}}(t, q, u^\dagger) := \|u(\cdot, t, q)\|_t^2 + \log \det K(t, q).$$

Here, $u(\cdot, t, q)$ is the GPR solution using the kernel function K_t and the observational data of $u^\dagger$ at X_q .

Remark 2.5. The formula (2.10) is the continuous formulation of the EB loss function, which is more convenient for theoretical analysis of consistency. The finite-sample formula (1.2) is more useful in numerical computation, and it can be derived from (2.10) by using the representer theorem.

Remark 2.6. As in Remark 2.4, we require the regularity parameter $t > d/2$. Here, furthermore, we introduce a number $\delta > 0$ and select the domain of the parameter to be $t \in [d/2 + \delta, 1/\delta]$; δ can be any arbitrary positive number, and this compactification of the parameter domain will simplify the subsequent analysis. The reader should not confuse real number δ with Dirac delta function δ .

For the KF loss function, we fix the subsampling operator to be equidistributed subsampling so that $\pi_{\chi_q} = \chi_{-1}$; for this choice, we can omit the dependence of the estimator on the subsampling operator π in the notation and write:

(2.11)

$$s^{\text{KF}}(q, u^\dagger) = \operatorname{argmin}_{t \in [d/2 + \delta, 1/\delta]} L^{\text{KF}}(t; \text{KF}) : \frac{\|u(\cdot, t, q) - u(\cdot, t, q - 1)\|^2}{\|u(\cdot, t, q)\|^2} \ll t$$

2.3. Consistency and implicit bias. In this subsection, we present our theory of consistency and characterize the implicit bias via numerical experiments. The sketch of proofs is given in the next subsections.

2.3.1. Main theorem. We have Theorem 2.7 regarding the consistency of the two statistical estimators in the large data limit:

Theorem 2.7. Fix $\delta > 0$. Suppose $u^\dagger$ is a sample drawn from the Gaussian process $\mathbf{N}(0, (-\Delta)^{-s})$. If $s \in [d/2 + \delta, 1/\delta]$ then, for the Empirical Bayesian estimator,

$$\lim_{q \rightarrow \infty} s^{\text{EB}}(q, u^\dagger) = s;$$

if $s = \frac{d}{2} \in [d/2 + \delta, 1/\delta]$ then for the Kernel Flow estimator,

$$\lim_{q \rightarrow \infty} s^{\text{KF}}(q, u^\dagger) = \frac{s - d/2}{2}.$$

In both cases the convergence is in probability with respect to randomly chosen $u^\dagger$.

Remark 2.8. Strictly speaking this theorem shows that EB consistently estimates the regularity parameter, whilst KF does not. However we make two observations about this. Firstly, the true value of s can be recovered from the KF estimator by a simple linear transformation. And, secondly, the value selected by KF is optimal with respect to minimizing a specific measure of generalization error (as we will show in the discussion of implicit bias in Subsection 2.3.3), and is of clear interest from this perspective.

Remark 2.9. The use of δ in the proof (and hence statement) of this theorem helps by compactifying the parameter space. In practice, numerics demonstrate that it is not intrinsic to the problem. We leave for future work the problem of a more refined theorem, and proof, which does not rely on it.

Remark 2.10. For economy of notation we will drop explicit reference to the dependence of the loss functions and the estimators on $u^\dagger$ in what follows; we will simply write $L^{\text{EB}}(t, q)$, $L^{\text{KF}}(t, q)$, $s^{\text{EB}}(q)$, $s^{\text{KF}}(q)$.

The remainder of this subsection is devoted to numerical experiments illustrating the theory, discussion of the implications of the theory (i.e. implicit bias), and an overview of the proof techniques we adopt.

2.3.2. Numerical illustration of theory. We present a numerical example to demonstrate the main theorem, and its consequences for regression. Consider the one dimensional case, i.e., $d = 1$. We set the ground truth $s = 2.5$ and so $\frac{s-d/2}{2} = 1$. The domain is discretized with $N = 2^{10}$ equidistributed grid points. For our first set of experiments we fix the resolution level of the data points to be $q = 9$, i.e., we have 2^9 equidistributed observations of the unknown function $u^\dagger$. In what follows the Laplacian is as defined in Subsection 2.1.2. Given a sample of $u^\dagger$ from $\mathcal{N}(0, (-\Delta)^{-s})$, we form the loss function for the EB and the KF estimators. We draw this sample using the formula (2.8) with $\sigma = 1$ and $\tau = 0$; we truncate the series to the grid resolution. A single realization of these loss functions is then shown in Figure 1.

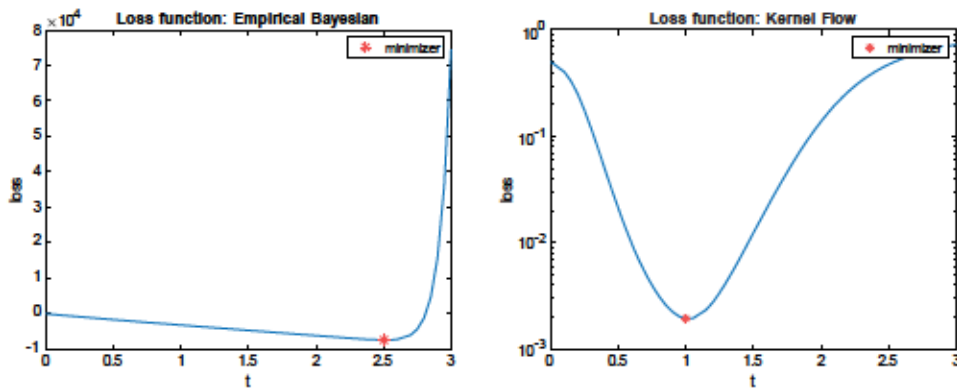


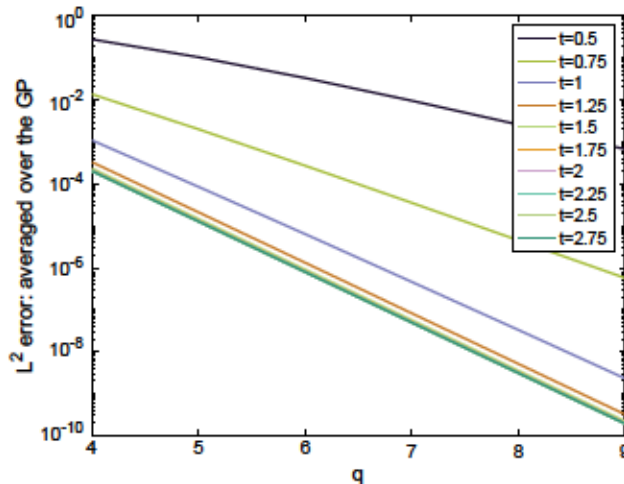
Figure 1. Left: EB loss; right: KF loss

We observe that the minimizer of the EB loss function is very close to $t = 2.5$, while the minimizer of the KF loss function is very close to $t = 1$, matching the predictions of Theorem 2.7. Furthermore, the loss functions exhibit some interesting features. Specifically, the EB loss function behaves as a linear function of t , for t less than s , and then blows up rapidly when t exceeds s . The KF loss function is more symmetric with respect to the minimizer $t = \frac{s-d/2}{2}$ in the logarithmic scale. We will make remarks that explain these observations in our theoretical analysis.

2.3.3. Implicit bias. We present here a second set of numerical experiments looking at the effect of the parameter value s selected by EB and KF on the approximation of the function $u^\dagger$, which is (typically) the primary goal of hierarchical parameter estimation. The experimental set-up is the same, but now we vary the resolution of the data points $q = 3, 4, \dots, 9$. We focus on the L^2 error between $u^\dagger$ and the GPR solution using learned parameters, i.e.,

$$\|u^\dagger(\cdot) - u(\cdot, t, q)\|_0^2.$$

We start, in Figure 2, by considering the error as a function of q , for different t . As we increase t , the regularity of the GP used for regression increases. In

Figure 2. L^2 error: averaged over the GP

order to illustrate clear trends, the L^2 error is averaged over the random draw of $u^\dagger \sim \mathcal{N}(0, (-\Delta)^{-s})$, so the effective error is $\mathbb{E}_{u^\dagger} \|u^\dagger(\cdot) - u(\cdot, t, q)\|_0^2$. From the figure, we can see that when t increases from 0.5 to 1, the convergence rate of the L^2 approximation error increases. Then, if we increase t further from 1 to 3, the slope of the convergence curve remains nearly the same. This demonstrates the fact that $s = \frac{s-d/2}{2}$ is the minimal t that suffices to achieve the fastest rate of L^2 error convergence. We have observed that this phenomenon is very stable with respect to the specific random draw: the general shape of the curves seen in Figure 2 is still observed when one specific draw of the true random process is used, although the resulting figure contains fluctuations and is not as clear as the average case that we show.

On the other hand, we can compute $\mathbb{E}_{u^\dagger} \|u^\dagger(\cdot) - u(\cdot, t, q)\|_0^2$ for $q = 9$ as a function of t ; see Figure 3. The optimality of the value $s = 2.5$ is clear. However, unlike the experiments in Figure 2, this result is not stable with respect to the random instance of the GP: the minimizer of the L^2 error fluctuates wildly in our experiments.

In summary, the second set of numerical experiments indicates the following implications for the regression accuracy of the EB and KF approaches to hierarchical parameter estimation. The KF estimator selects the minimal t that suffices to achieve the fastest rate of approximation error in the L^2 norm for a given fixed truth; in contrast, the EB estimator converges to the t that achieves the minimal L^2 error, averaged over the draw $u^\dagger \in \mathcal{N}(0, (-\Delta)^{-s})$. Note that KF is based on purely approximation theoretic considerations whilst EB is founded on statistical considerations — they attain very different implicit bias in selecting parameters.

2.3.4. Further discussion of the theory. We provide some further discussions of the implications of Theorem 2.7 in this subsection. The theory shows that the EB estimator recovers the ground truth parameter s of the statistical model. This is in line with expectations since the methodology is designed to recover the most likely value of s , given the data, and since the Gaussian measures occurring for different s are mutually singular. In the literature, such consistency results are primarily for

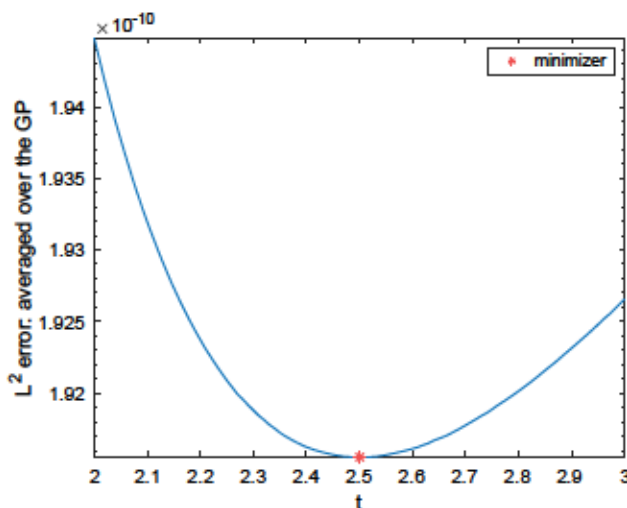


Figure 3. L^2 error: averaged over the GP, for $q = 9$

observational data in the Fourier domain; thus, the observation operator commutes with the prior. Here, our data model is in the physical domain, which leads to the need for considerably more sophisticated analysis, due to the noncommutativity of the observation operator and the prior operator, and yet is a much more practically useful setting, justifying the investment in the somewhat involved analysis. Our proof provides a novel sharp upper and lower bound on the terms $\|u(\cdot, t, q)\|_2^2$ and $\log \det K(t, q)$, based on techniques in approximation theory and the multiresolution analysis developed in [24]. Our techniques may have broader applications in analyzing the observational model in the physical domain.

Another interesting phenomenon shown in Theorem 2.7 is that the KF estimator, first proposed in [25] as a method to learn kernels for machine learning tasks, achieves a rather different consistency behavior, with the large data limit being $s = d/2$. This fact has the following consequence: if the ground truth function $u^\dagger$ has homogeneous critical regularity $s = d/2$, then the KF estimator will converge to half the critical regularity in the large data limit.

To understand the mechanism behind this effect, we observe that the KF loss is a surrogate for the (relative) $\|\cdot\|_t$ -norm approximation error between $u^\dagger$ and $u(\cdot, t, q)$. Furthermore, approximation theory implies that the GP regressor $u(\cdot, t, q)$ is also the optimal $\|\cdot\|_t$ -norm approximant of $u^\dagger$ in the linear span of the basis functions $\{\Delta^{-s/2}(x - x_j)\}_{j \in \mathbb{J}}$. Under this perspective, we see the KF loss incorporates two competing factors in the approximation: increasing t improves the approximation error by increasing the regularity of the basis functions while worsening the measurement of that approximation error by using a stronger norm. The balance between these two competing factors is achieved when t is half the critical regularity, which is the parameter that KF eventually picks. Our proof provides a detailed demonstration of this phenomenon.

In short, EB learns hierarchically based on statistical principles, whilst KF learns based on approximation theoretic ones. The consistency results presented here provide evidence that the interplay between statistical estimation and numerical

approximation can be very useful for parameter estimation and kernel learning in general, thus suggesting new ways of thinking hierarchically. This perspective is one of the main messages that we convey in this paper.

2.3.5. Proof strategy. Subsections 2.4, 2.5, 2.6 are devoted to proving Theorem 2.7. For the sake of understanding, we provide a high-level view of our proof strategies in this subsection. Fourier analysis plays an important role in the proof. It allows us to analyze the approximation error in a very precise way under this equidistributed design setting.

In our proof, we begin by establishing tight bounds on the terms that appear in the objective functions, i.e., $\|u(\cdot, t, q)\|_2^2$, $\log \det K(t, q)$ and $\|u(\cdot, t, q) - u(\cdot, t, q - 1)\|_2^2$, using the toolkit we develop in Subsection 2.4. The norms $\|u(\cdot, t, q)\|_2^2$ and $\|u(\cdot, t, q) - u(\cdot, t, q - 1)\|_2^2$ are expressed as random (as a function of $u^\dagger$) series and we carefully analyze the dependencies of the random variables to establish the convergence in probability. For $\log \det K(t, q)$, we employ the multiresolution approach introduced in [24] to establish a tight estimate of the spectrum of the Gram matrix from below and above. Given these estimates, we provide an intuitive understanding of how the loss functions behave and how the minimizers converge in Subsections 2.5, 2.6. In the rigorous treatment, the sharp bounds on the different components of the objective functions will be combined with the uniform convergence result of random series in [38] to obtain the convergence of minimizers.

2.3.6. Notations. In many parts of the analysis, we need to develop tight estimates on the terms appearing in the loss functions. Some useful notation for comparing different terms are introduced here. We write $A \lesssim B$ if there exists a constant C independent of q, t such that

$$\frac{1}{C}B \leq A \leq CB.$$

The constant may depend on the dimension d and on δ . Correspondingly, if we use $A \gtrsim B$ or $A \asymp B$, then only one side of the above inequality holds.

Fourier analysis plays a critical role in the analysis. We always use $u^\dagger$ for the ground truth function, while we omit the $\dagger$ symbol for ease of notation when discussing its Fourier transform, and write $\hat{u}$; we will also use $\hat{u}$, with more arguments, to denote the Fourier transform of the Gaussian process mean; see the discussion following Theorem 2.13. In the Fourier domain, we let $B_q := \{m \in \mathbb{Z} : -2^{q-1} \leq m \leq 2^{q-1} - 1\}$ and $B_q^{\otimes d} = B_q \otimes B_q \otimes \cdots \otimes B_q$ be the tensor product of d multiples of B_q . We have that $B_q^{\otimes d}$ is a box concentrating around the origin, so only the low-frequency part of the Fourier coefficients is considered.

2.4. Toolkit: Fourier series characterization. In this subsection, we prepare the necessary tools that are used to prove the main theorem of this paper.

We start by establishing a Fourier series characterization for $u(\cdot, t, q)$. This is a key ingredient in expressing the terms in the loss functions as random series. Our approach, using Fourier series, is motivated by the papers [8, 28], where the approximation power of shift-invariant subspaces of $L^2(\mathbb{R}^d)$ is studied; in our case we use related ideas in the $L^2(\mathbb{T}^d)$ setting.

To find the representation of the term $u(\cdot, t, q)$, we invoke its definition, i.e. $u(\cdot, t, q)$ is obtained by GP regression with the q -level data and the covariance

function $(-\Delta)^{-t}$. We use the representer theorem from GPR. Concretely, let the set of basis functions be

$$F_{t,q} = \text{span}_{j \in \mathbb{Z}^d} \{(-\Delta)^{-t} \phi(\cdot - x_j)\},$$

then, $u(\cdot, t, q)$ is the best approximation in $F_{t,q}$ to the true function under the $\|\cdot\|_t$ norm. Let us define

$$\hat{F}_{t,q} := \{g : Z^d \rightarrow \mathbb{C}, \text{ there exists an } f \in F_{t,q} \text{ such that } g = \hat{f}\},$$

the Fourier coefficients of functions in $F_{t,q}$. A quick observation is that for every $g \in \hat{F}_{t,q}$, we must have $g(0) = 0$ because of the mean zero property of $f \in F_{t,q}$. Proposition 2.11 gives a complete characterization of the basis functions in $\hat{F}_{t,q}$, for $t > d/2$.

Proposition 2.11. *For any $g \in \hat{F}_{t,q}$, there exists a 2^q -periodic function p on Z^d , such that*

$$g(m) = \begin{cases} |m|^{-2t} p(m), & m \neq 0 \\ 0, & m = 0. \end{cases}$$

The proof is in Subsection 6.1. Next, we define a 2^q -periodization operator, which will be used to compute the representation of $\hat{u}(m, t, q)$.

Definition 2.12. The operator T_q is defined as a mapping from the space of functions on Z^d to itself, such that

$$(T_q g)(m) := \sum_{\beta \in Z^d} g(m + 2^q \beta), \quad m \in Z^d,$$

whenever the right hand side series converges for the function $g : Z^d \rightarrow \mathbb{R}$. We also define

$$(2.12) \quad M_q^t(m) := \begin{cases} \sum_{\beta \in Z^d \setminus \{0\}} |2^q \beta|^{-2t}, & \text{if } m = j \cdot 2^q \text{ for some } j \in Z^d \\ \sum_{\beta \in Z^d} |m + 2^q \beta|^{-2t}, & \text{else.} \end{cases}$$

Both $T_q g$ and M_q^t are 2^q -periodic functions on Z^d . Based on this definition, Theorem 2.13 presents the explicit form of the Fourier transform of $u(\cdot, t, q)$; the proof is in Subsection 6.2. The proof relies on the Galerkin orthogonality property of $u(\cdot, t, q)$ due to its being the optimal approximate solution.

Theorem 2.13. *Let $\hat{u}(\cdot, t, q)$ be the Fourier coefficients of $u(\cdot, t, q)$, then for $m \in Z^d$, we have*

$$\hat{u}(m, t, q) = \begin{cases} 0, & \text{if } m = 0 \\ |m|^{-2t} \frac{(T_q u^*)(m)}{M_q^t(m)}, & \text{else,} \end{cases}$$

where $\hat{u}$ denotes the Fourier coefficients of $u^\dagger$.

This above representation is very useful for analyzing the terms $\|u(\cdot, t, q) - u^\dagger(\cdot, t, q)\|_{t,q}$ and $\|u(\cdot, t, q) - u^\dagger(\cdot, t, q)\|_{t,q}$. As well as studying the Fourier coefficients of $u(\cdot, t, q)$, which we denote by $\hat{u}(\cdot, t, q)$, we will also need to study the Fourier coefficients of $u^\dagger(\cdot)$ which, for ease of notation, we will denote by $\hat{u}(\cdot)$, henceforth, omitting the $\dagger$ symbol. It is thus important to look at the number of arguments of $\hat{u}$ to determine which object it is the Fourier transform of. Note also that $u(\cdot, t, q)$ is determined by $u^\dagger$; hence if $u^\dagger$ is random, so is $u(\cdot, t, q)$.

We will use the above Fourier analysis toolkit to study the consistency of EB and KF in the following two subsections.

2.5. Proof for the Empirical Bayesian estimator. In this subsection, we prove the consistency of the EB estimator. As explained before, our roadmap is to give a tight estimate of the loss functions first and then analyze the minimizers. For the norm term $\|u(\cdot, t, q)\|_t$ we invoke Theorem 2.13, based on which this term is expressed as a random series:

Proposition 2.14. *The $H^{-t}(\mathbb{T}^d)$ norm of $u(\cdot, t, q)$ has the representation*

$$\|u(\cdot, t, q)\|_t = (4\pi^2)^{-\frac{2-t}{2}} \sum_{m \in B_q^d} \frac{|T_q \hat{u}(m)|^2}{M_q^t(m)}.$$

Moreover, suppose $u^\dagger \sim \mathcal{N}(\mathbf{0}, (-\Delta)^{-s})$ for $s > \frac{d}{2}$, then

$$\|u(\cdot, t, q)\|_t^2 = (4\pi^2)^{t-s} \sum_{m \in B_q^d} \frac{M_q^s(m)}{M_q^t(m)} \xi_m^2,$$

where $\{\xi_m\}_{m \in B_q^d}$ are independent unit scalar Gaussian random variables.

Proof. Using Theorem 2.13, we get

$$\begin{aligned} \|u(\cdot, t, q)\|_t^2 &= \sum_{m \in \mathbb{Z}^d \setminus \{0\}} (4\pi^2)^t |m|^{2t} |\hat{u}(m, t, q)|^2 \\ &= (4\pi^2)^t \sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} \frac{|T_q \hat{u}(m)|^2}{|M_q^t(m)|^2} \\ &= (4\pi^2)^t \sum_{m \in B_q^d} \frac{M_q^t(m)}{|M_q^t(m)|^2} \frac{|T_q \hat{u}(m)|^2}{|M_q^t(m)|^2} \\ &= (4\pi^2)^t \sum_{m \in B_q^d} \frac{|T_q \hat{u}(m)|^2}{M_q^t(m)}, \end{aligned}$$

where in the third equality, we use the periodicity of the function $\frac{|T_q \hat{u}(m)|^2}{|M_q^t(m)|^2}$.

If we further assume $u^\dagger \sim \mathcal{N}(\mathbf{0}, (-\Delta)^{-s})$, then $\hat{u}(m) \sim \mathcal{N}(\mathbf{0}, (4\pi^2)^{-s} |m|^{-2s})$. For different m , these Gaussian random variables are independent. Thus, for different $m \in B_q^d$, we have $T_q \hat{u}(m) \sim \mathcal{N}(\mathbf{0}, (4\pi^2)^{-s} M_q^s(m))$, and they are independent. So we can write

$$\sum_{m \in B_q^d} \frac{|T_q \hat{u}(m)|^2}{M_q^t(m)} = (4\pi^2)^{-s} \sum_{m \in B_q^d} \frac{M_q^s(m)}{M_q^t(m)} \xi_m^2,$$

where $\{\xi_m\}_{m \in B_q^d}$ are independent unit scalar Gaussian random variables. $\square$

The independence of the random variables established in the preceding representation is crucial for the analysis. The terms $M_q^s(m), M_q^t(m)$ appear in the preceding; to analyze them we present a useful lemma below. The proof is in Subsection 6.3.

Lemma 2.15. *For $t \in [d/2 + \delta, 1/\delta]$ and $q \geq 0$, we have*

$$M_q^t(m) \sim \begin{cases} 2^{-2qt}, & \text{if } m = \mathbf{0} \\ |m|^{-2t}, & \text{if } m \in B_q^d \setminus \{\mathbf{0}\}. \end{cases}$$

Moreover, for $m \in B_q^d \setminus \{\mathbf{0}\}$, we have $M_q^t(m) - |m|^{-2t} \sim 2^{-2qt}$.

Now, we are ready to get the estimates of the loss function. Proposition 2.16 shows an upper and lower bound on the norm term.

Proposition 2.16 (Bound on the norm term). *Suppose $u^\dagger$ is a sample drawn from the Gaussian process $\mathbf{N}(\mathbf{o}, (-\Delta)^{-s})$ for $d/2 + \delta \leq s \leq 1/\delta$, then*

$$\|u(\cdot, t, q)\|_t^2 \sim 2^{-q(2s-2)\frac{1}{\delta}\zeta_0^2} + \sum_{m \in B_q^d \setminus \{\mathbf{o}\}} |m|^{2t-2s}\zeta_m^2,$$

where $\{\zeta_m\}_{m \in B_q^d}$ are independent unit scalar Gaussian random variables.

Proof. According to Lemma 2.15, for $m \in B_q^d \setminus \{\mathbf{o}\}$, we have $M_q^t(m) \sim |m|^{-2t}$; for $m = \mathbf{o}$, we have $M_q^t(m) \sim 2^{-2tq}$. Thus,

$$\begin{aligned} \|u(\cdot, t, q)\|_t^2 &= (4\pi^2)^{t-s} \left(\sum_{m \in B_q^d} \frac{M_q^s(m)}{M_q^t(m)} \zeta_m^2 \right) \\ &= (4\pi^2)^{t-s} \left(\sum_{m \in B_q^d \setminus \{\mathbf{o}\}} \frac{M_q^s(m)}{M_q^t(m)} \zeta_m^2 + \frac{M_q^s(\mathbf{o})}{M_q^t(\mathbf{o})} \zeta_0^2 \right) \\ &\sim 2^{-q(2s-2)\frac{1}{\delta}\zeta_0^2} + \sum_{m \in B_q^d \setminus \{\mathbf{o}\}} |m|^{2t-2s}\zeta_m^2. \end{aligned}$$

This completes the proof. $\square$

Proposition 2.16 states that the behavior of the norm term is nothing but a weighted sum of squares of independent Gaussian random variables, which is amenable to analysis. With this in mind, we state a lemma useful in the analysis of such random series, with proof deferred to Subsection 6.4.

Lemma 2.17. *Suppose $\{\zeta_m\}_{m \in \mathbb{Z}^d}$ are independent unit Gaussian random variables.*

- For $r > 0$, define the random series

$$a(r, q) = 2^{-qr} \sum_{m \in B_q^d \setminus \{\mathbf{o}\}} |m|^{r-d}\zeta_m^2.$$

Fix $\epsilon > 0$, then there exists a function $\gamma(r) > 0$ such that $\lim_{q \rightarrow \infty} a(r, q) = \gamma(r) > 0$ uniformly for $r \in [\epsilon, 1/\epsilon]$, where the convergence is in probability.

- For $r = 0$, define

$$a(\mathbf{o}, q) = \frac{1}{q} \sum_{m \in B_q^d \setminus \{\mathbf{o}\}} |m|^{-d}\zeta_m^2,$$

then there exists $\gamma(\mathbf{o}) \in (0, \infty)$ such that $\lim_{q \rightarrow \infty} a(\mathbf{o}, q) = \gamma(\mathbf{o})$ in probability.

We then move to the second term in the loss function, i.e., the log determinant term. It is deterministic and to study it we need a way of analyzing the spectrum of the Gram matrix. Proposition 2.18 gives upper and lower bounds on this term. The proof is in Subsection 6.5 and is motivated by analysis developed in the paper [24]. The idea is to use the Schur complement of the Gram matrix and rely on the variational characterization of the Schur complement to get a tight control on the spectrum. This technique is quite general and has been used in [24] to characterize

the spectrum of heterogeneous Laplacian operators; here we adapt it to fractional operators. On the other hand, for the homogeneous fractional Laplacian operators in this paper, it is also possible to calculate an explicit formula for the spectrum of $K(t, q)$, as has been used in Section 6.7 of [33]. We describe this simple proof in Subsection 6.5 but retain the proof employing the more general methodology as it may be useful for other problems.

Proposition 2.18 (Bound on the log det term). *For $d/2 + \delta \leq t \leq 1/\delta$, we have $(2t - d)g_1(q) - Cg_2(q) + K(t, 0) \leq \log \det K(t, q) \leq (2t - d)g_1(q) + Cg_2(q) + K(t, 0)$, where $g_1(q) = \sum_{k=1}^q (2^{kd} - 2^{(k-1)d})(-k \log 2)$ and $g_2(q) = (2^{qd} - 1)(2t - d)$. The constant C is independent of t, q . Moreover, $g_1(q) \sim -q2^{qd}$.*

With the loss function analyzed by the above results, the consistency of the EB estimator is readily stated as follows.

Theorem 2.19 (Consistency of Empirical Bayesian estimator). *Fix $\delta > 0$. Suppose $u^\dagger$ is a sample drawn from the Gaussian process $\mathcal{N}(0, (\Delta)^{-s})$. If $s \in [d/2 + \delta, 1/\delta]$ then*

$$\lim_{q \rightarrow \infty} s^{\text{EB}}(q) = s \quad \text{in probability.}$$

The detailed proof is in Subsection 6.6. We can understand the theorem intuitively by using the established results above. Recall there are two terms in the loss function: (1) the norm term $\|u(\cdot, t, q)\|_t^2$; (2) the log det term. For the norm term, from Proposition 2.16 and Lemma 2.17, its behavior for $q \rightarrow \infty$ is roughly

- Growing like $2^{q(2t-2s+d)}$ if $t > s - d/2$;
- Growing like q if $t = s - d/2$;
- Remaining bounded if $t < s - d/2$.

The log det term decreases like $-(2t - d)q2^{qd}$ according to Proposition 2.18. Noticing that the EB loss function has the form

$$L^{\text{EB}}(t, q) = \|u(\cdot, t, q)\|_t^2 + \log \det K(t, q),$$

we arrive at the following intuitive observations:

- When $t < s$, the dominant behavior of $L^{\text{EB}}(t, q)$ is controlled by the log determinant term, since the growth rate of the norm term $2^{q(2t-2s+d)} = o(q2^{qd})$. As a consequence, $L^{\text{EB}}(t, q)$ exhibits the overall behavior $-(2t - d)q2^{qd}$. Therefore, the loss function decreases linearly with t in this regime. This is consistent with what is observed in Figure 1.
- When $t \geq s$, the increasing speed of the norm term beats the decreasing rate of the log det term, so the norm term dominates the behavior of $L^{\text{EB}}(t, q)$. Overall, it is like $2^{q(2t-2s+d)}$, which increases exponentially with t ; again this is consistent with what is observed in Figure 1.

According to the above observations, the minimizer of $L^{\text{EB}}(t, q)$ will converge to s . To make the intuition leading to this conclusion rigorous, we need to use techniques of uniform convergence for random series. For details we refer to Subsection 6.6.

2.6. Proof for the Kernel Flow estimator. In this subsection, we establish the consistency of the KF estimator. As before, we start by estimating the growth behavior of terms that appear in the loss function. We begin with the interaction term $\mu(\cdot, t, q) - \mu(\cdot, t, q - 1)$. Similar to the analysis of the norm term in the preceding subsection, we represent it by using Fourier series.

Proposition 2.20. *The $H^t(T^d)$ norm of $u(\cdot, t, q) - u(\cdot, t, q-1)$ has the representation*

$$(2.13) \quad \|u(\cdot, t, q) - u(\cdot, t, q-1)\|_t^2 = (4\pi^2)^t \sum_{m \in B_q^d} M_q^t(m) \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2.$$

Proof. By Theorem 2.13, we have

$$\hat{u}(m, t, q) - \hat{u}(m, t, q-1) = \begin{cases} 0, & \text{if } m = 0 \\ |m|^{-2t} \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right), & \text{else.} \end{cases}$$

Thus,

$$\begin{aligned} \|u(\cdot, t, q) - u(\cdot, t, q-1)\|_t^2 &= (4\pi^2)^t \sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{2t} |\hat{u}(m, t, q) - \hat{u}(m, t, q-1)|^2 \\ &= (4\pi^2)^t \sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2 \\ &= (4\pi^2)^t \sum_{m \in B_q^d} M_q^t(m) \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2. \quad \square \end{aligned}$$

By carefully studying the correlation between the random variables appearing in the preceding proposition, we obtain lower and upper bounds in the following two propositions; proofs can be found in Subsections 6.7 and 6.8.

Proposition 2.21 (Lower bound on the interaction term). *Suppose $u^\dagger$ is a sample drawn from the Gaussian process $N(0, (-\Delta)^{-s})$ for $d/2 + \delta \leq s \leq 1/\delta$, then*

$$\|u(\cdot, t, q) - u(\cdot, t, q-1)\|_t^2 \gtrsim \sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{-2tq} |m|^{4t-2s} \zeta_{m,2}^2,$$

where $\{\xi_m\}_{m \in B_{q-1}^d \setminus \{0\}}$ are independent unit scalar Gaussian random variables.

The upper bound has a more complex form. We introduce the notation $Z_2^d = \{0, 1\}^d$ comprising d dimensional vectors with each component being in $\{0, 1\}$. In Proposition 2.22, we also use the convention that $|m|^a = 0$ for $m = 0$ and any $a \in \mathbb{R}$ to make the notation more compact.

Proposition 2.22 (Upper bound on the interaction term). *Suppose $u^\dagger$ is a sample drawn from the Gaussian process $N(0, (-\Delta)^{-s})$ for $d/2 + \delta \leq s \leq 1/\delta$, then*

$$\|u(\cdot, t, q) - u(\cdot, t, q-1)\|_t^2 \leq \sum_{k \in Z_2^d} \sum_{m \in B_{q-1}^d} (2^{-q(2s-2t)} + 2^{-2tq} |m|^{4t-2s}) \zeta_{k,m}^2,$$

where for a fixed $k \in Z_2^d$, $\{\xi_{k,m}\}_{m \in B_{q-1}^d}$ are independent unit scalar Gaussian random variables.

We remark that in the upper bound, the random variables for different k may exhibit correlation. However, since the term $\sum_{m \in B_{q-1}^d} (2^{-q(2s-2t)} + 2^{-2tq} |m|^{4t-2s}) \zeta_{k,m}^2$ has the same form for each k , and the number of different k is finite, it suffices to analyze the random series for a single k , in which we have the independence of random variables. The theorem is stated below.

Theorem 2.23 (Consistency of the Kernel Flow estimator). *Fix $\delta > 0$. Suppose u is a sample drawn from the Gaussian process $N(0, (-\Delta)^{-s})$. If $\frac{s-d/2}{2} \in [d/2+\delta, 1/\delta]$ then for the Kernel Flow estimator,*

$$\lim_{q \rightarrow \infty} s^{\text{KF}}(q) = \frac{s-d/2}{2} \quad \text{in probability.}$$

The idea behind the proof of the theorem is to combine Propositions 2.21, 2.22 and Lemma 2.17. Together they imply the growth behavior of the loss function

$$L^{\text{KF}} = \frac{\|u(\cdot, t, q) - u(\cdot, t, q-1)\|^2}{\|u(\cdot, t, q)\|_1^2}$$

as follows:

- When $t < \frac{s-d/2}{2}$, the numerator decays like 2^{-2tq} since $4t-2s < -d$, in which case the summation $\sum_{m \in B_q^d \setminus \{0\}} |m|^{4t-2s} \zeta_m^2$ remains bounded. The denominator remains bounded. So the overall behavior is 2^{-2tq} .
- When $\frac{s-d/2}{2} < t < s-d/2$, the numerator decays like $2^{-2tq} \times 2^{q(4t-2s+d)} = 2^{q(2t-2s+d)}$ according to Lemma 2.17. The denominator remains bounded. The overall behavior is $2^{q(2t-2s+d)}$.
- When $t > s-d/2$, the numerator behaves like $2^{q(2t-2s+d)}$, while the denominator behaves like $2^{q(2t-2s+d)}$. The overall behavior is of order 1.

These observations are consistent with what is observed in Figure 1. Based on them we deduce that the minimizer converges to $\frac{s-d/2}{2}$. The loss function exhibits symmetric behavior with respect to $\frac{s-d/2}{2}$ for $t \in (d/2, s-d)$. The detailed rigorous treatment is presented in Subsection 6.9.

2.7. Discussions. In the preceding three subsections, we have presented the consistency theory, its implication for implicit bias, as well as the tools and strategies underlying our proofs. This subsection adds to several discussions on the theory and proofs.

First, our theory applies to the torus domain. One may wonder whether these techniques can be applied to boundary conditions beyond the periodic ones. The main tool used in the proofs is Fourier's series (based on the eigenfunctions of the Laplacian operator). These are used to characterize the norm term and determinant term. We expect these techniques to generalize to other problems, such as the box with Dirichlet or Neumann boundary conditions in which the Fourier sine or cosine series are natural; the detailed analysis is left as future work. However, we need to point out that the limitation of this proof idea is that it requires a clear analytic understanding of the spectral properties of the kernel operator, i.e., its eigenfunctions. In Subsection 3.2.1, we present numerical experiments beyond this setting, which involves more challenging Laplacians with discontinuous coefficients that can model more complicated heterogeneous random fields.

Second, this section considers the regularity parameter only. In spatial statistics literature, consistency results on this parameter (for general Mat'ern type model) are very scarce and difficult. Here, we obtain a proof for the torus model, which is the main technical contribution of this paper. We will discuss the learning of other parameters in the next section, to make the story of the Mat'ern-like model on the torus more complete.

Finally, as we get two algorithms that can “consistently” learn the information of the regularity parameter when the number of data is large, a natural question is when to choose which. To answer this question, we present numerical study of the variances of both estimators for the Mat’ern-like model in the next subsection.

2.8. Variance of regularity parameter estimation. In this subsection, we compare the variance of the two estimators for recovering the regularity parameter s . We return to the experimental set-up in Subsection 2.3.2. We form the EB and KF estimators for 50 instances of different draws of the GP, normalized by the limiting optimum values s and $\frac{s-d/2}{2}$ respectively. The statistics of the two estimators are summarized in the histogram (see Figure 4). Clearly, EB exhibits smaller variance

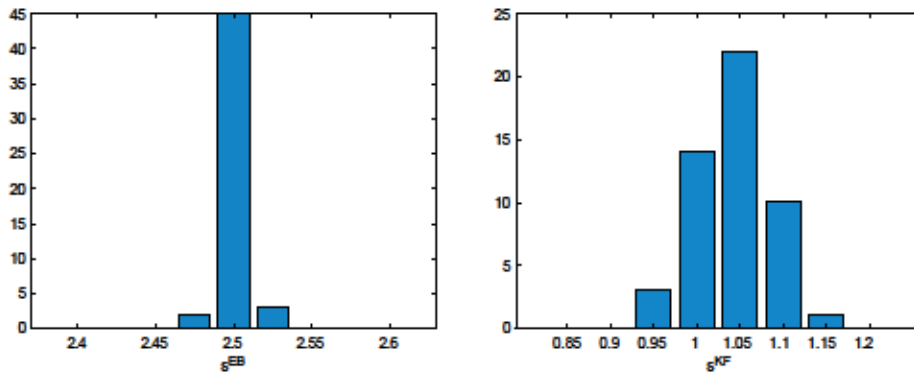


Figure 4. Histogram of the regularity estimators for the Mat’ern-like process. Left: EB; right: KF.

than KF. We compute the estimated variance using the 50 instances. Finally we get

$$\frac{\text{Var}(s^{EB})}{s^2} \approx 1.44 \times 10^{-5} \quad \text{and} \quad \frac{\text{Var}(s^{KF})}{((s - d/2)/2)^2} \approx 3.6 \times 10^{-3}.$$

Since the variance of EB is smaller, if our target is to estimate s for the exact GP model, then this suggests that the EB method is preferable.

3. More well-specified examples

The setting in Section 2 concerns regularity parameter of the Mat’ern-like model only. This section aims to extend this discussion to a wider range of settings by means of numerical experiments. First, we study the learning of lengthscale and amplitude parameters in the Mat’ern-like model in Subsection 3.1; these experiments lead to a more complete story for the Mat’ern-like model on the torus. Then, in Subsection 3.2, we consider other well-specified models, extending beyond the Mat’ern-like process example. In Subsection 3.3, we also discuss some computational aspects of the EB and KF approaches.

3.1. Recovery of amplitude and lengthscale. We start with the learning of amplitude and lengthscale parameters in the Mat’ern-like model, via either EB or KF method.

In spatial statistics, an important general principle in looking at the recovery of hyperparameters via EB is to determine whether or not the family of measures are mutually singular with respect to changes in the parameter to be estimated; learning parameters which give rise to mutually singular families is usually easy, since different almost sure properties can often be used to distinguish measures and this can be achieved without an abundance of data; in contrast those parameters that do not give rise to mutually singular measures typically require an abundance of realizations to be accurately learned. We illustrate this issue in the context of estimating one parameter by EB, the changing of which leads to mutually singular measures, and estimating two parameters by EB, changing one of which leads to mutual singularity, and the other to equivalence, for the Mat'ern-like process. We also study analogous questions about identifiability for the KF method. In all cases we work with loss functions that are natural generalizations of (2.10), (2.11).

3.1.1. Recovery of σ . A first observation is that the KF loss function is invariant under change of σ , so it cannot recover this parameter. We also note that measures are mutually singular with respect to changes in σ , and so we do expect to be able to recover σ by EB. For the EB estimator, we design the experiment as follows. We study whether the EB method can recover σ while s, τ are fixed. In detail, we consider a problem with domain the one dimensional torus T^1 . The Mat'ern-like kernel has regularity $s = 2.5$, amplitude $\sigma = 1$ and lengthscale $\tau = 0$. We assume the values of s, τ are known, but not σ . We want to recover σ by seeing a single discretized realization $u^t \sim N(0, \sigma^2(-\Delta + \tau^2 I)^{-s})$. The domain T^1 is discretized into $N = 2^{10}$ equidistributed grid points. The data we observe is the values of u^t in 2^9 equidistributed points. We build the EB loss function (see equation (3.1)) and plot the figure for a single instance; see Figure 5. We introduce ς as the variable

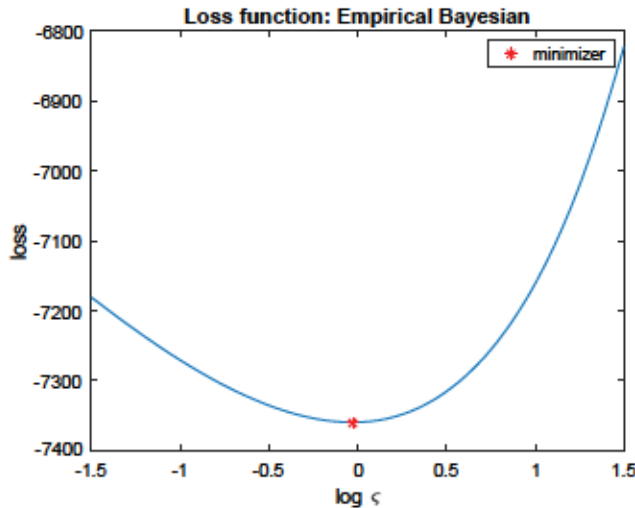


Figure 5. EB loss function for recovering σ

to be maximized over to determine our estimate of σ . In our experiments we work with the parameterization $\varsigma = \exp(\varsigma')$ in order to ensure that the estimated σ is positive. Hence, the x-axis of Figure 5 is ς' . The figure shows that the minimizer of

the loss function is close to the point $\zeta^l = 0$ ($\zeta = 1$), so the estimator σ^{EB} is close to the ground truth σ .

We can theoretically analyze the convergence. The same set-up in Subsection 2.1 is adopted, except now we assume the function is drawn from $\mathcal{N}(0, \sigma^2(-\Delta)^{-s})$ with s known and we want to recover σ by seeing the equidistributed spatial samples on the torus. After calculating the likelihood in such a case, we get the EB estimator below. Here we abuse the notation to write

$$(3.1) \quad \begin{aligned} \sigma^{\text{EB}}(q, u^\dagger) &= \underset{\zeta > 0}{\operatorname{argmin}} L^{\text{EB}}(\zeta, q, u^\dagger), \\ L^{\text{EB}}(\zeta, q, u^\dagger) &:= \frac{\sigma^2 \|u(\cdot, s, q)\|_s^2}{\zeta^2} + \log \det K(s, q) + 2^{qd} \log \zeta^2. \end{aligned}$$

The definition of $u(\cdot, s, q), K(s, q)$ is the same as in Subsection 2.1. Recall that $u(\cdot, s, q)$ is the mean of the GP found by conditioning a prior measure $\mathcal{N}(0, \sigma^2(-\Delta)^{-s})$ on observations of $u^\dagger$ at the observation data with level q . The definition of $L^{\text{EB}}(\zeta, q, u^\dagger)$ also follows from Subsection 2.1. We abuse notation to write $L^{\text{EB}}(\zeta, q, u^\dagger)$ for the EB loss function used in the estimation of σ ; the reader should not confuse this with $L^{\text{EB}}(t, q, u^\dagger)$ in Subsection 2.1 which is used for recovering the regularity parameter s .

In this setting we have the following consistency result:

Theorem 3.1. Fix $\delta > 0$. Suppose $u^\dagger$ is a sample drawn from the Gaussian process $\mathcal{N}(0, \sigma^2(-\Delta)^{-s})$ for some $s \in [d/2 + \delta, 1/\delta]$. Then, for the Empirical Bayesian estimator of σ , it holds that

$$\lim_{q \rightarrow \infty} \sigma^{\text{EB}}(q, u^\dagger) = \sigma,$$

where the convergence is in probability with respect to randomly chosen $u^\dagger$.

Proof. By taking the derivative of $L^{\text{EB}}(\zeta, q, u^\dagger)$ with respect to ζ and setting it to 0, we get the explicit formula:

$$(3.2) \quad \sigma^{\text{EB}}(q, u^\dagger) = \sigma \frac{1}{2^{qd}} \frac{\|u(\cdot, s, q)\|_s^2}{\sigma^2}.$$

Due to Proposition 2.14, we get our $\|u(\cdot, s, q)\|_s^2 = \frac{\sigma^2 \zeta^2}{m \in B_q} \frac{1}{m}$. By the Law of Large Numbers, we have

$$\lim_{q \rightarrow \infty} \frac{\|u(\cdot, s, q)\|_s^2}{2^{qd}} = 1,$$

from which the consistency follows. $\square$

Remark 3.2. We note that consistency results for the amplitude parameter have been well studied in the literature; see [33]. The purpose of this subsection is to tie those results to the rather explicit setting of our paper. One important feature of the torus model is that we are able to get an explicit and simple formula for σ^{EB} , so the consistency results are very clear. Moreover, since σ^{EB} is the average of i.i.d. Gaussian random variables, one can also easily read off other statistical properties of this estimator (although the result of asymptotic distribution is also not completely new; see for example the discussion on page 201 in [33]).

3.1.2. Recovery of s, σ simultaneously. We now build on the previous experiment to study whether the EB method can recover s, σ simultaneously when τ is fixed. We reemphasize that since the measures are mutually singular with respect to changes in σ and s we do expect to be able to recover (σ, s) by EB. The basic set-up is the same as the last subsection, and now we minimize the EB loss function to recover s, σ where, again, $\sigma = \exp(\sigma')$. We run 50 instances (each instance corresponds to a random draw of ζ), and collect the estimators $(s^{\text{EB}}, \log \sigma^{\text{EB}})$ of the EB loss function for each instance. We present the histogram of the two values obtained in the experiments as follows (Figure 6).

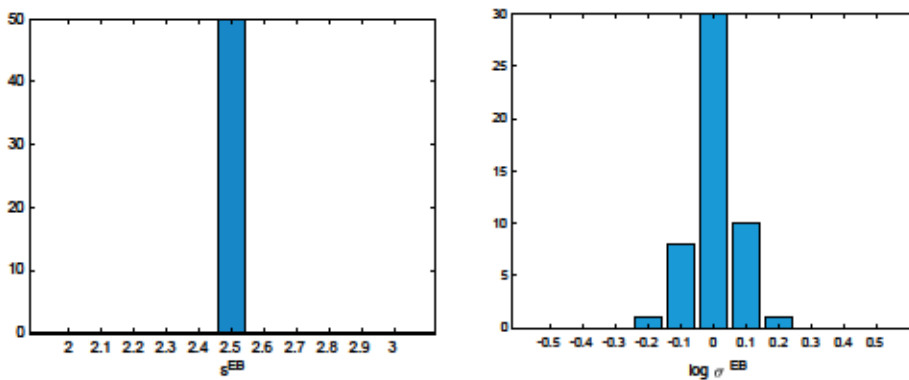


Figure 6. Left: histogram of the s^{EB} ; right: histogram of the $\log \sigma^{\text{EB}}$.

From the figure, we observe that in the 50 runs, the minimizer $(s^{\text{EB}}, \sigma^{\text{EB}})$ is close to the ground truth $(2.5, 1)$. We conclude that the EB method can recover the two parameters simultaneously in such a context.

3.1.3. Recovery of τ . We consider whether EB and KF can recover the inverse lengthscale parameter τ . We assume that σ is fixed at 1, s is chosen to be 2.5, and sample $u^t \sim \mathcal{N}(0, (\Delta + \tau^2 I)^{-s})$ with $\tau = 1$. As in the preceding experiments we consider the one dimensional torus example, and the same discretization precision and data acquisition setting as before. We draw 50 instances of u^t , and for each of them, calculate the minimizers of the EB and KF loss function. We write $\tau = \exp(\tau')$ and the estimator is $\log \tau^{\text{EB}}$ for τ' , which we constrain to be in the interval $[-2, 2]$. In the EB loss function we fix $t = s$ within the loss function; for the KF method, we select $t = s$ (case 1) and $t = \frac{s-d/2}{2}$ (case 2) respectively within the loss function. The histograms of the minimizers of the resulting EB loss function and KF loss functions (in both cases) are presented in Figure 7, expressed in terms of $\log \tau^{\text{EB}}$ and $\log \tau^{\text{KF}}$. In the 50 runs, the EB estimator takes many different values with no apparent pattern. For both case 1 and case 2, the KF estimator of τ' takes the value 2 very often, which is the maximal value of the constrained decision variable. None of the estimators recover the true $\tau' = 0$.

The behavior of the KF estimator can be explained by the observation that when τ increases, the function drawn from the Gaussian prior becomes smoother, and hence the subsampling step in the KF loss does not sacrifice too much information. Therefore, the KF loss exhibits a tendency to get smaller as τ increases. We

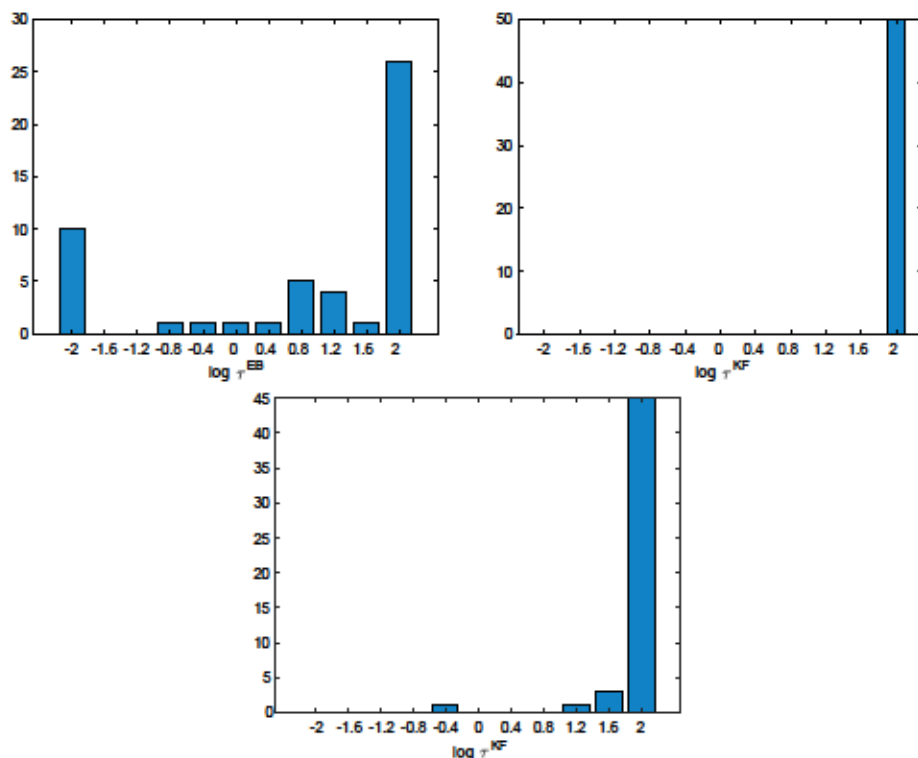


Figure 7. Histogram of the $\log \tau^{\text{EB}}$ or $\log \tau^{\text{KF}}$. Upper left: EB loss; upper right: KF loss (case 1); bottom: KF loss (case 2).

can understand why EB cannot recover τ by studying the equivalence of Gaussian measures. As shown in [9], when dimension $d \leq 3$, the Gaussian measures $\mathcal{N}(\mathbf{0}, (-\Delta + \tau^2 I)^{-s})$ for different τ are equivalent; thus one cannot expect to recover τ using the information from one sample.

We can also consider the problem of recovering s, τ simultaneously, i.e., we solve a joint minimization problem to get $s^{\text{EB}}, \log \tau^{\text{EB}}$ and $s^{\text{KF}}, \log \tau^{\text{KF}}$. The set-up is the same as above, with the sample drawn from $\mathcal{N}(\mathbf{0}, (-\Delta + \tau^2 I)^{-s})$ for $\tau = 1$ and $s = 2.5$. We form the EB and KF loss for 50 instances of different draws and find the minimizers as corresponding estimators. The histograms of the estimators are shown in Figures 8 and 9. These figures show that in this joint optimization, the EB method picks the correct value $s^{\text{EB}} = 2.5$ for estimating s , and exhibit no patterns for $\log \tau^{\text{EB}}$; the KF method finds values close to 1 for s^{KF} , as it would in the absence of simultaneous estimation of τ , and selects the largest possible value in the constraint for $\log \tau^{\text{KF}}$, here being 2. The conclusion is that the fact that τ cannot be learned accurately does not influence the estimation of the regularity parameter s in a context in which the two are learned simultaneously. Indeed, this conclusion also holds when we are recovering the three parameters (s, σ, τ) simultaneously.

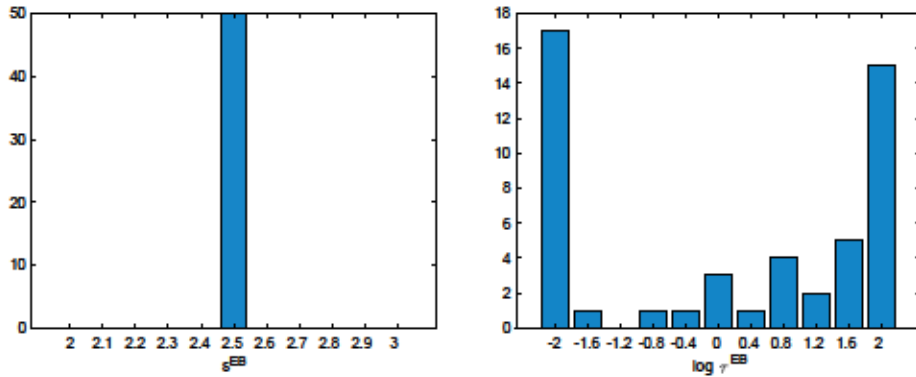


Figure 8. EB approach. Left: histogram of the s^{EB} ; right: histogram of the $\log \tau^{\text{EB}}$.

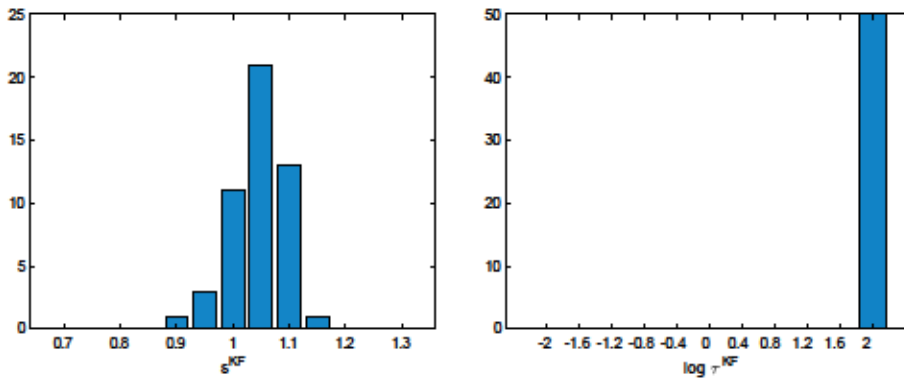


Figure 9. KF approach. Left: histogram of s^{KF} ; right: histogram of the $\log \tau^{\text{KF}}$.

3.2. Other well-specified examples. In this subsection, we consider numerical examples for recovering parameters of a random field in the well-specified case, going beyond the Mat'ern process studied thus far.

3.2.1. Recovery of regularity parameter for variable coefficient elliptic operator. Set $D = [0, 1]$ so that $d = 1$. The theoretical result in Section 2 assumes the function observed $u^\dagger$ is drawn from $\mathcal{N}(0, (-\Delta)^{-\beta})$ on a torus. In this subsection, we assume $u^\dagger$ is drawn from $\mathcal{N}(0, (-\nabla \cdot (a \nabla))^{-\beta})$ for some non-constant function a , and that the elliptic operator implicit in this definition of a Gaussian measure is equipped with homogeneous Dirichlet boundary condition on D . We observe its values on the 2^9 equidistributed points of the total 2^{10} grid points used for discretization.

Here we select a coefficient $a(x)$ that exhibits a discontinuity at $x = 1/2$:

$$(3.3) \quad a(x) = \begin{cases} 1 & x \in [0, 1/2] \\ 2 & x \in (1/2, 1] \end{cases}$$

As a consequence the induced operator is not the Laplacian. We pick $\beta = 2.5$ to draw a sample $u^\dagger$.

In the well-specified case, the GP used in defining the EB and KF estimators is parameterized by $\mathcal{N}(\mathbf{o}, (-\nabla \cdot (a \nabla \cdot))^{-1})$ and we aim to learn parameter t given a data calculated using a draw from the same measure with $t = s$. We consider the well-specified case here (the misspecified case will be considered in Subsection 4.1). We output the histogram of the EB and KF estimators for 50 different draws of $u^\dagger$ in Figure 10. The experiments show that for the variable coefficient elliptic

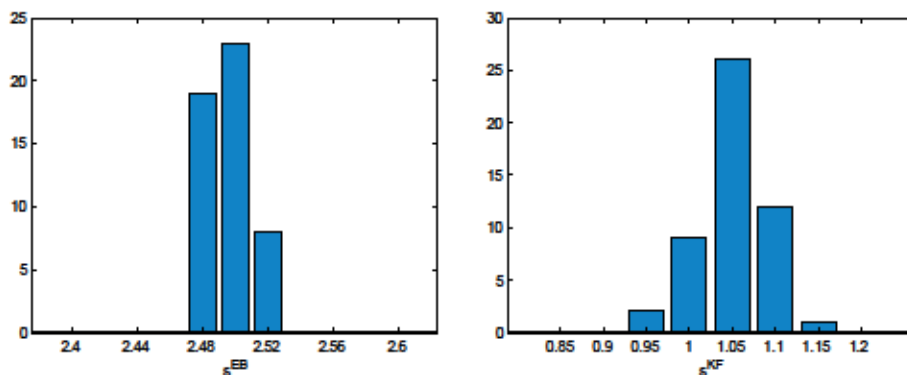


Figure 10. Histogram of the regularity estimators for the variable coefficient covariance case. Left: EB; right: KF.

operator model, EB and KF succeed in converging to the correct limits. We can calculate the (normalized) variance of the two estimators based on the histograms:

$$\frac{\text{Var}(s^{\text{EB}})}{s^2} \approx 7.8 \times 10^{-5} \quad \text{and} \quad \frac{\text{Var}(s^{\text{KF}})}{((s - d/2)/2)^2} \approx 4 \times 10^{-3}.$$

The relative magnitude is similar to the one in Subsection 2.8.

3.2.2. Recovery of discontinuity position for conductivity field. Define the conductivity field $a_\theta : [0, 1] \rightarrow \mathbb{R}$, and parameterized by $\theta \in [0, 1]$, via

$$(3.4) \quad a_\theta(x) = \begin{cases} 1 & x \in [0, \theta] \\ 2 & x \in (\theta, 1]. \end{cases}$$

In this subsection, we assume that our data $u^\dagger$ is obtained by solving the SPDE

$$-\nabla \cdot (a_{1/2} \nabla u^\dagger) = \xi,$$

subject to a homogeneous Dirichlet boundary condition on $[0, 1]$. We choose ξ as a random draw from $\mathcal{N}(\mathbf{o}, (-\Delta)^{-1})$. We can view $u^\dagger$ as a sample drawn from $\mathcal{N}(\mathbf{o}, C_a)$ where

$$(3.5) \quad C_a = (-\nabla \cdot (a_{1/2} \nabla \cdot))^{-1} (-\Delta)^{-1} (-\nabla \cdot (a_{1/2} \nabla \cdot))^{-1}.$$

We observe the value of $u^\dagger$ on the 2^9 equidistributed points of the total 2^{10} grid points used for discretization. We use EB and KF to estimate θ from the partial observation of the function $u^\dagger$ based on the GP model $\mathcal{N}(\mathbf{o}, C_{a,s})$ where

$$(3.6) \quad C_{a,s} = (-\nabla \cdot (a_\theta \nabla \cdot))^{-1} (-\Delta)^{-1} (-\nabla \cdot (a_\theta \nabla \cdot))^{-1}.$$

The model is well-specified for $s = 1$ and misspecified for $s \neq 1$. Here consider the well-specified case in this subsection, i.e., $s = 1$, and $C_{\alpha,s} = C_\alpha$; the misspecified case is covered in Subsection 4.2.

We let the domain for θ be $[0.3, 0.7]$ in the definition of EB and KF estimators. We compute the estimators for 50 different draws of $u^\dagger$. The histograms of the EB and KF estimators are shown in Figure 11. The loss functions for one random instance are shown in Figure 12.

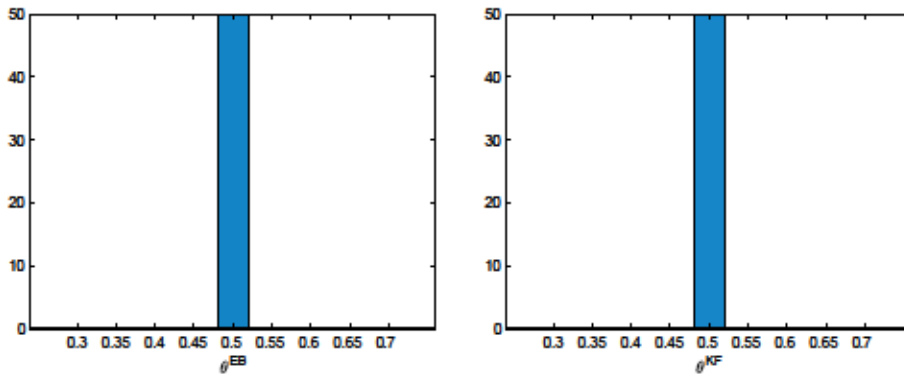


Figure 11. Histogram of the discontinuity position estimators (well-specified). Left: EB; right: KF.

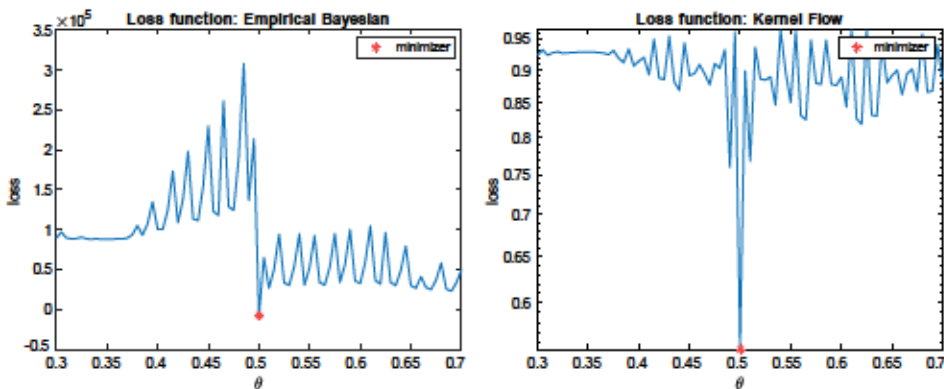


Figure 12. Loss function for recovering the discontinuity (well-specified). Left: EB; right: KF.

Our experiments show that both EB and KF can recover $\theta = 1/2$, and the recovery is very stable with respect to different draws of $u^\dagger$ from the SPDE. We conclude that the EB and KF can go beyond the Mat'ern-like kernel model in practice; recovering the point of discontinuity of the conductivity field is an example of this fact.

3.3. Computational aspects. In this subsection, we add some discussions about the computational aspects. We start by remarking on how to compute the kernel function and sample the GP realization generally. Every kernel operator we consider

involves certain differential operators. We discretize these differential operators and perform an eigenfunction decomposition of the obtained matrix. Then we use these eigenfunctions and eigenvalues to compute approximation of the kernel matrix, and draw samples from the GP with the covariance matrix being the kernel matrix; see also discussions in Remark 2.2. This is similar to the spectral expansion of a kernel function and the Mercer decomposition of a GP.

Practical applications of hierarchical GPR require weighing statistical efficiency against computational complexity. Although the regularity models covered in this paper appear to produce well-behaved EB and KF loss functions with easily identifiable global minimizers, models with high dimensional parameter space typically require using algorithms such as gradient descent which do not come with theoretical guarantees on the identification of global minimizers. Furthermore, when the size of the data is large, computation becomes a limiting factor, and subsampling offers a traditional remedy when combined with gradient descent, but again theoretical guarantees are not typically to be expected. The stochastic algorithm presented in [25] for KF can be interpreted as an SGD algorithm aimed at minimizing the average loss

$$\mathbb{E}_{\pi_1} \mathbb{E}_{\pi_2} L^{\text{KF}}(\theta, \pi_1 X, \pi_2 \pi_1 X, u^\dagger),$$

via draws from the distribution of π_1 and π_2 ($\pi_1 X$ is a random subsampling of X , and $\pi_2 \pi_1 X$ is a further random subsampling of $\pi_1 X$). The efficacy of an analogous strategy for EB remains unclear due to the presence of the log determinant term in the loss. It is of future interest to explore further the computational aspects of the EB and KF approaches to hierarchical learning.

4. Model misspecification

All our preceding experiments are focused on the well-specified case: the function $u^\dagger$ is drawn from the GP model assumed in the estimation, or equivalently, the model for $u^\dagger$ and for the kernel family K_θ in defining the loss functions are matched. This subsection studies model misspecification. We consider two possible ways to misspecify the model: (1) the function $u^\dagger$ is drawn from a GP which is different from that used in defining the loss function; (2) the function $u^\dagger$ is a fixed deterministic function. The second case may arise, for example, if the function comes from a solution of a PDE with some physical data, and there is no natural stochastic context for its provenance. The aim of this subsection is to study the behavior of the EB and KF estimators to compare their robustness to model misspecification.

4.1. Stochastic model misspecification for recovering regularity. In this subsection, we assume $u^\dagger$ is drawn from $\mathcal{N}(\mathbf{0}, (\nabla \cdot (\mathcal{Q} \cdot))^{-1})$, while the GP used in defining the EB and KF estimators is still $\mathcal{N}(\mathbf{0}, (-\Delta)^{-1})$. This results in a model misspecification corresponding to the well-specified model in Subsection 3.2.1. As in Subsection 3.2.1, we select σ as in (3.3) and we set $s = 2.5$ to draw the sample $u^\dagger$. Figure 13 shows the histograms of the minimizers of the EB and KF loss functions obtained from 50 independent draws from the Gaussian Process. Despite misspecification, the EB and KF estimators are still concentrated around 2.5 and 1, respectively. We also calculate the variance:

$$\frac{\text{Var}(s^{\text{EB}})}{s^2} \approx 5.9 \times 10^{-4} \quad \text{and} \quad \frac{\text{Var}(s^{\text{KF}})}{((s - d/2)/2)^2} \approx 6.8 \times 10^{-4}.$$

In this example, the (normalized) variances of KF of EB are of similar magnitude. This is different from the well-specified case in Subsection 3.2.1 where the variance of EB is much smaller than KF.

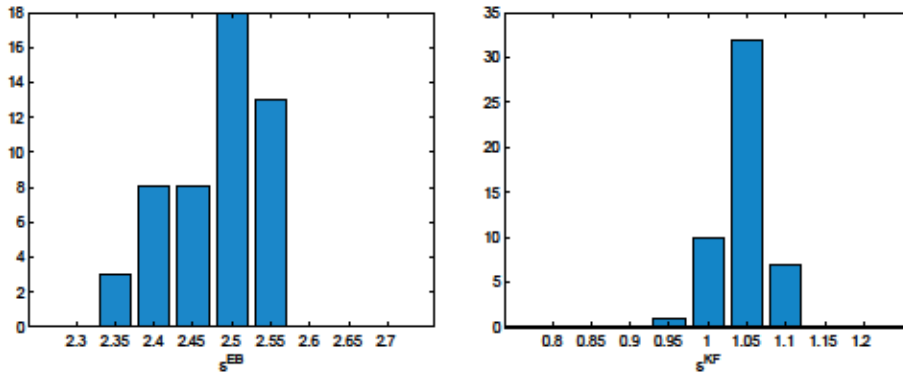


Figure 13. Histogram of the regularity estimators under model misspecification. Left: EB; right: KF.

4.2. Stochastic model misspecification for recovering discontinuity. In this subsection, we consider the model misspecifications that correspond to the well-specified case in Subsection 3.2.2. For the GP defining the EB and KF estimators we use the centred Gaussian with covariance operator given by (3.6) with $s = 5$; meanwhile u^t is drawn from the centred Gaussian with covariance operator given by (3.5); thus we are in a misspecified version of the setting arising in Subsection 3.2.2 and, as there, our aim is to recover the point of discontinuity. We illustrate the loss functions for a single draw of u^t in Figure 14. These plots are not sensitive to the particular draw of u^t and illustrate the robustness of KF (and the lack of robustness of EB) to this misspecification. Indeed, the EB estimator gives 0.3 which is the lower boundary of the compact parameter space used in the minimization, while the KF estimator picks the true parameter 0.5. The loss function of KF, shown in Figure 14, exhibits a sharp global minimizer at $\theta = 0.5$.

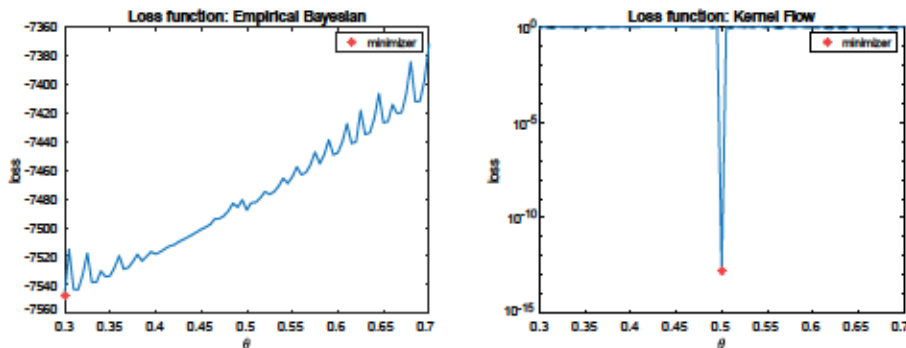


Figure 14. Loss function for estimating the discontinuity parameter under model misspecification. Left: EB; right: KF.

4.3. Deterministic model. In this subsection, we consider the EB and KF estimators for the parameter t in the GP model $\mathcal{N}(0, (-\Delta)^{-s})$ where Δ is equipped with homogeneous Dirichlet boundary conditions on $[0, 1]$. However, rather than choosing $u^\dagger$ that is drawn from the $\mathcal{GP}_{\mathcal{N}}(0, (-\Delta)^{-s})$ for some s (as we did in Section 2), we choose it be the solution to the equation $(-\Delta)^s u^\dagger(\cdot) = \delta(\cdot - 1/2)$, i.e., $u^\dagger$ is the Green function corresponding to the differential operator $(-\Delta)^s$ and evaluated at $y = 1/2$. Since $u^\dagger$ has no stochastic background, we understand this situation as a deterministic model misspecification.

We observe the value of $u^\dagger$ on the 2^9 equidistributed points of the total 2^{10} grid points used for discretization. We conduct numerical experiments to find the value of the EB and KF estimators. Our experiments show that the EB estimator returns $2s$ and the KF estimator returns s for this one dimensional example. The loss function in the case $s = 1.2$ is shown in Figure 15.

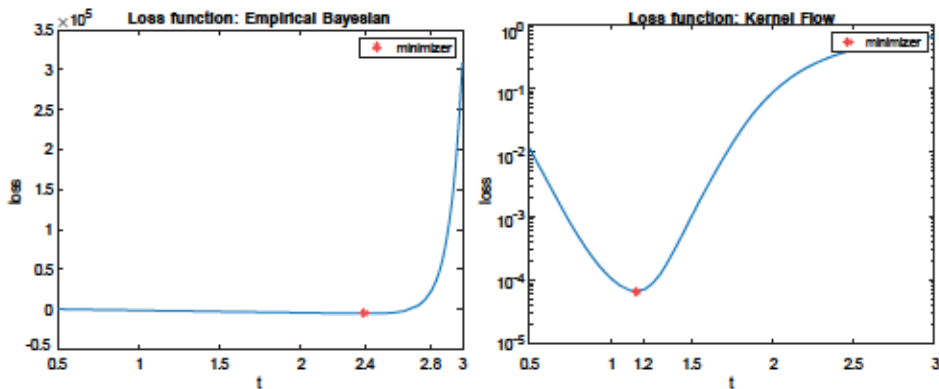


Figure 15. Loss function for estimating the regularity parameter under deterministic $u^\dagger$. Left: EB; right: KF.

We now describe some regularity considerations in order to understand the observed phenomenon. In this one dimensional example, $\delta(\cdot - 1/2)$ belongs to $H^\eta([0, 1])$ for any $\eta < -1/2$, so the solution $u \in H^{2s+\eta}([0, 1])$ for any $\eta < -1/2$. It is of critical regularity $2s-1/2$, but this criticality is not homogeneous: it is caused by the presence of a singularity induced by the Dirac function.

The discussion in Section 2 implies KF will recover $s - 1/4$ while EB recovers $2s$ for a function with homogeneous critical regularity $2s - 1/2$. However, the experiments here show that KF recovers s while EB recovers $2s$, for this function with critical regularity $2s - 1/2$; unlike the setting in Section 2, here the ground truth lacks spatial homogeneity. This suggests that the KF estimator for the regularity parameter is sensitive to whether the regularity of the target function is spatially homogeneous or not. This fact is not surprising, considering the vast literature on adaptive approximation for functions with singularities, which implies the presence of a singularity will exert considerable influence on the approximation error resulting from minimizing the KF loss function. In this example, the optimal approximation in KF error comes at $t = s$. We can understand this phenomenon as follows. Recall $u^\dagger = (-\Delta)^{-s} \delta(\cdot - 1/2)$. Using $\mathcal{N}(0, (-\Delta)^{-s})$ in the GPR is equivalent to using the basis functions $\text{span}_{j \in J} \{(-\Delta)^{-s} \delta(\cdot - x_j)\}$ (as in Section 2.1) with x_i being the data points indexed by $j \in J_q$ to approximate $u^\dagger$. When $t = s$ and one of the $x_j = 1/2$,

the ground truth will just be in the basis functions set, so it is straightforward to imagine $t = s$ leads to the smallest approximation error, and KF picks this value.

We understand the fact that EB still picks $t = 2s$ by making the following observation: there are only two terms in the EB loss function. The log determinant term remains the same for each t when $u^\dagger$ changes. For the norm term $\|u(\cdot, t, q) - u^\dagger\|_t^2$, the blow-up rate depends on the regularity of $u^\dagger$. Here, it makes no difference whether the regularity of $u^\dagger$ is spatially homogeneous or not.

4.4. Discussions. The above numerical experiments reveal complicated behavior of EB and KF with respect to model misspecification. In the second experiment, we found that KF is robust while EB is not, for a certain type of GP model misspecification. This appears natural since EB is based on probabilistic modeling whilst KF is purely based on approximation theoretic criteria. In Subsection 4.2 the prior used in EB is mutually singular with respect to the GP that $u^\dagger$ is drawn from and it is not surprising that EB is fragile. On the other hand, KF does not require probabilistic modeling to motivate it, and so its robustness to misspecifications behaves differently. Indeed, in the second experiment, the discontinuity point influences the approximation accuracy a lot, and even the kernel used in defining KF is misspecified, KF still succeeds in selecting the correct parameter, as it focuses on the approximation accuracy rather than statistical inference.

In the well-specified cases, e.g. experiments in Section 2, EB outperforms KF in terms of the variance of estimators. Therefore, if $u^\dagger$ is a random object and we know the prior correctly, then EB should be a preferable choice for estimating parameters. If this is not the case and misspecification occurs, EB might be vulnerable and KF could be a potential alternative.

5. Concluding remarks

In this paper, we have studied the Empirical Bayes and Kernel Flow approaches to hyperparameter learning. The first approach is based on statistical considerations, while the second approach originates from an approximation theoretic viewpoint. Their distinct objectives lead them to different behaviors and different interpretations of optimality.

For the Mat'ern-like process model, we made a detailed theoretical study of the recovery of the regularity parameter. We proved the EB estimator converges to s , while the KF estimator converges to $\frac{s-d/2}{2}$, both results holding in probability in the large data limit if the regularity of the GP that $u^\dagger$ draws from is s . Our experiments illustrate that, in terms of the L^2 error $\|u(\cdot, t, q) - u^\dagger\|_0^2$, the parameter $t = \frac{s-d/2}{2}$ relates to the minimal t that achieves the fast error rate while $t = s$ relates to the t that achieves the smallest error, averaged over the GP $\mathcal{M}_N^s(\mathbf{o}, \mathbf{K}(\Delta)^{-s})$. This demonstrates the different drivers that guide the EB and KF methods in selecting the parameters. The statistical and approximation theoretic principles behind them lead to the differences between them.

In the theoretical study, we developed a Fourier analysis toolkit for this problem, and as a byproduct, we showed the consistency of recovering σ in the Mat'ern-like process for the EB method. Recovery of the lengthscale parameter and recovery of several parameters simultaneously was studied via numerical experiments. It is of future interest to perform theoretical studies explaining these empirically observed phenomena. Furthermore, the theory in this paper is based on an equidistributed

design for the data location, and the generalization to randomized design remains a potential further direction. Also, our focus in this paper is on the noiseless observation setting, and an extension to the noisy case is of future theoretical interest.

Our numerical experiments for additional well-specified and misspecified models extend the scope of this paper beyond the Mat'ern-like kernels. Both the two estimators work very well in the well-specified models we consider; we would like to explore this more in the future, both theoretically and numerically, potentially in more complex models that are present in machine learning. The variance and robustness of the estimators behave differently for the misspecified models. The variabilities in robustness are in line with our expectation since these estimators follow from different decision rules; these rules can vary considerably in sensitivity to model mismatches of different kinds. In practice, users should choose the correct approach to avoid high sensitivity to likely model errors present.

As a summary, this paper demonstrates some basic aspects of the difference between Bayesian and approximation theoretic approaches for hierarchical learning. Generally, it is of interest to study EB and KF for other types of models and to study other parameter selection criteria based on the two principles beyond EB and KF, such as a fully Bayesian approach or another choice of d for the approximation, and identify their pros and cons under different scenarios. We are interested in exploring the theoretical and practical performance of methods under such a framework, and we believe that a diversity in such methods will enable users to deal with the model misspecification that is to be expected in many applications.

6. Appendix: Proofs

6.1. Proof of Proposition 2.11.

Proof. Let $\phi_j(x) = (-\Delta)^{-t} \delta(x - x_j)$ and in particular $\phi_0(x) = (-\Delta)^{-t} \delta(x)$. We have for $m \in \mathbb{Z}^d$,

$$\hat{\phi}_0(m) = \begin{cases} (4\pi^2)^{-t} |m|^{-2t}, & m \neq 0 \\ 0, & m = 0. \end{cases}$$

We introduce the translation operator $\tau_{j2^{-q}}$ which acts on function $u : \mathbb{T}^d \mathbb{R}$ and is defined by

$$(\tau_{j2^{-q}} u)(x) = u(x_1 - j_1 2^{-q}, x_2 - j_2 2^{-q}, \dots, x_d - j_d 2^{-q})$$

for $j = (j_1, j_2, \dots, j_d) \in \mathbb{Z}^d$ and $x = (x_1, x_2, \dots, x_d) \in \mathbb{R}^d$. Then, for $j \in J_q$, we have the relation $\delta(\cdot - x_j) = \tau_{j2^{-q}} \delta(\cdot)$. Using the property of the Fourier coefficients, we obtain

$$\hat{\phi}_j(m) = \hat{\phi}_0(m) e^{-2\pi i(j2^{-q}, m)} = \begin{cases} (4\pi^2)^{-t} |m|^{-2t} e^{-2\pi i(j2^{-q}, m)}, & m \neq 0 \\ 0, & m = 0. \end{cases}$$

By definition, $\hat{F}_{t,q}$ is the span of such $\hat{\phi}_j$ for $j \in J_q$. Hence, for any $g \in \hat{F}_{t,q}$, it can be written as a linear combination of these functions. Equivalently, there exists a 2^q -periodic function p such that

$$g(m) = \begin{cases} |m|^{-2t} p(m), & m \neq 0 \\ 0, & m = 0. \end{cases}$$

This gives the desired representation of g . □

6.2. Proof of Theorem 2.13.

Proof. By Proposition 2.11, there exists a 2^q -periodic function $p_1(m)$ on Z^d , such that

$$\hat{u}(m, t, q) = \begin{cases} |m|^{-2t} p_1(m), & m \neq \mathbf{o} \\ \mathbf{o}, & m = \mathbf{o}. \end{cases}$$

By the definition of GPR, we have $[u^\dagger(\cdot) - u(\cdot, t, q), \delta(\cdot - x_j)] = \mathbf{o}$ for every data point x_j . In the Fourier domain, according to the characterization of $\hat{f}_{t,q}$, this orthogonality leads to

$$(6.1) \quad \sum_{m \in Z^d} (\hat{u}(m) - \hat{u}(m, t, q)) p(m) = \mathbf{o}$$

for $p : Z^d \rightarrow \mathbb{C}$ being any 2^q -periodic function. Recalling Definition 2.12, we have

$$(6.2) \quad (T_q \hat{u})(m) = \sum_{\beta \in Z^d} \hat{u}(m + 2^q \beta).$$

The fact that the above sum converges may be seen as a consequence of the Cauchy–Schwarz inequality and the regularity of u (recall $\frac{d}{2} + \delta$). Using (6.2) and the representation of $\hat{u}(m, t, q)$, we reformulate (6.1) as

$$\sum_{m \in B_q^d} (T_q \hat{u})(m) - M_q^t(m) p_1(m) p(m) = \mathbf{o}.$$

The above formula holds for any 2^q -periodic function p . Let $g(m) = (T_q \hat{u})(m) - M_q^t(m) p_1(m)$, then we get that g is a 2^q -periodic function on Z^d and that

$$\sum_{m \in B_q^d} g(m) p(m) = \mathbf{o}$$

holds for any 2^q -periodic function p . This implies that $g(m) = \mathbf{o}$. Hence, we get

$$p_1(m) = \frac{(T_q \hat{u})(m)}{M_q^t(m)}.$$

Plugging this expression into the above representation formula for $\hat{u}(m, t, q)$ leads to

$$\hat{u}(m, t, q) = \begin{cases} \mathbf{o}, & \text{if } m = \mathbf{o} \\ |m|^{-2t} \frac{(T_q \hat{u})(m)}{M_q^t(m)}, & \text{else.} \end{cases}$$

This completes the proof. $\square$

6.3. Proof of Lemma 2.15.

Proof. Recall the definition

$$M_q^t(m) := \begin{cases} \sum_{\beta \in Z^d \setminus \{0\}} |2^q \beta|^{-2t}, & \text{if } m = j \cdot 2^q \text{ for some } j \in Z^d \\ \sum_{\beta \in Z^d} |m + 2^q \beta|^{-2t}, & \text{else.} \end{cases}$$

Because of the periodicity of M_q^t , we need only to study $m \in B_q^d$. We split it into two cases.

$$(1) \text{ If } m = \mathbf{o}, \text{ then } M_q^t(m) = \sum_{\beta \in Z^d \setminus \{0\}} |2^q \beta|^{-2t} = 2^{-2qt} \sum_{\beta \in Z^d \setminus \{0\}} |\beta|^{-2t} \sim 2^{-2qt}.$$

(2) If $m \in B_q^d \setminus \{0\}$, then $M_q^t(m) = \prod_{\beta \in \mathbb{Z}^d} |m + 2^q \beta|^{-2t} = |m|^{-2t} \prod_{\beta \in \mathbb{Z}^d \setminus \{0\}} |m + 2^q \beta|^{-2t}$. Since $B_q^d = [-2^{q-1}, 2^{q-1} - 1]^d$, each component of $m \in B_q^d$ is bounded by 2^{q-1} in amplitude, and therefore each component of $2^{-q}m$ is bounded by $1/2$ in amplitude. So, it follows that

$$\prod_{\beta \in \mathbb{Z}^d \setminus \{0\}} |m + 2^q \beta|^{-2t} = 2^{-2qt} \prod_{\beta \in \mathbb{Z}^d \setminus \{0\}} |2^{-q}m + \beta|^{-2t} \sim 2^{-2qt}.$$

Then, we get $|m|^{-2t} \leq M_q^t(m) \leq |m|^{-2t} + 2^{-2qt} \leq |m|^{-2t}$ where we have used the fact that $|m| \leq 2^q$. Therefore, it holds that $M_q^t(m) \sim |m|^{-2t}$.

As a byproduct of the above proof, we also get $M_q^t(m) - |m|^{-2t} \sim 2^{-2qt}$. $\square$

6.4. Proof of Lemma 2.17.

Proof. First, we prove the pointwise convergence (i.e., for each fixed r), then move on to prove uniform convergence. To achieve this, we calculate the variance:

$$\begin{aligned} \text{Var}(a(r, q)) &\sim 2^{-2rq} \prod_{m \in B_q^d \setminus \{0\}} |m|^{2r-2d} \\ &\leq 2^{-2rq} \prod_{m \in B_q^d \setminus \{0\}} \int_1^{2^q} x^{2r-2d+1} dx = 2^{-2rq} \prod_{m \in B_q^d \setminus \{0\}} \int_1^{2^q} x^{2r-d-1} dx. \end{aligned}$$

For $r = d/2$, the integral gives $\log(2^q) = q \log 2$; for $r \neq d/2$, it is $\frac{1}{2r-d} (2^{q(2r-d)} - 1)$.

In both cases, we have $\lim_{q \rightarrow \infty} \text{Var}(a(r, q)) = 0$. Thus, $a(r, q)$ converges in L^2 to the limit of its expectation, which we may calculate as follows:

$$\lim_{q \rightarrow \infty} \mathbb{E} a(r, q) = \lim_{q \rightarrow \infty} \prod_{m \in B_q^d \setminus \{0\}} (2^{-q})^d |2^{-q}m|^{r-d} = \int_{[-1/2, 1/2]^d} |y|^{r-d} dy := \gamma(r) > 0.$$

Hence, we get $\lim_{q \rightarrow \infty} a(r, q) = \gamma(r) > 0$ in L^2 for every $r \in [e, 1/e]$, and the convergence also holds in probability. We may now proceed to show uniform convergence. We rely on Exercise 3.2.3 in [38]. Based on that, it suffices to prove $a(r, q)$ is uniformly Lipschitz continuous as a function of r for $q \in \mathbb{N}$. Pick any $r_1, r_2 \in [e, 1/e]$, then

$$\begin{aligned} &|a(r_1, q) - a(r_2, q)| \\ &= \prod_{m \in B_q^d \setminus \{0\}} |2^{-q}| (|2^{-q}m|^{r_1-d} - |2^{-q}m|^{r_2-d})| \\ &\leq \prod_{m \in B_q^d \setminus \{0\}} 2^{-qd} |r_1 - r_2| (|2^{-q}m|^{-d} + |2^{-q}m|^{1/-d}) |\log(2^{-q}m)| \zeta_{mr}^2 \end{aligned}$$

where in the last step we have used the fact that $||2^{-q}m|^{r_1-d} - |2^{-q}m|^{r_2-d}| = ||2^{-q}m|^{r-d} \log(2^{-q}m)(r_1 - r_2)|$ for some η that lies between r_1 and r_2 , and we use the bound $r_1, r_2 \in [e, 1/e]$. Now, we define the random series:

$$L(q) := 2^{-qd} \prod_{m \in B_q^d \setminus \{0\}} (|2^{-q}m|^{-d} + |2^{-q}m|^{1/-d}) |\log(2^{-q}m)| \zeta_m^2.$$

We calculate its variance as follows:

$$\begin{aligned} \text{Var}(L(q)) &\sim 2^{-2dq} \left(|2^{-q}m|^{2-2d} + |2^{-q}m|^{2/-2d} \log^2 |2^{-q}m| \right) \\ &\quad \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} \frac{1}{2^{-q}} t^{2-2d+d-1} \log^2 t \, dt + \frac{1}{2^{-q}} t^{2/-2d+d-1} \log^2 t \, dt \\ &= 2^{-q} \int_{\mathbb{R}^d} \frac{1}{2^{-q}} (t^{2-d-1} + t^{2/-d-1}) \log^2 t \, dt \\ &\quad \int_{\mathbb{R}^d} \frac{1}{2^{-q}} (t^{-d-1} + t^{1/-d-1}) \, dt \lesssim 2^{-q}. \end{aligned}$$

The last term will go to 0 as q goes to infinity. Thus, $L(q)$ converges in L^2 (and thus in probability) to $L^* = \lim_{q \rightarrow \infty} EL(q)$, which is

$$\begin{aligned} \lim_{q \rightarrow \infty} EL(q) &= \lim_{q \rightarrow \infty} 2^{-qd} \int_{\mathbb{R}^d} \frac{1}{2^{-q}} (|2^{-q}m|^{-d} + |2^{-q}m|^{1/-d}) \log^2 |2^{-q}m| \\ &= \int_{\mathbb{R}^d} (|y|^{-d} + |y|^{1/-d}) \log^2 |y| \, dy \\ &\lesssim \int_{[-1/2, 1/2]^d} (|y|^{-2-d} + |y|^{1/(2)-d}) \, dy < \infty. \end{aligned}$$

Using Markov's inequality we deduce that, for any $\epsilon' > 0$, it holds that

$$\mathbb{P}(|L(q) - L^*| \geq \epsilon') \leq \frac{\mathbb{E}|L(q) - L^*|^2}{(\epsilon')^2} \leq \frac{2^{-q}}{(\epsilon')^2}.$$

Thus,

$$\sum_{q=1}^{\infty} \mathbb{P}(|L(q) - L^*| \geq \epsilon') \leq \sum_{q=1}^{\infty} \frac{2^{-q}}{(\epsilon')^2} < \infty.$$

From the Borel-Cantelli lemma it follows that $\lim_{q \rightarrow \infty} L(q) = L^*$ almost surely, and therefore $L(q)$ is bounded uniformly for q almost surely. Since $|a(r_1, q) - a(r_2, q)| \leq L(q)|r_1 - r_2|$, it follows that $a(r, q)$ is uniformly Lipschitz continuous as a function of r for $q \in \mathbb{N}$. Invoking Exercise 3.2.3 in [38] concludes this case.

For the case $r = 0$, we follow the same strategy as in the previous case. First, we calculate the corresponding variance:

$$\begin{aligned} \text{Var}(a(0, q)) &\sim \frac{1}{q^2} \int_{\mathbb{R}^d} \frac{1}{2^q} |m|^{-2d} \\ &\quad \int_{\mathbb{R}^d} \frac{1}{q^2} x^{-2d+d-1} \, dx \lesssim \frac{1}{q^2}, \end{aligned}$$

where the last term goes to 0 as q goes to infinity. Then, we calculate the expectation:

$$\mathbb{E}a(0, q) = \frac{1}{q} \int_{\mathbb{R}^d} \frac{1}{2^q} |m|^{-d} \, dm.$$

The limit when $q \rightarrow \infty$ is identified through the following calculations:

$$\begin{aligned} \lim_{q \rightarrow \infty} \frac{1}{q} \int_{m \in B_q^d \setminus \{0\}} |m|^{-d} &= \lim_{q \rightarrow \infty} \int_{m \in B_{q+1}^d \setminus B_q^d} |m|^{-d} \\ &= \lim_{q \rightarrow \infty} \int_{m \in B_{q+1}^d \setminus B_q^d} |2^{-q} m|^{-d} \\ &= \int_{[-1,1]^d \setminus [-1/2,1/2]^d} |x|^{-d} dx < \infty; \end{aligned}$$

here we have used the definition of the Riemann integral. Finally, we conclude that $\lim_{q \rightarrow \infty} a(o, q) = \gamma(o)$ in probability for $\gamma(o) \in (o, \infty)$. $\square$

6.5. Proof of Proposition 2.18.

Proof. First, we have the relation

$$\log \det K(t, q) = \log \det K(t, q-1) + \log \det(K(t, q)/K(t, q-1)),$$

where $K(t, q)/K(t, q-1)$ is the Schur complement of $K(t, q-1)$ in $K(t, q)$. Due to the variational property of the Schur complement (see Lemma 13.24 in [24]), the smallest and largest eigenvalues of $K(t, q)/K(t, q-1)$ satisfy (in the dual norm $\|\cdot\|_{-t}$)

$$\begin{aligned} \lambda_{\min}(K(t, q)/K(t, q-1)) &\geq \inf_{y \in \mathbb{R}^{|J_q|}} \frac{\|\sum_{j \in J_q} y_j \delta(x - x_j)\|_{-t}^2}{\|y\|^2}, \quad \text{and} \\ (6.3) \quad \lambda_{\max}(K(t, q)/K(t, q-1)) &= \sup_{y \in \mathbb{R}^{|J_q|}} \inf_{z \in \mathbb{R}^{|J_{q-1}|}} \frac{\|\sum_{j \in J_q} y_j \delta(x - x_j) - \sum_{j \in J_{q-1}} z_j \delta(x - x_j)\|_{-t}^2}{\|y\|^2}. \end{aligned}$$

These two formulae will be crucial in the subsequent analysis. We start by estimating the smallest and largest eigenvalues of the Schur complement. Let $w = (-\Delta)^{-t} \sum_{j \in J_q} y_j \delta(x - x_j)$, whose Fourier coefficients are

$$(6.4) \quad \hat{w}(m) = \begin{cases} 0, & \text{if } m = 0 \\ (4\pi^2)^{-t} |m|^{-2t} g(m), & \text{else,} \end{cases}$$

where the function $g(m)$ is defined by

$$(6.5) \quad g(m) = \sum_{j \in J_q} y_j \exp(2\pi i(j \cdot 2^{-q}, m)).$$

For the smallest eigenvalue, we write

$$\begin{aligned} \inf_{j \in J_q} \|\sum_{j \in J_q} y_j \delta(x - x_j)\|_{-t}^2 &= \|w\|_{-t}^2 = (4\pi^2)^t \inf_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{2t} |\hat{w}(m)|^2 \\ &= (4\pi^2)^{-t} \inf_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2. \end{aligned}$$

Notice that

$$\inf_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2 = \inf_{m \in B_q^d} M_q^t(m) |g(m)|^2 \geq 2^{2tq} \inf_{m \in B_q^d} |g(m)|^2$$

and

$$\begin{aligned}
 |g(m)|^2 &= \left| \sum_{m \in B_q^d} \sum_{j \in J_q} y_j \exp(2\pi i(j - l)2^{-q}, m) \right|^2 \\
 &= \sum_{m \in B_q^d} \sum_{j \in J_q} \sum_{l \in J_q} y_j y_l \exp(2\pi i((j - l)2^{-q}, m)) \\
 (6.6) \quad &= \sum_{\substack{j \in J_q, l \in J_q \\ j \neq l}} y_j y_l \sum_{m \in B^d} \exp(2\pi i((j - l)2^{-q}, m)) \\
 &\sim 2^{-q} |y|.
 \end{aligned}$$

In the last line we have used the fact that

$$\exp(2\pi i((j - l)2^{-q}, m)) = \begin{cases} 0, & \text{if } j - l \neq 0 \\ \sum_{m \in B_q^d} 1 \sim 2^{qd}, & \text{if } j - l = 0. \end{cases}$$

Thus, combining the above results, we obtain the bound on the smallest eigenvalue

$$\lambda_{\min}(K(t, q)/K(t, q - 1)) \gtrsim 2^{-q(2t-d)}.$$

We then move to consider the largest eigenvalue. First, notice that

$$\inf_{x \in \mathbb{R}^{|J_{q-1}|}} \left\| \sum_{j \in J_q} y_j \delta(x - x_j) - \sum_{j \in J_{q-1}} z_j \delta(x - x_j) \right\|_{\mathcal{F}_t}^2 = \inf_{v \in \mathcal{F}_{t,q-1}} \|w - v\|_{\mathcal{F}_t}^2.$$

Naturally, one can express the optimal v in the above variational formulation using the Fourier series representation explained before. However, this will lead to many interactions between different frequencies. To make the analysis cleaner, we adopt another strategy. We first approximate the function w by a band-limited function, whose projection into $\mathcal{F}_{t,q-1}$ will be more concise. Precisely, define a band limited version of w , written as w_h , by

$$(6.7) \quad \hat{w}_h(m) = \begin{cases} \hat{w}(m), & \text{if } m \in B_{q-1}^d \\ 0, & \text{if } m \in (B_{q-1}^d)^c. \end{cases}$$

To estimate $\inf_{v \in \mathcal{F}_{t,q-1}} \|w - v\|_{\mathcal{F}_t}^2$ we follow the two steps below:

Step 1. We prove $\|w - w_h\|_{\mathcal{F}_t}^2 \lesssim 2^{-q(2t-d)} |y|^2$. Let us calculate the quantity directly:

$$\begin{aligned}
 \|w - w_h\|_{\mathcal{F}_t}^2 &= (4\pi^2)^{-t} \left(\sum_{m \in (B_{q-1}^d)^c} |m|^{-2t} |g(m)|^2 \right) \\
 &= (4\pi^2)^{-t} \left(\sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2 - \sum_{m \in B_{q-1}^d} |m|^{-2t} |g(m)|^2 \right) \\
 &= (4\pi^2)^{-t} \left(\sum_{m \in B_q^d} M_q^t(m) |g(m)|^2 - \sum_{m \in B_{q-1}^d} |m|^{-2t} |g(m)|^2 \right) \\
 &\lesssim 2^{-2qt} \sum_{m \in B_q^d} |g(m)|^2 \lesssim 2^{-q(2t-d)} |y|^2.
 \end{aligned}$$

Here we have used the fact that $M_q^t(m) = |m|^{-2t} \lesssim 2^{-2qt}$ for $m \in B_{q-1}^d$ and $M_q^t(m) \lesssim 2^{-2qt}$ for $m \in B_q^d \setminus B_{q-1}^d$ according to the results in Lemma 2.15. In the last line, the bound (6.6) is applied.

Step 2. We prove $\inf_{v \in F_{t,q-1}} \|w_h - v\|_{L^2(\mathbb{T}^{2t-d})}^2 |y|^2$. Based on Theorem 2.13, we know the optimal v for this variational problem has the Fourier coefficients

$$\hat{v}(m) = \begin{cases} 0, & \text{if } m = 0 \\ |m|^{-2t} \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)}, & \text{else.} \end{cases}$$

Then, using the Fourier representation of the norm, we get

$$\begin{aligned} & \|w_h - v\|_{L^2}^2 \\ & \sim \sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{2t} |\hat{w}_h(m) - \hat{v}(m)|^2 \\ & = \sum_{m \in B_{q-1}^d \setminus \{0\}} |m|^{-2t} |g(m) - \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)}|^2 + \sum_{m \in (B_{q-1}^d)^c} |m|^{-2t} \left| \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)} \right|^2. \end{aligned}$$

For the first term, since w_h is band-limited, we know if $m \in B_{q-1}^d \setminus \{0\}$, then $(T_{q-1} \hat{w}_h)(m) = |m|^{-2t} g(m)$. Thus, we can write this term as

$$\begin{aligned} & \sum_{m \in B_{q-1}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2 \left(1 - \frac{|m|^{-2t}}{M_{q-1}^t(m)} \right)^2 \\ & = \sum_{m \in B_{q-1}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2 \frac{(M_{q-1}^t(m) - |m|^{-2t})^2}{M_{q-1}^t(m)^2} \\ & \stackrel{(a)}{\leq} \sum_{m \in B_{q-1}^d \setminus \{0\}} |m|^{-2t} |g(m)|^2 \frac{2^{-4tq}}{|m|^{-4t}} \\ & \stackrel{(b)}{\leq} \sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{-2tq} |g(m)|^2 \leq 2^{-q(2t-d)} |y|^2, \end{aligned}$$

where in (a), we have used the fact that $M_{q-1}^t(m) - |m|^{-2t} \sim 2^{-2tq}$ and $M_{q-1}^t(m) \sim |m|^{-2t}$ for $m \in B_{q-1}^d \setminus \{0\}$ based on Lemma 2.15. In (b), we have used $|m| \leq 2^q$. The last inequality is obtained by recalling (6.6).

For the second term, we write

$$\begin{aligned} & \sum_{m \in (B_{q-1}^d)^c} |m|^{-2t} \left| \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)} \right|^2 \\ & = \sum_{m \in \mathbb{Z}^d \setminus \{0\}} |m|^{-2t} \left| \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)} \right|^2 - \sum_{m \in B_{q-1}^d \setminus \{0\}} |m|^{-2t} \left| \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)} \right|^2 \\ & \stackrel{(c)}{=} \sum_{m \in B_{q-1}^d \setminus \{0\}} (M_{q-1}^t(m) - |m|^{-2t}) \left| \frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)} \right|^2 \\ & = \sum_{m \in B_{q-1}^d \setminus \{0\}} (M_{q-1}^t(m) - |m|^{-2t}) \frac{|m|^{-2t} |g(m)|^2}{M_{q-1}^t(m)^2} \\ & \leq 2^{-2tq} \sum_{m \in B_{q-1}^d \setminus \{0\}} |g(m)|^2 \leq 2^{-q(2t-d)} |y|^2, \end{aligned}$$

where in (c), we have used the periodicity of the function $\frac{(T_{q-1} \hat{w}_h)(m)}{M_{q-1}^t(m)}$.

Now, combining Steps 1 and 2 leads to the conclusion

$$\inf_{v \in F_{t,q-1}} \|w - v\|_t^2 \leq 2^{-q(2t-d)} \|y\|^2,$$

and in particular, it implies

$$\lambda_{\max}(K(t, q)/K(t, q-1)) \leq 2^{-q(2t-d)}.$$

As a consequence of the upper and lower bounds for the eigenvalues of the matrix $K(t, q)/K(t, q-1)$, we deduce that they are all on the scale of $2^{-q(2t-d)}$. Let C be a constant independent of t, q such that $C^{-1}2^{-q(2t-d)} \leq K(t, q)/K(t, q-1) \leq C2^{-q(2t-d)}$. Then,

$$\begin{aligned} (2^{qd} - 2^{(q-1)d})((2t-d)(-q) \log 2 - C) &\leq \log \det K(t, q)/K(t, q-1) \\ &\leq (2^{qd} - 2^{(q-1)d})((2t-d)(-q) \log 2 + C). \end{aligned}$$

Using the implied bounds on the recursion relation, we get

$$(2t-d)g_1(q) - Cg_2(q) + K(t, 0) \leq \log \det K(t, q) \leq (2t-d)g_1(q) + Cg_2(q) + K(t, 0),$$

where $g_1(q) = \sum_{k=1}^q (2^{kd} - 2^{(k-1)d})(-k \log 2)$ and $g_2(q) = (2^{qd} - 1)(2t-d)$. Summing the series in $g_1(q)$ leads to $g_1(q) \sim -q2^{qd} \log 2 \sim -q2^{qd}$. The proof of Proposition 2.18 is completed.

Remark 6.1. The above technique of using the Schur complements is quite general and could be potentially applied to other operators such as heterogeneous Laplacians; see [24]. However, for the homogeneous Laplacian on the torus in this paper, we may also prove the result via a simpler approach. The key observation is that there is an explicit formula for the spectrum of $K(t, q)$, as also exploited in [33, Sec. 6.7]. Indeed, using the formula for the spectrum given in Lemma 6.2, we get

$$\begin{aligned} \log \det K(t, q) &= \sum_{m \in B_q^d} \log (2^{qd} (4\pi^2)^{-t} M_q^t(m)) \\ &= qd2^{qd} \log 2 - 2^{qd} t \log(4\pi^2) + \sum_{m \in B_q^d} \log M_q^t(m). \end{aligned}$$

By Lemma 2.15, it holds that

$$M_q^t(m) \sim \begin{cases} 2^{-2qt}, & \text{if } m = 0 \\ |m|^{-2t}, & \text{if } m \in B_q^d \setminus \{0\}. \end{cases}$$

That is, there exists a constant C independent of t such that

$$\begin{aligned} -2t \log |m| - \log C &\leq \log M_q^t(m) \leq -2t \log |m| + \log C \\ \text{for } m \in B_q^d \setminus \{0\}, \text{ and } -2^{qt} \log 2 - \log C &\leq \log M_q^t(0) \leq -2^{qt} \log 2 + \log C. \end{aligned}$$

$$\log |m| \sim \int_0^{2^q} r^{d-1} \log r \, dr \sim q2^{qd},$$

and $2^{qd} = o(q2^{qd})$, we get

$$-(2t-d)q2^{qd} - C2^{qd} \leq \log \det K(t, q) \leq -(2t-d)q2^{qd} + C2^{qd}.$$

This completes the alternative proof of Proposition 2.18. $\square$

Lemma 6.2. *The eigenvalues of $K(t, q)$ are $2^{qd}(4\pi^2)^{-t}M_q^t(m)$ for $m \in B_q^d$, where $M_q^t(m)$ is defined in (2.12), with the corresponding eigenfunctions $\varphi_m(x) \in \mathbb{R}^{2^{qd}}$.*

Proof. We can prove this claim using Mercer's decomposition as follows. First, for $x_i, x_j \in X_q$, it holds that

$$\begin{aligned} K(t, q)_{ij} &= \sum_{m \in \mathbb{Z}^d \setminus \{0\}} (4\pi^2)^{-t} |m|^{-2t} \varphi_m(x_i) \varphi_m^*(x_j) \\ &= \sum_{m \in B_q^d} (4\pi^2)^{-t} M_q^t(m) \varphi_m(x_i) \varphi_m^*(x_j), \end{aligned}$$

where we have used the fact that $\varphi_{m+2q\beta}(x_i) = \varphi_m(x_i)$ for any $\beta \in \mathbb{Z}^d$ and $x_i \in X_q$. Thus, for every $n \in B_q^d$, we get

$$\begin{aligned} K(t, q)_{ij} \varphi_n(x_j) &= \sum_{m \in B_q^d} (4\pi^2)^{-t} M_q^t(m) \varphi_m(x_i) \varphi_m^*(x_j) \varphi_n(x_j) \\ &= \sum_{m \in B_q^d} (4\pi^2)^{-t} M_q^t(m) \varphi_m(x_i) 2^{qd} \delta_{mn} \\ &= 2^{qd} (4\pi^2)^{-t} M_q^t(m) \varphi_n(x_i), \end{aligned}$$

where in the second equality we used the property of Fourier series. This implies $\varphi_n(x_q)$ is an eigenfunction. The proof of the lemma is completed. $\square$

6.6. Proof of Theorem 2.19.

Proof. Recall the definition,

$$s^{\text{EB}}(q) = \operatorname{argmin}_t L^{\text{EB}}(t, q) := \|u(\cdot, t, q)\|_t^2 + \log \det K(t, q).$$

Define a rescaled version of the loss function by

$$\tilde{L}_{\text{EB}}(t, q) = \frac{1}{|g_1(q)|} L^{\text{EB}}(t, q) = \underbrace{\frac{1}{|g_1(q)|} \|u(\cdot, t, q)\|_t^2}_{(1)} + \underbrace{\frac{1}{|g_1(q)|} \log \det K(t, q)}_{(2)}.$$

We note that by Proposition 2.18, we have $|g_1(q)| \sim q^{2qd}$. Now, we estimate the growth rate of (1) and (2) separately. From Propositions 2.16 and 2.18, we get

$$\underbrace{(1)}_{(3)} \sim \underbrace{\frac{1}{q} 2^{-q(2s-2t+d)} \zeta_0^2}_{(3)} + \underbrace{\frac{1}{q} 2^{-q(2s-2t)} \sum_{m \in B_q^d \setminus \{0\}} 2^{-q(2t-2s+d)} |m|^{2t-2s} \zeta_m^2}_{(4)},$$

and for the log det part, it holds that

$$d - 2t + \frac{-Cg_2(q) + K(t, 0)}{|g_1(q)|} \leq (2) \leq d - 2t + \frac{Cg_2(q) + K(t, 0)}{|g_1(q)|}.$$

It follows that $\lim_{q \rightarrow \infty} (2) = d - 2t$. Thus, our remaining task is to analyze terms (3), (4) in (1). We split the problem into four cases.

Case 1 ($t = s$). It is easy to see $\lim_{q \rightarrow \infty} \textcircled{3} = 0$ and

$$\textcircled{4} = \frac{1}{q} 2^{-qd} \quad \zeta_m^2 = \frac{1}{q} a(d, q),$$

$$m \in B_q^d \setminus \{0\}$$

so that $\lim_{q \rightarrow \infty} \textcircled{4} = 0$. Here we use the definition of a in Lemma 2.17. Therefore, $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(s, q) = d - 2s$.

Case 2 ($1/\delta \geq t \geq s + e$). We have $\textcircled{3} \geq 0$. The term $\textcircled{4}$ can be written as

$$\textcircled{4} = \frac{1}{q 2^{-q(2t-2s)}} a(2t - 2s + d, q),$$

where we recall the definition of the function a in Lemma 2.17. According to this lemma, we get the uniform convergence

$$\lim_{q \rightarrow \infty} a(2t - 2s + d, q) = \gamma(2t - 2s + d) > 0$$

in probability. In the meantime, $\lim_{q \rightarrow \infty} q 2^{-q(2t-2s)} = 0$. So, $\lim_{q \rightarrow \infty} \textcircled{4} = 0$ in probability, and uniformly in $1/\delta \geq t \geq s + e$. In terms of $\tilde{L}_{EB}(t, q)$, this corresponds to $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(t, q) = \infty$.

Case 3 ($s - e \geq t \geq s - d/2 + e$). In this case, $2t - 2s + d \geq e$ so Lemma 2.17 can be applied. We write the term

$$\textcircled{4} = \frac{2^{-q(2s-2t)}}{q} a(2t - 2s + d, q).$$

This will converge to 0 as q goes to infinity, since $\lim_{q \rightarrow \infty} \frac{2^{-q(2s-2t)}}{q} = 0$ and $\lim_{q \rightarrow \infty} a(2t - 2s + d, q) = \gamma(2t - 2s + d) \in (0, \infty)$. The term $\textcircled{3}$ also converges to 0. Thus, $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(t, q) = d - 2t$ in probability, and uniformly for $s - e \geq t \geq s - d/2 + e$.

Case 4 ($s - d/2 + e \geq t \geq d/2 + \delta$). We still have that $\textcircled{3}$ converges to 0. For term $\textcircled{4}$, we have

$$\textcircled{4} = \frac{2^{-qd}}{q} \quad |m|^{2t-2s} \zeta_m^2 \leq \frac{2^{-qd}}{q} \quad |m|^{2(s-d/2+)-2s} \zeta_m^2,$$

$$m \in B_q^d \setminus \{0\} \quad m \in B_q^d \setminus \{0\}$$

where we have used the monotonicity of the function $|m|^{2t-2s}$ with respect to t . Then, it reduces to the case $t = s - d/2 + \delta$, which is covered by Case 3. Hence, we have $\lim_{q \rightarrow \infty} \textcircled{4} = 0$ uniformly for $s - d/2 + \delta \geq t \geq d/2 + \delta$. Therefore, we get $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(t, q) = d - 2t$ in probability, and uniformly for $s - d/2 + \delta \geq t \geq d/2 + \delta$.

Let us make a summary of the arguments above. We have established that, for any small $e > 0$, $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(t, q) = \infty$ uniformly for $1/\delta \geq t \geq s + e$, and $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(t, q) = d - 2t$ uniformly for $s - e \geq t \geq d/2 + \delta$, and $\lim_{q \rightarrow \infty} \tilde{L}_{EB}(s, q) = d - 2s$. All the convergence is in probability. Note that s^{EB} is the minimizer of $L_{EB}(t, q)$, hence also of $\tilde{L}_{EB}(t, q)$. The above convergence results for $\tilde{L}_{EB}(t, q)$ imply that $s^{EM} \notin [s - e, s + e]$ with probability 1 as q goes to infinity, for any $e > 0$. Thus, we must have

$$\lim_{q \rightarrow \infty} s^{EB}(q) = s.$$

The proof is complete. $\square$

6.7. Proof of Proposition 2.21.

Proof. In order to write the interaction terms as a random series with some desired independence pattern for the random variables involved, we need to consider the geometry of the lattice carefully. We introduce another set $S_q := \{m \in \mathbb{Z} : -2^{q-2} \leq m \leq 3 \times 2^{q-2} - 1\}$ and let $S_q^d = S_q \otimes S_q \otimes \cdots \otimes S_q$ denote the tensor product of d multiples of S_q . The set S_q is a shift of B_{q-1} and S_q^d is a shift of B_{q-1}^d .

Define the set $B_{q-1}^d + 2^{q-1}k := \{m + 2^{q-1}k : m \in B_{q-1}^d\}$ for $k \in \mathbb{Z}^d$. We have the relation

$$S_q^d = \bigcup_{k \in \mathbb{Z}_2^d} (B_{q-1}^d + 2^{q-1}k),$$

where $\mathbb{Z}_2^d = \{0, 1\}^d$. Note that for $k_1 \neq k_2$, the intersection between $B_{q-1}^d + 2^{q-1}k_1$ and $B_{q-1}^d + 2^{q-1}k_2$ is empty.

Using (2.13) and the periodicity of the functions involved, we get

$$\begin{aligned} & \|u(\cdot, t, q) - u(\cdot, t, q-1)\|_{\ell^2}^2 \\ &= (4\pi^2)^t \sum_{m \in B_q^d} M_q^t(m) \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2 \\ &= (4\pi^2)^t \sum_{m \in S_q^d} M_q^t(m) \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2 \\ &= (4\pi^2)^t \sum_{k \in \mathbb{Z}_2^d} \sum_{m \in (B_{q-1}^d + 2^{q-1}k)} M_q^t(m) \left(\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \right)^2. \end{aligned}$$

Recall the relation

$$T_{q-1} \hat{u}(m) = \sum_{l \in \mathbb{Z}_2^d} T_q \hat{u}(m + 2^{q-1}l),$$

based on which we get

$$\begin{aligned} & \frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)} \\ &= \left(\frac{1}{M_q^t(m)} - \frac{1}{M_{q-1}^t(m)} \right) T_q \hat{u}(m) - \frac{1}{M_{q-1}^t(m)} \sum_{l \in \mathbb{Z}_2^d \setminus \{0\}} T_q \hat{u}(m + 2^{q-1}l). \end{aligned}$$

Since $u^\dagger \sim \mathcal{N}(0, (-\Delta)^{-s})$, it holds $\hat{u}(m) \sim \mathcal{N}(0, (4\pi^2)^{-s} |m|^{-2s})$. Moreover, for different m , these Gaussian random variables are independent from each other. Thus, for a fixed k and for $m \in (B_{q-1}^d + 2^{q-1}k)$, the Gaussian random variables

$$\frac{T_q \hat{u}(m)}{M_q^t(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^t(m)}$$

are independent from each other. Furthermore, by calculating their variance, we can write

$$\begin{aligned}
 & \frac{M_q^s(m)}{q} \left[\frac{T_q \hat{u}(m)}{M_q^s(m)} - \frac{T_{q-1} \hat{u}(m)}{M_{q-1}^s(m)} \right]^2 \\
 &= (4\pi)^{-s} \mathbb{E} \left[\left(\frac{1}{M_q^s(m)} - \frac{1}{M_{q-1}^s(m)} \right) M_q(m) M_q^s(m) + \frac{M_q^s(m)}{(M_{q-1}^s(m))^2} M_q^s(m+2^{q-1}) \right] \zeta_{k,m}^2 \\
 &= (4\pi^2)^{-s} \mathbb{E} \left[\frac{M_q^s(m)(M_q^s(m) - M_{q-1}^s(m))^2}{M_q^s(m)(M_{q-1}^s(m))^2} + \frac{M_q^s(m)}{(M_{q-1}^s(m))^2} M_q^s(m+2^{q-1}) \right] \zeta_{k,m}^2 \\
 &=: A_{k,m} \zeta_{k,m}^2,
 \end{aligned}$$

where $\{\zeta_{k,m}\}_k$ are independent unit scalar Gaussian random variables. Clearly, we have the lower bound

$$A_{k,m} \geq (4\pi^2)^{-s} \frac{M_q^s(m)}{(M_{q-1}^s(m))^2} M_q^s(m - 2^{q-1}k).$$

Thus, denoting $e_1 = (1, 0, \dots, 0) \in \mathbb{Z}^d$, we get

$$\begin{aligned}
 & \|u(\cdot, t, q) - u(\cdot, t, q-1)\|_t^2 \\
 & \geq (4\pi^2)^{-2s} \sum_{k \in \mathbb{Z}^d} \sum_{m \in (B_{q-1}^d + 2^{q-1}k)} \frac{M_q^s(m)}{(M_{q-1}^s(m))^2} M_q^s(m - 2^{q-1}k) \zeta_{k,m}^2 \\
 & \geq (4\pi^2)^{-2s} \sum_{m \in (B_{q-1}^d + 2^{q-1}e_1)} \frac{M_q^s(m)}{(M_{q-1}^s(m))^2} M_q^s(m - 2^{q-1}e_1) \zeta_{e_1,m}^2 \\
 & = (4\pi^2)^{-2s} \sum_{m \in (B_{q-1}^d + 2^{q-1}e_1)} \frac{M_q^s(m)}{(M_{q-1}^s(m - 2^{q-1}e_1))^2} M_q^s(m - 2^{q-1}e_1) \zeta_{e_1,m}^2 \\
 & \stackrel{\diamond}{\geq} \sum_{m \in (B_{q-1}^d \setminus \{0\} + 2^{q-1}e_1)} \frac{2^{-2qt}}{|m - 2^{q-1}e_1|^{-4t}} |m - 2^{q-1}e_1|^{2s} \zeta_{e_1,m}^2 \\
 & = \sum_{m \in B_{q-1}^d \setminus \{0\}} \frac{2^{-2qt}}{2^{2qt} |m|^{4t-2s}} \zeta_{e_1, m+2^{q-1}e_1}^2.
 \end{aligned}$$

In the above derivation, we have used the fact that for $m \notin B_{q-1}^d$, it holds that $M_q^s(m) \sim |m|^{-2s}$, $M_{q-1}^s(m) \sim |m|^{-2t}$, and in particular, $M_q^s(m) \sim |m|^{-2t} \sim 2^{-2qt}$ for $m \in (B_{q-1}^d \setminus \{0\} + 2^{q-1}e_1)$. Renaming the subscripts in $\zeta_{e_1, m+2^{q-1}e_1}^2$ completes the proof. $\square$

6.8. Proof of Proposition 2.22.

Proof. We need to upper bound $A_{k,m}$ for $k \in \mathbb{Z}^d$, $m \in B_{q-1}^d + 2^{q-1}k$, which is defined in the proof of Proposition 2.21. First, we have

$$M_q^s(m + 2^{q-1}l) = M_{q-1}^s(m) - M_q^s(m),$$

$$l \in \mathbb{Z}_2^d \setminus \{0\}$$

and the estimate $0 \leq M_{q-1}^s(m) - M_q^s(m) \leq M_{q-1}^s(m)$ for any $d/2 + \delta \leq t \leq 1/\delta$. Based on this observation, for $k \in \mathbb{Z}^d \setminus \{0\}$ and $m \in B_{q-1}^d \setminus \{0\} + 2^{q-1}k$, we have the

bound

$$A_{k,m} \lesssim \frac{M_q^s(m)}{M_q^t(m)} + M_q^t(m) \frac{M_{q^{-1}}^s(m)}{(M_{q^{-1}}^t(m))^2} \\ \lesssim 2^{-q(2s-2t)} + 2^{-2tq} |m - 2^{q-1}k|^{4t-2s},$$

where we have used the fact that for $m \in B^d \setminus \{0\} + 2^{q-1}k$, it holds that $M^s(m) \sim 2^{-2sq}$, $M_q^t(m) \sim 2^{-2tq}$, $M_{q^{-1}}^s(m) \sim |m - 2^{q-1}k|^{-2s}$, $M_{q^{-1}}^t(m) \sim |m - 2^{q-1}k|^{-2t}$, according to Lemma 2.15. For $m = 2^{q-1}k$, we get $A_{k,m} \lesssim 2^{-q(2s-2t)}$. So in general, we can write $A_{k,m} \lesssim 2^{-q(2s-2t)} + 2^{-2tq} |m - 2^{q-1}k|^{4t-2s}$ for $m \in B_{q^{-1}}^d + 2^{q-1}k$ where

we use the convention that $|m|^a = 0$ for $m = 0$ and any $a \in \mathbb{R}$ to make the notation more compact.

When $k = 0$, using Lemma 2.15 again, we get for $m \in B_{q^{-1}}^d \setminus \{0\}$,

$$A_{k,m} \lesssim \frac{M_q^s(m)(M_q^t(m) - M_q^t(m))^2}{M_q^t(m)(M_{q^{-1}}^t(m))^2} + \frac{M_q^t(m)}{(M_{q^{-1}}^t(m))^2} (M_{q^{-1}}^s(m) - M_q^s(m)) \\ \lesssim \frac{|m|^{-2s} 2^{-4tq}}{|m|^{-2t}} + \frac{|m|^{-2t}}{|m|^{-4t}} 2^{-2sq} \\ = |m|^{\frac{-6t}{6t-2s} 2^{-4tq} + |m|^{2t} 2^{-2sq}} \\ \lesssim 2^{-2tq} |m|^{4t-2s} + 2^{-q(2s-2t)},$$

where in the last line we used the relation $|m| \lesssim 2^q$. For $m = 0$, based on the above calculation, we can get $A_{k,m} \lesssim 2^{-q(2s-2t)}$. Thus, generally, we can write $A_{k,m} \lesssim 2^{-2tq} |m|^{4t-2s} + 2^{-q(2s-2t)}$ for $m \in B_{q^{-1}}^d$ by using the notational convention above.

Combining these estimates, we arrive at

$$\|u(\cdot, t, q) - u(\cdot, t, q-1)\|_{L^2}^2 \\ \lesssim \sum_{k \in \mathbb{Z}_2^d, m \in (B_{q^{-1}}^d + 2^{q-1}k)} A_{k,m} \zeta_{k,m}^2 \\ \lesssim \sum_{k \in \mathbb{Z}_2^d, m \in (B_{q^{-1}}^d + 2^{q-1}k)} (2^{-q(2s-2t)} + 2^{-2tq} |m - 2^{q-1}k|^{4t-2s}) \zeta_{k,m}^2 \\ = \sum_{k \in \mathbb{Z}_2^d, m \in B_{q^{-1}}^d} (2^{-q(2s-2t)} + 2^{-2tq} |m|^{4t-2s}) \zeta_{k,m+2^{q-1}k}^2.$$

After a change of notation, we get the desired estimate. $\square$

6.9. Proof of Theorem 2.23.

Proof. Recall

$$s^{\text{KF}}(q) = \operatorname{argmin}_{t \in [d/2+\delta, 1/\delta]} L^{\text{KF}} : \frac{\|u(\cdot, t, q) - u(\cdot, t, q-1)\|_{L^2}^2}{\|u(\cdot, t, q)\|_{L^2}^2}.$$

We analyze the denominator and numerator separately. We start with the numerator. Let

$$V_1(t, q) = \frac{1}{q} 2^{q(s-d/2)} \|u(\cdot, t, q) - u(\cdot, t, q-1)\|_{L^2}^2.$$

Case 1 ($t = \frac{s-d/2}{2}$). We derive an upper bound on V_1 . By Proposition 2.22,

$$|u(\cdot, t, q) - u(\cdot, t, q-1)|_t^2 \leq \sum_{k \in Z_{\frac{d}{2}}^d} \sum_{m \in B_{q-1}^d} (2^{-q(2s-2t)} + 2^{-2tq} |m|^{4t-2s}) \zeta_{k,m}^2$$

Take $t = \frac{s-d/2}{2}$. For each $k \in Z_{\frac{d}{2}}^d$ consider the term

$$\begin{aligned} V_1^k(t, q) &= \frac{1}{q} \sum_{m \in B_{q-1}^d} 2^{q(s-d/2)} (2^{-q(2s-2t)} + 2^{-2tq} |m|^{4t-2s}) \zeta_{k,m}^2 \\ &= \frac{1}{q} \sum_{m \in B_{q-1}^d} 2^{q(s-d/2)} (2^{-q(s+d/2)} + 2^{-q(s-d/2)} |m|^{-d}) \zeta_{k,m}^2 \\ &= \frac{1}{q} \sum_{m \in B_{q-1}^d} (2^{-qd} + |m|^{-d}) \zeta_{k,m}^2 \\ &\leq \frac{1}{q} \sum_{m \in B_{q-1}^d} |m|^{-d} \zeta_{k,m}^2. \end{aligned}$$

By Lemma 2.17, $\lim_{q \rightarrow \infty} \frac{1}{q} \sum_{m \in B_{q-1}^d} |m|^{-d} \zeta_{k,m}^2 = \gamma(o) \in (0, \infty)$. Thus, $V_1^k(t, q)$ remains bounded for $q \in \mathbb{N}$. Since $V_1(t, q) = \sum_{k \in Z_{\frac{d}{2}}^d} V_1^k(t, q)$, it follows that $V_1(t, q)$ remains bounded for $q \in \mathbb{N}$, in the case $t = \frac{s-d/2}{2}$.

Case 2 ($1/\delta \geq t \geq \frac{s-d/2}{2} + e$). We provide a lower bound of V_1 here. Using Proposition 2.21, we get

$$\begin{aligned} V_1(t, q) &\gtrsim \frac{1}{q} \sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{-2tq} |m|^{4t-2s} \zeta_m^2 \\ &= \frac{1}{q} \sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{q(s-d/2-2t)} |m|^{4t-2s} \zeta_m^2 \\ &= \frac{1}{q} 2^{q(s-d/2-2t)} 2^{(q-1)(4t-2s+d)} \left(\sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{-(q-1)(4t-2s+d)} |m|^{4t-2s} \zeta_m^2 \right) \\ &= \frac{1}{q} 2^{(q/2-1)(4t-2s+d)} a(4t-2s+d, q-1). \end{aligned}$$

By Lemma 2.17, $\lim_{q \rightarrow \infty} a(4t-2s+d, q-1) = \gamma(4t-2s+d) > 0$ uniformly for $1/\delta \geq t \geq \frac{s-d/2}{2} + e$. Since $\lim_{q \rightarrow \infty} \frac{1}{q} 2^{(q/2-1)(4t-2s+d)} = \infty$, we get $\lim_{q \rightarrow \infty} V_1(t, q) = \infty$ and its growth rate is $\gtrsim \frac{1}{q} 2^{(q/2-1)(4t-2s+d)}$.

Case 3 ($\frac{s-d/2}{2} - e \geq t \geq d/2 + \delta$). We provide a lower bound on V_1 here. Similarly to our analysis in Case 2, we have

$$\begin{aligned} V_1(t, q) &\gtrsim \frac{1}{q} \sum_{m \in B_{q-1}^d \setminus \{0\}} 2^{q(s-d/2-2t)} |m|^{4t-2s} \zeta_m^2 \\ &\gtrsim \frac{1}{q} 2^{q(s-d/2-2t)} \zeta_1^2. \end{aligned}$$

Then, it holds that

$$\mathbb{P}\left(\frac{1}{q^{2q(s-d/2-2\eta)/2}} \xi_1^2 \geq 2 \right) = \mathbb{P}(\xi_1^2 \geq q2^{-q(s-d/2-2\eta)/2}) \rightarrow 1$$

as $q \rightarrow \infty$. Thus, we get $\lim_{q \rightarrow \infty} V_1(t, q) = \infty$ uniformly for this range of t and the growth rate is $\diamond 2^{q(s-d/2-2\eta)/2}$. We have finished the analysis of the numerator. Now we proceed to analyze the denominator, which comprises the norm term. From Proposition 2.16, we have

$$(6.8) \quad \|u(\cdot, t, q)\|_t^2 \sim 2^{-q(2s-2\eta)\xi_0^2} + \sum_{m \in B_{\frac{s}{2}} \setminus \{0\}} |m|^{2t-2s\xi_m^2},$$

where $\{\xi_m\}_{m \in B_{\frac{s}{2}}}$ are independent unit scalar Gaussian random variables. Recall that our final target in this theorem is to show that, for any $\epsilon > 0$,

$$\lim_{q \rightarrow \infty} \mathbb{P}[S^{KF}(q) \in (\frac{s-d/2}{2} - \epsilon, \frac{s-d/2}{2} + \epsilon)] = 1$$

Let $I = [d/2 + \delta, 1/\delta] \setminus [\frac{s-d/2}{2} - \epsilon, \frac{s-d/2}{2} + \epsilon]$. By rewriting the loss function, it suffices to show

$$\lim_{q \rightarrow \infty} \mathbb{P}\left[\frac{V_1(\frac{s-d/2}{2}, q)}{\|u(\cdot, \frac{s-d/2}{2}, q)\|_{\frac{s-d/2}{2}}^2} \geq \inf_{t \in I} \frac{V_1(t, q)}{\|u(\cdot, t, q)\|_t^2}\right] = 0.$$

Let us write

$$(6.9) \quad r(t, q) = \frac{V_1(t, q)}{V_1(\frac{s-d/2}{2}, q)} \cdot \frac{\|u(\cdot, \frac{s-d/2}{2}, q)\|_{\frac{s-d/2}{2}}^2}{\|u(\cdot, t, q)\|_t^2},$$

then all we need is to show

$$\lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in I} r(t, q) \leq 1] = 0.$$

For $t \in I^1 = [d/2 + \delta, \frac{s-d/2}{2} - \epsilon]$, according to the analysis for the numerator, we have that for some constant C independent of q ,

$$(6.10) \quad \lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in I^1} \frac{V_1(t, q)}{2^{q(s-d/2-2\eta)/2}} \geq C] = 1,$$

and also, $V_1(\frac{s-d/2}{2}, q)$ remains uniformly bounded for $q \in \mathbb{N}$. Furthermore, equation (6.8) implies the following relation:

$$(6.11) \quad \inf_{t \in I^1} \frac{\|u(\cdot, \frac{s-d/2}{2}, q)\|_{\frac{s-d/2}{2}}^2}{\|u(\cdot, t, q)\|_t^2} \diamond 1,$$

due to the inequality $t \leq \frac{s-d/2}{2} - \epsilon$. Combining the above two estimates in (6.10) and (6.11), and recalling the expression for $r(t, q)$ in (6.9), we get

$$(6.12) \quad \lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in I^1} r(t, q) \leq 1] = 0.$$

Then, let $I^2 = [\frac{s-d/2}{2} + \epsilon, 1/\delta]$. We also need to show $\lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in I^2} r(t, q) \leq 1] = 0$, or equivalently,

$$\lim_{q \rightarrow \infty} \mathbb{P}\left[\frac{\|u(\cdot, \frac{s-d/2}{2}, q)\|_{\frac{s-d/2}{2}}^2}{V_1(\frac{s-d/2}{2}, q)} \leq \sup_{t \in I^2} \frac{\|u(\cdot, t, q)\|_t^2}{V_1(t, q)}\right] = 0.$$

Since $V_1(\frac{s-d/2}{2}, q)$ remains bounded according to the result in Case 1, it suffices to show

$$\lim_{q \rightarrow \infty} \sup_{t \in \mathbb{R}_s^d} \frac{|u(\cdot, t, q)|^2}{V_1(t, q)} = 0$$

in probability. Using the estimate of $V_1(t, q)$ in Case 2 that

$$V_1(t, q) \asymp \frac{1}{q} 2^{-(q/2-1)(4t-2s+d)},$$

it suffices to show

$$\lim_{q \rightarrow \infty} \sup_{t \in \mathbb{R}_s^d} q 2^{-(q/2-1)(4t-2s+d)} |u(\cdot, t, q)|^2 = 0.$$

To achieve this, we recall the expression of the norm term and write

$$q 2^{-(q/2-1)(4t-2s+d)} |u(\cdot, t, q)|^2 = \frac{q(s+d)+4t-2s+d}{2} \sum_{m \in B_q^d \setminus \{0\}} |m|^{2t-2s} \zeta_m^2$$

Clearly, the first term on the right hand side converges to 0, so we only need to deal with the second term. Let

$$\beta(t, q) = q 2^{-(q/2-1)(4t-2s+d)} \sum_{m \in B_q^d \setminus \{0\}} |m|^{2t-2s} \zeta_m^2.$$

Consider $t \in [s - d/2 + e', 1/\delta]$ where e' is a parameter to be tuned. We have $2t - 2s + d \geq e' > 0$ so we are able to write

$$\begin{aligned} \beta(t, q) &= q 2^{-(q/2-1)(4t-2s+d)} 2^{q(2t-2s+d)} a(2t - 2s + d, q) \\ &= q 2^{-q(s-d/2)+4t-2s+d} a(2t - 2s + d, q). \end{aligned}$$

By Lemma 2.17, $\lim_{q \rightarrow \infty} a(2t - 2s + d, q) = \gamma(2t - 2s + d)$ in probability uniformly for $t \in [s - d/2 + e', 1/\delta]$. Since $\lim_{q \rightarrow \infty} q 2^{-(q/2-1)(4t-2s+d)} 2^{q(2t-2s+d)} = 0$, we get $\lim_{q \rightarrow \infty} \sup_{t \in [s-d/2+e', 1/\delta]} \beta(t, q) = 0$.

For $t \in [\frac{s-d/2}{2} + e, s - d/2 + e']$, we have the estimate

$$q 2^{-(q/2-1)(4t-2s+d)} \leq q 2^{-(q/2-1)(4t-2s+d)} \leq \frac{1}{t - \frac{s-d/2}{2} + e} = q 2^{-2q+4}$$

and

$$\sum_{m \in B_q^d \setminus \{0\}} |m|^{2t-2s} \zeta_m^2 \leq \sum_{m \in B_q^d \setminus \{0\}} |m|^{-d+2} \zeta_m^2$$

where we have used the fact that t is upper bounded by $s - d/2 + e'$. Hence,

$$\begin{aligned} \sup_{t \in [\frac{s-d/2}{2} + e, s-d/2+e']} \beta(t, q) &\leq q 2^{-2q+4} \sum_{m \in B_q^d \setminus \{0\}} |m|^{-d+2} \zeta_m^2 \\ &= q 2^{-2q+4} 2^{2q'} a(2e', q). \end{aligned}$$

Now, we set $e' = e/2$ such that $\lim_{q \rightarrow \infty} q 2^{-2q+4} 2^{2q'} = 0$. Lemma 2.17 leads to $\lim_{q \rightarrow \infty} a(2e', q) = \gamma(2e') < \infty$, from which we can conclude $\lim_{q \rightarrow \infty} \sup_{t \in \mathbb{R}_s^d} \beta(t, q) = 0$. Therefore, we get

$$(6.13) \quad \lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in \mathbb{R}_s^d} r(t, q) \leq 1] = 0.$$

Combining (6.12) and (6.13) gives

$$(6.14) \quad \lim_{q \rightarrow \infty} \mathbb{P}[\inf_{t \in I_q} r(t, q) \leq 1] = 0.$$

Based on the definition of $r(t, q)$ in (6.9) and the arguments therein, we obtain

$$\lim_{q \rightarrow \infty} \mathbb{P}[s^{KF}(q) \in (\frac{s-d/2}{2} - e, \frac{s-d/2}{2} + e)] = 1,$$

from which the consistency of the KF estimator follows. $\square$

References

- [1] D. M. Allen, *The relationship between variable selection and data augmentation and a method for prediction*, *Technometrics* **16** (1974), 125–127, DOI 10.2307/1267500. MR343481
- [2] S.-i. Amari and S. Wu, *Improving support vector machine classifiers by modifying kernel functions*, *Neural Net.* **12** (1999), no. 6, 783–789.
- [3] F. Bachoc, *Cross validation and maximum likelihood estimations of hyper-parameters of Gaussian processes with model misspecification*, *Comput. Statist. Data Anal.* **66** (2013), 55–69, DOI 10.1016/j.csda.2013.03.016. MR3064023
- [4] F. Bachoc, A. Lagnoux, and T. M. N. Nguyen, *Cross-validation estimation of covariance parameters under fixed-domain asymptotics*, *J. Multivariate Anal.* **160** (2017), 42–67, DOI 10.1016/j.jmva.2017.06.003. MR3688689
- [5] V. I. Bogachev, *Gaussian measures*, *Mathematical Surveys and Monographs*, vol. 62, American Mathematical Society, Providence, RI, 1998, DOI 10.1090/surv/062. MR1642391
- [6] C. Cortes, M. Kloft, and M. Mohri, *Learning kernels using local rademacher complexity*, *Advances in Neural Information Processing Systems (NeurIPS Proceedings)* **26**, pp. 2760–2768, 2013.
- [7] M. Dashti and A. M. Stuart, *The Bayesian approach to inverse problems*, *Handbook of uncertainty quantification*. Vol. 1, 2, 3, Springer, Cham, 2017, pp. 311–428. MR3839555
- [8] C. de Boor, R. A. DeVore, and A. Ron, *Approximation from shift-invariant subspaces of $L_2(\mathbb{R}^d)$* , *Trans. Amer. Math. Soc.* **341** (1994), no. 2, 787–806, DOI 10.2307/2154583. MR1195508
- [9] M. M. Dunlop, T. Helin, and A. M. Stuart, *Hyperparameter estimation in Bayesian MAP estimation: parameterizations and consistency*, *SMAI J. Comput. Math.* **6** (2020), 69–100, DOI 10.5802/smai-jcm.62. MR4100532
- [10] S. Geisser, *The predictive sample reuse method with applications*, *Journal of the American statistical Association*, **70** (1975), no. 350, 320–328.
- [11] J. K. Ghosh and R. V. Ramamoorthi, *Bayesian nonparametrics*, *Springer Series in Statistics*, Springer-Verlag, New York, 2003. MR1992245
- [12] P. Guttorp and T. Gneiting, *Studies in the history of probability and statistics. XLIX. On the Mat'ern correlation family*, *Biometrika* **93** (2006), no. 4, 989–995, DOI 10.1093/biomet/93.4.989. MR2285084
- [13] B. Hamzi and H. Owhadi, *Learning dynamical systems from data: a simple cross-validation perspective, part I: parametric kernel flows*, *Phys. D* **421** (2021), 132817, DOI 10.1016/j.physd.2020.132817. MR4233447
- [14] K. He, X. Zhang, S. Ren, and J. Sun, *Deep residual learning for image recognition*, *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 770–778.
- [15] N. L. Hjort, C. Holmes, P. Müller, and S. G. Walker, *Bayesian Nonparametrics*, volume 28. Cambridge University Press, 2010.
- [16] T. Hofmann, B. Schölkopf, and A. J. Smola, *Kernel methods in machine learning*, *Ann. Statist.* **36** (2008), no. 3, 1171–1220, DOI 10.1214/009053607000000677. MR2418654
- [17] A. Jacot, F. Gabriel, and C. Hongler, *Neural tangent kernel: convergence and generalization in neural networks*, *Advances in Neural Information Processing Systems*, 2018, pp. 8571–8580.
- [18] B. T. Knapik, B. T. Szabó, A. W. van der Vaart, and J. H. van Zanten, *Bayes procedures for adaptive inference in inverse problems for the white noise model*, *Probab. Theory Related Fields* **164** (2016), no. 3–4, 771–813, DOI 10.1007/s00440-015-0619-7. MR3477780

- [19] B. T. Knapik, A. W. van der Vaart, and J. H. van Zanten, *Bayesian inverse problems with Gaussian priors*, *Ann. Statist.* **39** (2011), no. 5, 2626–2657, DOI 10.1214/11-AOS920. MR2906881
- [20] R. Kohavi, et al, *A study of cross-validation and bootstrap for accuracy estimation and model selection*, *IJCAI*, Montreal, Canada, vol. 14, 1995, pp. 1137–1145.
- [21] R. Kohn, C. F. Ansley, and D. Tharm, *The performance of cross-validation and maximum likelihood estimators of spline smoothing parameters*, *J. Amer. Statist. Assoc.* **86** (1991), no. 416, 1042–1050. MR1146351
- [22] F. Lindgren, H. Rue, and J. Lindström, *An explicit link between Gaussian fields and Gaussian Markov random fields: the stochastic partial differential equation approach*, *J. R. Stat. Soc. Ser. B Stat. Methodol.* **73** (2011), no. 4, 423–498, DOI 10.1111/j.1467-9868.2011.00777.x. With discussion and a reply by the authors. MR2853727
- [23] H. Owhadi, *Do ideas have shape? Plato's theory of forms as the continuous limit of artificial neural networks*, Preprint, arXiv:2008.03920, 2020.
- [24] H. Owhadi and C. Scovel, *Operator-adapted wavelets, fast solvers, and numerical homogenization*, *Cambridge Monographs on Applied and Computational Mathematics*, vol. 35, Cambridge University Press, Cambridge, 2019. From a game theoretic approach to numerical approximation and algorithm design. MR3971243
- [25] H. Owhadi and G. R. Yoo, *Kernel flows: from learning kernels from data into the abyss*, *J. Comput. Phys.* **389** (2019), 22–47, DOI 10.1016/j.jcp.2019.03.040. MR3934896
- [26] O. Perrin and P. Monestiez, *Modeling of non-stationary spatial structure using parametric radial basis deformations*, *GeoENV II—Geostatistics for Environmental Applications*, Springer, 1999, pp. 175–186.
- [27] C. E. Rasmussen, *Gaussian processes in machine learning*, *Summer School on Machine Learning*, Springer, 2003, pp. 63–71.
- [28] A. Ron, *The L_2 -approximation orders of principal shift-invariant spaces generated by a radial basis function*, *Numerical methods in approximation theory*, Vol. 9 (Oberwolfach, 1991), *Internat. Ser. Numer. Math.*, vol. 105, Birkhäuser, Basel, 1992, pp. 245–268, DOI 10.1007/978-3-0348-8619-2_14. MR1269365
- [29] P. D. Sampson and P. Guttorp, *Nonparametric estimation of nonstationary spatial covariance structure*, *J. Amer. Statist. Assoc.*, **87** (1992), no. 417, 108–119.
- [30] M. Scheuerer, R. Schaback, and M. Schlather, *Interpolation of spatial data—a stochastic or a deterministic problem?*, *European J. Appl. Math.* **24** (2013), no. 4, 601–629, DOI 10.1017/S0956792513000016. MR3082868
- [31] A. M. Schmidt and A. O'Hagan, *Bayesian inference for non-stationary spatial covariance structure via spatial deformations*, *J. R. Stat. Soc. Ser. B Stat. Methodol.* **65** (2003), no. 3, 743–758, DOI 10.1111/1467-9868.00413. MR1998632
- [32] M. L. Stein, *A comparison of generalized cross validation and modified maximum likelihood for estimating the parameters of a stochastic process*, *Ann. Statist.* **18** (1990), no. 3, 1139–1157, DOI 10.1214/aos/1176347743. MR1062702
- [33] M. L. Stein, *Interpolation of spatial data*, *Springer Series in Statistics*, Springer-Verlag, New York, 1999. Some theory for Kriging, DOI 10.1007/978-1-4612-1494-6. MR1697409
- [34] C. J. Stone, *An asymptotically optimal window selection rule for kernel density estimates*, *Ann. Statist.* **12** (1984), no. 4, 1285–1297, DOI 10.1214/aos/1176346792. MR760688
- [35] A. M. Stuart and A. L. Teckentrup, *Posterior consistency for Gaussian process approximations of Bayesian posterior distributions*, *Math. Comp.* **87** (2018), no. 310, 721–753, DOI 10.1090/mcom/3244. MR3739215
- [36] A. L. Teckentrup, *Convergence of Gaussian process regression with estimated hyperparameters and applications in Bayesian inverse problems*, *SIAM/ASA J. Uncertain. Quantif.* **8** (2020), no. 4, 1310–1337, DOI 10.1137/19M1284816. MR4164077
- [37] M. J. van der Laan, S. Dudoit, and A. W. van der Vaart, *The cross-validated adaptive epsilon-net estimator*, *Statist. Decisions* **24** (2006), no. 3, 373–395, DOI 10.1524/std.2006.24.3.373. MR2305113
- [38] A. W. van der Vaart and J. A. Wellner, *Weak convergence and empirical processes*, *Springer Series in Statistics*, Springer-Verlag, New York, 1996. With applications to statistics, DOI 10.1007/978-1-4757-2545-2. MR1385671

- [39] A. W. van der Vaart, S. Dudoit, and M. J. van der Laan, *Oracle inequalities for multi-fold cross validation*, *Statist. Decisions* **24** (2006), no. 3, 351–371, DOI 10.1524/std.2006.24.3.351. MR2305112
- [40] G. Wahba and J. Wendelberger, *Some new mathematical methods for variational objective analysis using splines and cross validation*, *Monthly Weather Rev.*, **108** (1980), no. 8, 1122–1143.
- [41] J. J. Warnes and B. D. Ripley, *Problems with likelihood estimation of covariance functions of spatial Gaussian processes*, *Biometrika* **74** (1987), no. 3, 640–642, DOI 10.1093/biomet/74.3.640. MR909370
- [42] H. Wendland, *Scattered data approximation*, *Cambridge Monographs on Applied and Computational Mathematics*, vol. 17, Cambridge University Press, Cambridge, 2005. MR2131724
- [43] P. Whittle, *On stationary processes in the plane*, *Biometrika* **41** (1954), 434–449, DOI 10.1093/biomet/41.3-4.434. MR67450
- [44] A. G. Wilson, Z. Hu, R. Salakhutdinov, and E. P. Xing, *Deep kernel learning*, *Proceedings of the 19th International Conference on Artificial Intelligence and Statistics* **51**, 2016, pp. 370–378.
- [45] Y. Yang, *Consistency of cross validation for comparing regression procedures*, *Ann. Statist.* **35** (2007), no. 6, 2450–2473, DOI 10.1214/009053607000000514. MR2382654
- [46] Z. Ying, *Asymptotic properties of a maximum likelihood estimator with data from a Gaussian process*, *J. Multivariate Anal.* **36** (1991), no. 2, 280–296, DOI 10.1016/0047-259X(91)90062-7. MR1096671
- [47] G. R. Yoo and H. Owhadi, *Deep regularization and direct training of the inner layers of neural networks with kernel flows*, *Preprint*, arXiv:2002.08335, 2020.
- [48] H. Zhang and Y. Wang, *Kriging and cross-validation for massive spatial data*, *Environmetrics* **21** (2010), no. 3–4, 290–304, DOI 10.1002/env.1023. MR2842244

Applied and Computational Mathematics, Caltech, Pasadena, California 91106
 Email address: yifan@caltech.edu

Applied and Computational Mathematics, Caltech, Pasadena, California 91106
 Email address: owhadi@caltech.edu

Applied and Computational Mathematics, Caltech, Pasadena, California 91106
 Email address: astuart@caltech.edu