

Distributed Bayesian Inference in Massive Spatial Data

Rajarshi Guhaniyogi*

Department of Statistics, Texas A & M University, College Station

Cheng Li†

Department of Statistics and Data Science, National University of Singapore

Terrance Savitsky‡

U.S. Bureau of Labor Statistics

Sanvesh Srivastava§

Department of Statistics and Actuarial Science, The University of Iowa

Abstract. Gaussian process (GP) regression is computationally expensive in spatial applications involving massive data. Various methods address this limitation, including a small number of Bayesian methods based on distributed computations (or the divide-and-conquer strategy). Focusing on the latter literature, we achieve three main goals. First, we develop an extensible Bayesian framework for distributed spatial GP regression that embeds many popular methods. The proposed framework has three steps that partition the entire data into many subsets, apply a readily available Bayesian spatial process model in parallel on all the subsets, and combine the posterior distributions estimated on all the subsets into a pseudo posterior distribution that conditions on the entire data. The combined pseudo posterior distribution replaces the full data posterior distribution in prediction and inference problems. Demonstrating our framework's generality, we extend posterior computations for (non-distributed) spatial process models with a stationary full-rank and a nonstationary low-rank GP priors to the distributed setting. Second, we contrast the empirical performance of popular distributed approaches with some widely used non-distributed alternatives and highlight their relative advantages and shortcomings. Third, we provide theoretical support for our numerical observations and show that the Bayes L_2 -risks of the combined posterior distributions obtained from a subclass of the divide-and-conquer methods achieves the near-optimal convergence rate in estimating the true spatial surface with various types of covariance functions. Additionally, we provide upper bounds on the number of subsets to achieve these near-optimal rates.

Key words and phrases: Distributed Bayesian inference, Gaussian process, low-rank Gaussian process, massive spatial data, Wasserstein barycenter.

*rajguhaniyogi@tamu.edu

†stalic@nus.edu.sg

‡savitsky.terrance@bls.gov

§sanvesh-srivastava@uiowa.edu Corresponding author.

1. INTRODUCTION

A fundamental challenge in geostatistics is the analysis of massive spatially-referenced data. Such data sets provide scientists with an unprecedented opportunity to hypothesize and test complex theories, see for example [Cressie and Wikle \(2011\)](#), [Banerjee et al. \(2014\)](#). This has led to the development of complex and flexible GP-based models that are computationally intractable for a large number of spatial locations, denoted as n , due to the $O(n^3)$ computational cost and the $O(n^2)$ storage cost. An overwhelming number of methods exists to address this issue that develop either efficient alternatives to the GP model or efficient approximations of the likelihood. We broadly refer to these approaches as the *non-distributed* methods. An emerging class of Bayesian methods addresses this problem using distributed computations, where the scalability of an existing, possibly non-distributed, spatial GP regression model is enhanced multiple folds by suitably distributing the computations and storage of data subsets across many machines. This article proposes a novel class of distributed Bayesian framework for process-based geostatistical models that contains many popular approaches, presents a comparative study of important approaches within this class, and contrasts their performance with representative non-distributed methods.

1.1 Non-distributed Methods for GP Modeling of Massive Spatial Data

Efficient GP-based models for massive spatial data have received extensive attention due to their great practical importance ([Heaton et al., 2019](#)). A common idea in GP-based modeling is to seek dimension-reduction by endowing the spatial covariance matrix either with a low-rank or a sparse structure. Low-rank structures on the spatial covariance matrix are the most widely used tool for efficient spatial computation. They represent the spatial surface using *r priori* chosen basis functions with associated computational complexity of $O(nr^2 + r^3)$ ([Cressie and Johannesson, 2008](#), [Banerjee et al., 2008](#), [Finley et al., 2009](#), [Guhaniyogi et al., 2011](#), [Wikle, 2010](#)); however, a major shortcoming of the above methods is that a small (r/n) -ratio yields inaccurate GP approximations, resulting in the propensity to oversmooth the spatial data ([Stein, 2014](#), [Simpson et al., 2012](#)).

A specific form of sparse structure, called covariance tapering, uses compactly supported covariance functions to create sparse spatial covariance matrices that approximate the full covariance matrix ([Kaufman et al., 2008](#), [Furrer et al., 2006](#), [Daley et al., 2015](#), [Bevilacqua et al., 2020](#)). Covariance tapering still requires expensive determinant evaluation of the massive covariance matrix, and the choice of the taper range can be difficult for spatial data over irregularly spaced locations ([Anderes et al., 2013](#)). An alternative approach is to introduce sparsity in the inverse covariance (precision) matrix of the GP likelihoods using products of lower dimensional conditional distributions ([Vecchia, 1988](#), [Rue et al., 2009](#), [Stein et al., 2004](#)), or via composite likelihoods ([Eidsvik et al., 2014](#), [Bai et al., 2012](#), [Bevilacqua and Gaetan, 2015](#)). Extending these ideas, recent approaches introduce sparsity in the inverse covariance (precision) matrix of process realizations and hence enable “kriging” at arbitrary locations ([Datta et al., 2016](#), [Guinness, 2018](#), [Finley et al., 2019a](#)). In related literature on computer experiments, localized approximations of GP models are proposed; see, for example, [Gramacy and Apley \(2015\)](#), [Gramacy and Haaland \(2016\)](#). These methods scale well with the sample size and are able to capture local spatial variations.

The remaining variants of dimension-reduction methods combine the benefits of low-rank and sparse covariance functions. Examples include non-stationary models (Banerjee et al., 2014) and multi-resolution models (Nychka et al., 2015, Katzfuss, 2017, Guinness, 2019, Katzfuss and Guinness, 2021, Guhaniyogi and Sanso, 2017). Multi-resolution models are difficult to implement and lack large sample theoretical guarantees, but they successfully capture spatial variation at multiple scales and are computationally efficient. The GP with Matérn covariance can be viewed as the solution of a stochastic partial differential equation. This observation has motivated GP approximations (Lindgren et al., 2011, Bolin and Lindgren, 2013), including a recent extension to multivariate non-Gaussian models with marginal Matérn covariance functions (Bolin and Wallin, 2020). This class of methods work well for Matérn covariance functions but are inapplicable in scaling GP with low-rank or non-stationary covariance functions.

1.2 Distributed Bayes

Rooted in the divide-and-conquer technique, the distributed Bayesian methods do not belong to any of the classes of methods in Section 1.1. They instead fit an existing model on different data subsets exploiting the distributed computing architecture. The results from the subsets are combined using an aggregation algorithm. These methods were first proposed in machine learning, including Consensus Monte Carlo (Scott et al., 2016), the Weierstrass sampler (Wang and Dunson, 2013), the semiparametric density product (Neiswanger et al., 2014), the median posterior (Minsker et al., 2014) and the Wasserstein posterior (Srivastava et al., 2015). Most of these methods are developed only for independent data. Recently, distributed Bayes has been applied to a variety of statistical problems in both modeling and computation, such as density estimation (Su, 2020), modeling of multivariate binary data (Mehrotra et al., 2021), sequential Monte Carlo (Lindsten et al., 2017), random partition trees (Wang et al., 2015), etc. For GP models, Zhang and Williamson (2019) proposes to combine GP fitted on different data subsets via an importance-sampled mixture-of-experts model. Theoretical results on distributed GP inference have been developed recently (Cheng and Shang, 2017, Szabo and van Zanten, 2019, Shang et al., 2019). Nevertheless, these theoretical works and applications have mainly focused on univariate domains for nonparametric regression and have not considered the GP-based models used in spatial applications such as GP with Matérn covariance on a spatial domain.

On the spatial front, Barbian and Assunção (2017) propose combining point estimates of spatial parameters obtained from different subsets, but they do not provide combined inference on the spatial processes or predictions. Similarly, Heaton et al. (2017) partition the spatial domain and assume independence between the data in different partitions. Guhaniyogi and Banerjee (2018, 2019) propose the idea of “meta-posterior,” a computationally efficient approximation to the full data posterior. This approach does not assume independence across data blocks and enables accurate prediction with uncertainty (Heaton et al., 2019); however, Guhaniyogi and Banerjee (2018) does not offer any theoretical guidance on choosing the number of subsets for optimal inference on the spatial surface.

Instead of developing a new spatial GP regression model, we describe a general class of three-step distributed Bayesian approaches for extending an existing process-based geostatistical model, which includes a number of important special

cases. To implement the general approach, the n spatial locations are divided into k subsets such that each subset has representative data samples from all regions of the spatial domain with the j th subset containing m_j data samples. Second, posterior computations are implemented in parallel on the k subsets using any chosen spatial process model after raising the model likelihood to a power of n/m_j in the j th subset. The pseudo posterior distribution obtained using the modified likelihood is called the “subset pseudo posterior distribution”. Since j th subset pseudo posterior distribution conditions on (m_j/n) -fraction of the full data, the modification of the likelihood by raising it to the power of n/m_j ensures that variance of each subset pseudo posterior is of the same order (as a function of n) as that of the full data posterior distribution. Third, the k subset pseudo posterior distributions are combined into a single pseudo probability distribution, called the combined pseudo posterior, that conditions on the full data and replaces the computationally expensive full data posterior distribution for prediction and inference. Our distributed framework leverages existing spatial GP regression models and enhances their scalability by embedding them within the three-step framework. For example, Section 3.1 embeds full-rank and low-rank spatial GP regression models within the distributed framework and Section 3.3 discusses various methods for combining the k subset pseudo posteriors.

The proposed framework builds on the recent works that combine the subset pseudo posterior distributions through their geometric centers (e.g., mean, median) and guarantee wide applicability under general assumptions (Minsker et al., 2014, Srivastava et al., 2015, Li et al., 2017, Minsker et al., 2017, Savitsky and Srivastava, 2018, Srivastava et al., 2018, Minsker, 2019). The theory and practice of such distributed approaches are limited to parametric models. In contrast, the framework proposed here is tuned for accurate and computationally efficient posterior inference in nonparametric Bayesian models based on GP priors. In particular, we develop a new approach to modify the likelihood for computing the subset pseudo posterior distribution of an unknown function, an infinite-dimensional parameter, that subsumes the parametric distributed methods. We offer novel theoretical results on the convergence rate of the combined pseudo posterior to the true function. Finally, we also provide guidance on choosing k depending on the covariance function and n such that the combined pseudo posterior maintains near minimax optimal performance as $n \rightarrow \infty$. The proposed distributed framework delivers principled Bayesian inference and predictions without any restrictive data- or model-specific assumptions, such as the independence between data subsets or independence between blocks of parameters.

A related focus of this article is to present a comparative study of the proposed class of distributed approaches with important non-distributed approaches for modeling massive spatial data. We illustrate the application of the distributed framework for enhancing the scalability of spatial models with a low-rank non-stationary GP prior called the modified predictive process (MPP) prior (Finley et al., 2009). This prior is commonly used for estimating spatial surfaces in applications with massive sample size, but it struggles to provide accurate inference in a manageable time beyond (approximately) 10^4 observations. We embed MPP within our distributed framework and scale it to spatial applications of much bigger sizes and assess its performance relative to other distributed and state-of-the-art non-distributed alternatives for efficient spatial GP modeling. Unfortunately,

there is no theoretical guarantee for convergence of the Markov chain to its stationary distribution, where MCMC samples are drawn from the subset pseudo posteriors with an MPP prior on spatial effects; however, we find strong empirical evidence for it and propose to develop the theoretical support elsewhere.

2. BAYESIAN INFERENCE IN GP-BASED SPATIAL MODELS

Consider the model for the data observed at location $\mathbf{s}$ in a compact domain $\mathcal{D}$,

$$(1) \quad y(\mathbf{s}) = \mathbf{x}(\mathbf{s})^T \boldsymbol{\beta} + w(\mathbf{s}) + \epsilon(\mathbf{s}),$$

where $y(\mathbf{s})$ and $\mathbf{x}(\mathbf{s})$ are the response and a $p \times 1$ predictor vector respectively at $\mathbf{s}$, $\boldsymbol{\beta}$ is a $p \times 1$ predictor coefficient, $w(\mathbf{s})$ is the value of an unknown spatial function $w(\cdot)$ at $\mathbf{s}$, and $\epsilon(\mathbf{s})$ is the value of a white-noise process $\epsilon(\cdot)$ at $\mathbf{s}$, which is independent of $w(\cdot)$. The Bayesian implementation of the model in (1) customarily assumes that (a) $\boldsymbol{\beta}$ apriori follows $N(\boldsymbol{\mu}_\beta, \boldsymbol{\Sigma}_\beta)$ and (b) $w(\cdot)$ and $\epsilon(\cdot)$ apriori follow mean 0 GPs with covariance functions $C_\alpha(\mathbf{s}_1, \mathbf{s}_2)$ and $D_\alpha(\mathbf{s}_1, \mathbf{s}_2)$ that model $\text{cov}\{w(\mathbf{s}_1), w(\mathbf{s}_2)\}$ and $\text{cov}\{\epsilon(\mathbf{s}_1), \epsilon(\mathbf{s}_2)\}$, respectively, where α are the process parameters indexing the two families of covariance functions and $\mathbf{s}_1, \mathbf{s}_2 \in \mathcal{D}$; therefore, the parameters are $\boldsymbol{\Omega} = \{\alpha, \beta\}$. The training data consists of predictors and responses observed at n spatial locations, denoted as $\mathcal{S} = \{\mathbf{s}_1, \dots, \mathbf{s}_n\}$.

Standard Markov chain Monte Carlo (MCMC) algorithms exist for performing posterior inference on $\boldsymbol{\Omega}$ and $w(\cdot)$ at a set of locations $\mathcal{S}^* = \{\mathbf{s}_1^*, \dots, \mathbf{s}_l^*\}$, where $\mathcal{S}^* \cap \mathcal{S} = \emptyset$, and for predicting $y(\mathbf{s}^*)$ for any $\mathbf{s}^* \in \mathcal{S}^*$ (Banerjee et al., 2014). Given $\mathcal{S}$, the prior assumptions on $w(\cdot)$ and $\epsilon(\cdot)$ imply that $\mathbf{w}^T = \{w(\mathbf{s}_1), \dots, w(\mathbf{s}_n)\}$ and $\boldsymbol{\epsilon}^T = \{\epsilon(\mathbf{s}_1), \dots, \epsilon(\mathbf{s}_n)\}$ are independent and follow $N\{\mathbf{0}, \mathbf{C}(\alpha)\}$ and $N\{\mathbf{0}, \mathbf{D}(\alpha)\}$, respectively, with the (i, j) th entries of $\mathbf{C}(\alpha)$ and $\mathbf{D}(\alpha)$ are $C_\alpha(\mathbf{s}_i, \mathbf{s}_j)$ and $D_\alpha(\mathbf{s}_i, \mathbf{s}_j)$, respectively. The hierarchy in (1) is completed by assuming that α apriori follows a distribution with density $\pi(\alpha)$. The MCMC algorithm for sampling $\boldsymbol{\Omega}$, $\mathbf{w}^{*T} = \{w(\mathbf{s}_1^*), \dots, w(\mathbf{s}_l^*)\}$, and $\mathbf{y}^{*T} = \{y(\mathbf{s}_1^*), \dots, y(\mathbf{s}_l^*)\}$ cycle through the following three steps until sufficient MCMC samples are drawn post convergence:

1. Integrate over $\mathbf{w}$ in (1) and

(a) sample β given $\mathbf{y}, \mathbf{X}, \alpha$ from $N(\mathbf{m}_\beta, \mathbf{V}_\beta)$, where

$$(2) \quad \mathbf{V}_\beta = \left\{ \mathbf{X}^T \mathbf{V}(\alpha)^{-1} \mathbf{X} + \boldsymbol{\Sigma}_\beta^{-1} \right\}^{-1}, \quad \mathbf{m}_\beta = \mathbf{V}_\beta \left\{ \mathbf{X}^T \mathbf{V}(\alpha)^{-1} \mathbf{y} + \boldsymbol{\Sigma}_\beta^{-1} \boldsymbol{\mu}_\beta \right\},$$

$\mathbf{X} = [\mathbf{x}(\mathbf{s}_1) : \dots : \mathbf{x}(\mathbf{s}_n)]^T$ is the $n \times p$ matrix of predictors, with $p < n$, $\mathbf{V}(\alpha) = \mathbf{C}(\alpha) + \mathbf{D}(\alpha)$; and

(b) sample α given $\mathbf{y}, \mathbf{X}, \beta$ using the Metropolis-Hastings algorithm with a normal random walk proposal.

2. Sample $\mathbf{w}^*$ given $\mathbf{y}, \mathbf{X}, \alpha, \beta$ from $N(\mathbf{m}_*, \mathbf{V}_*)$, where

$$(3) \quad \mathbf{V}_* = \mathbf{C}_{*,*}(\alpha) - \mathbf{C}_*(\alpha) \mathbf{V}(\alpha)^{-1} \mathbf{C}_*(\alpha)^T, \quad \mathbf{m}_* = \mathbf{C}_*(\alpha) \mathbf{V}(\alpha)^{-1} (\mathbf{y} - \mathbf{X} \beta),$$

$\mathbf{C}_*(\alpha)$ and $\mathbf{C}_{*,*}(\alpha)$ are $l \times n$ and $l \times l$ matrices, respectively, and the (i, j) th entries of $\mathbf{C}_{*,*}(\alpha)$ and $\mathbf{C}_*(\alpha)$ are $C_\alpha(\mathbf{s}_i^*, \mathbf{s}_j^*)$ and $C_\alpha(\mathbf{s}_i^*, \mathbf{s}_j)$, respectively.

3. Sample $\mathbf{y}^*$ given $\alpha, \beta, \mathbf{w}^*$ from $N\{\mathbf{X}^* \beta + \mathbf{w}^*, \mathbf{D}(\alpha)\}$, where $\mathbf{X}^{*T} = [\mathbf{x}(\mathbf{s}_1^*) : \dots : \mathbf{x}(\mathbf{s}_l^*)]$.

Many spatial models can be formulated in terms of (1) by assuming different forms of $C_\alpha(\mathbf{s}_1, \mathbf{s}_2)$ and $D_\alpha(\mathbf{s}_1, \mathbf{s}_2)$; see Banerjee et al. (2014) and supplementary material for details on the MCMC algorithm. Irrespective of the form of $\mathbf{D}(\alpha)$, if no additional assumptions are made on the structure of $\mathbf{C}(\alpha)$, then the three steps require $O(n^3)$ flops in computation and $O(n^2)$ memory units in storage in every MCMC iteration. Spatial models with this form of posterior computations are based on a *full-rank* GP prior, which are infeasible to compute for big data.

There are methods which either impose a low-rank structure or a sparse structure on $\mathbf{C}(\alpha)$ to address this computational issue (Banerjee et al., 2014). Methods with a low-rank structure on $\mathbf{C}(\alpha)$ expresses $\mathbf{C}(\alpha)$ in terms of $r \ll n$ basis functions, in turn inducing a *low-rank* GP prior. Again, a class of sparse structure uses compactly supported covariance functions to create $\mathbf{C}(\alpha)$ with overwhelming zero entries (Kaufman et al., 2008, Furrer et al., 2006), where as another variety of sparse structure imposes a Markov random field model on the joint distribution of $\mathbf{y}$ (Vecchia, 1988, Rue et al., 2009, Stein et al., 2004) or $\mathbf{w}$ (Datta et al., 2016, Guinness, 2018). We use the MPP prior as a representative example of this broad class of computationally efficient methods. Let $\mathcal{S}^{(0)} = \{\mathbf{s}_1^{(0)}, \dots, \mathbf{s}_r^{(0)}\}$ be a set of r locations, known as the “knots,” which may or may not intersect with $\mathcal{S}$. Let $\mathbf{c}(\mathbf{s}, \mathcal{S}^{(0)}) = \{C_\alpha(\mathbf{s}, \mathbf{s}_1^{(0)}), \dots, C_\alpha(\mathbf{s}, \mathbf{s}_r^{(0)})\}^T$ be an $r \times 1$ vector and $\mathbf{C}(\mathcal{S}^{(0)})$ be an $r \times r$ matrix whose (i, j) th entry is $C_\alpha(\mathbf{s}_i^{(0)}, \mathbf{s}_j^{(0)})$. Using $\mathbf{c}(\mathbf{s}_1, \mathcal{S}^{(0)}), \dots, \mathbf{c}(\mathbf{s}_n, \mathcal{S}^{(0)})$ and $\mathbf{C}(\mathcal{S}^{(0)})$, define the diagonal matrix $\boldsymbol{\delta} = \text{diag}\{\delta(\mathbf{s}_1), \dots, \delta(\mathbf{s}_n)\}$ with $\delta(\mathbf{s}_i) = C_\alpha(\mathbf{s}_i, \mathbf{s}_i) - \mathbf{c}^T(\mathbf{s}_i, \mathcal{S}^{(0)}) \mathbf{C}(\mathcal{S}^{(0)})^{-1} \mathbf{c}(\mathbf{s}_i, \mathcal{S}^{(0)})$, $i = 1, \dots, n$. Let $\mathbf{1}(\mathbf{a} = \mathbf{b}) = 1$ if $\mathbf{a} = \mathbf{b}$ and 0 otherwise. Then, MPP is a GP with covariance function

$$(4) \quad \tilde{C}_\alpha(\mathbf{s}_1, \mathbf{s}_2) = \mathbf{c}^T(\mathbf{s}_1, \mathcal{S}^{(0)}) \mathbf{C}(\mathcal{S}^{(0)})^{-1} \mathbf{c}(\mathbf{s}_2, \mathcal{S}^{(0)}) + \delta(\mathbf{s}_1) \mathbf{1}(\mathbf{s}_1 = \mathbf{s}_2),$$

where $\mathbf{s}_1, \mathbf{s}_2 \in \mathcal{D}$, $\tilde{C}_\alpha(\mathbf{s}_1, \mathbf{s}_2)$ depends on the covariance function of the parent GP and the selected r knots, which define $\mathbf{C}(\mathcal{S}^{(0)})$, $\mathbf{c}^T(\mathbf{s}_1, \mathcal{S}^{(0)})$, and $\mathbf{c}^T(\mathbf{s}_2, \mathcal{S}^{(0)})$. We have used a $\sim$ in (4) to distinguish the covariance function of a low-rank GP prior from that of its parent full-rank GP. If $\tilde{\mathbf{C}}(\alpha)$ is a matrix with (i, j) th entry $\tilde{C}_\alpha(\mathbf{s}_i, \mathbf{s}_j)$, then the posterior computations using MPP, a low-rank GP prior, replace $\mathbf{C}(\alpha)$ by $\tilde{\mathbf{C}}(\alpha)$ in the steps 1(a), 1(b), and 2. The (low) rank r structure imposed by $\mathbf{C}(\mathcal{S}^{(0)})$ implies that $\tilde{\mathbf{C}}(\alpha)^{-1}$ computation requires $O(nr^2)$ flops using the Woodbury formula (Harville, 1997); however, massive spatial data require that $r = O(\sqrt{n})$, leading to the computational inefficiency of low-rank methods.

The next section discusses a general three-step distributed framework to scale the posterior computations in spatial GP regression models with full-rank and low-rank GP priors. Briefly, the first and second steps divide the full data and fit a low-rank or full-rank spatial GP regression model on each data subset after modifying the subset likelihood, respectively, and the third step combines draws from the all the subset pseudo posteriors. We discuss a few popular alternatives for combining draws from the subset pseudo posteriors and offer novel convergence rate results for an important subclass of combination approaches.

3. DISTRIBUTED FRAMEWORK FOR BAYESIAN INFERENCE IN SPATIAL REGRESSION MODELS

3.1 First Step: Partitioning of Spatial Locations

We partition the n spatial locations into k non-overlapping subsets. The default partitioning scheme is to randomly allocate the locations into k possibly non-overlapping subsets (referred to as the random partitioning scheme hereon) to ensure that each subset has representative data samples from all subregions of the domain. We provide discussion on the choice of k later.

Let $\mathcal{S}_j = \{\mathbf{s}_{j1}, \dots, \mathbf{s}_{jm_j}\}$ denote the set of m_j spatial locations in subset j ($j = 1, \dots, k$). Conceptually, a spatial location can belong to multiple subsets, though for this work we have assumed disjoint subsets, so that $\sum_{j=1}^k m_j = n$ and $\cup_{j=1}^k \mathcal{S}_j = \mathcal{S}$, where $\mathbf{s}_{ji} = \mathbf{s}_{i'}$ for some $\mathbf{s}_{i'} \in \mathcal{S}$ and for every $i = 1, \dots, m_j$ and $j = 1, \dots, k$. Denote the data in the j th partition as $\{\mathbf{y}_j, \mathbf{X}_j\}$ ($j = 1, \dots, k$), where $\mathbf{y}_j = \{y(\mathbf{s}_{j1}), \dots, y(\mathbf{s}_{jm_j})\}^T$ is a $m_j \times 1$ vector and $\mathbf{X}_j = [\mathbf{x}(\mathbf{s}_{j1}) : \dots : \mathbf{x}(\mathbf{s}_{jm_j})]^T$ is a $m_j \times p$ matrix of predictors corresponding to the spatial locations in $\mathcal{S}_j$ with $p < m_j$. In modern grid or cluster computing environments, all the machines in the network have similar computational power, so the performances of distributed Bayesian methods are optimized by choosing similar values of $m_1, \dots, m_k$.

One can choose more sophisticated partitioning schemes than random partitioning. For example, it is possible to cluster the data based on centroid clustering (Knorr-Held and Raßer, 2000) or hierarchical clustering based on spatial gradients (Anderson et al., 2014, Heaton et al., 2017), and then construct subsets such that each subsets contains representative data samples from each cluster. Detailed exploration later shows that even random partitioning leads to desirable inference in the various simulation settings and in the sea surface data example. Perhaps more sophisticated blocking methods may provide further improvement in the cases where spatial locations are drawn based on specific designs; for example, sophisticated partitioning schemes have inferential benefits when a sub-domain shows substantial local behavior compared to the others (Guhaniyogi and Sanso, 2017), or sampled locations are chosen based on a specific survey design. Since they are atypical examples in the spatial context, we will pursue them elsewhere.

The univariate spatial GP regression model for any location $\mathbf{s}_{ji} \in \mathcal{S}_j \subset \mathcal{D}$ is

$$(5) \quad y(\mathbf{s}_{ji}) = \mathbf{x}(\mathbf{s}_{ji})^T \boldsymbol{\beta} + w(\mathbf{s}_{ji}) + \epsilon(\mathbf{s}_{ji}), \quad i = 1, \dots, m_j.$$

Let $\mathbf{w}_j^T = \{w(\mathbf{s}_{j1}), \dots, w(\mathbf{s}_{jm_j})\}$ and $\boldsymbol{\epsilon}_j^T = \{\epsilon(\mathbf{s}_{j1}), \dots, \epsilon(\mathbf{s}_{jm_j})\}$ be the realizations of GP $w(\cdot)$ and white-noise process $\epsilon(\cdot)$, respectively, in the j th subset. After marginalizing over $\mathbf{w}_j$ in the GP-based model for the j th subset, the likelihood of $\boldsymbol{\Omega} = \{\boldsymbol{\alpha}, \boldsymbol{\beta}\}$ is given by $\ell_j(\boldsymbol{\Omega}) = N\{\mathbf{y}_j \mid \mathbf{X}_j \boldsymbol{\beta}, \mathbf{V}_j(\boldsymbol{\alpha})\}$, where $\mathbf{V}_j(\boldsymbol{\alpha}) = \mathbf{C}_j(\boldsymbol{\alpha}) + \mathbf{D}_j(\boldsymbol{\alpha})$ and $\mathbf{V}_j(\boldsymbol{\alpha}) = \tilde{\mathbf{C}}_j(\boldsymbol{\alpha}) + \mathbf{D}_j(\boldsymbol{\alpha})$ for full-rank and low-rank GP priors, respectively, and $\mathbf{C}_j(\boldsymbol{\alpha}), \tilde{\mathbf{C}}_j(\boldsymbol{\alpha}), \mathbf{D}_j(\boldsymbol{\alpha})$ are obtained by extending the definitions of $\mathbf{C}(\boldsymbol{\alpha}), \tilde{\mathbf{C}}(\boldsymbol{\alpha}), \mathbf{D}(\boldsymbol{\alpha})$ to the j th subset. The likelihood of $\mathbf{w}_j$ given $\mathbf{y}_j, \mathbf{X}_j$, and $\boldsymbol{\Omega}$ is $\ell_j(\mathbf{w}_j) = N\{\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta} \mid \mathbf{w}_j, \mathbf{D}_j(\boldsymbol{\alpha})\}$. The likelihoods in $\ell_j(\boldsymbol{\Omega})$ and $\ell_j(\mathbf{w}_j)$ yield the posterior distributions for $\boldsymbol{\beta}, \boldsymbol{\alpha}, \mathbf{w}^*, \mathbf{y}^*$ ($\mathbf{w}^*$ and $\mathbf{y}^*$ have already been defined in the second paragraph of Section 2) based on full-rank or low-rank GP priors and are called j th subset pseudo posterior distributions.

3.2 Second Step: Sampling From Subset Pseudo Posterior Distributions

We define subset pseudo posterior distributions by modifying the likelihoods in $\ell_j(\boldsymbol{\Omega})$ and $\ell_j(\mathbf{w}_j)$. More precisely, the density of the j th subset pseudo posterior distribution of $\boldsymbol{\Omega}$ is given by

$$(6) \quad \pi_{m_j}(\boldsymbol{\Omega} \mid \mathbf{y}_j) = \frac{\{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\Omega})}{\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\Omega}) d\boldsymbol{\Omega}},$$

where we assume that $\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\Omega}) d\boldsymbol{\Omega} < \infty$, and the subscript ‘ m_j ’ denotes that the density conditions on m_j data samples in the j th subset. The modification of likelihood to yield the subset pseudo posterior density in (6) is called *stochastic approximation* (Minsker et al., 2014). Raising the likelihood to the power of n/m_j is equivalent to replicating every $y(\mathbf{s}_{ji})$ n/m_j times ($i = 1, \dots, m_j$), so stochastic approximation accounts for the fact that the j th subset pseudo posterior distribution conditions on a (m_j/n) -fraction of the full data and ensures that its variance is of the same order (as a function of n) as that of the full data posterior distribution. Unlike parametric models, stochastic approximation in spatial regression models has not been studied previously in the literature.

We address this gap using the proposed stochastic approximation in (6). The full conditional densities of j th subset pseudo posterior distributions for prediction and inference follow from their full data counterparts. The j th full conditional densities of $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$ in the GP-based models are

$$\pi_{m_j}(\boldsymbol{\beta} \mid \mathbf{y}_j, \boldsymbol{\alpha}) = \frac{\{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\beta})}{\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\beta}) d\boldsymbol{\beta}}, \quad \pi_{m_j}(\boldsymbol{\alpha} \mid \mathbf{y}_j, \boldsymbol{\beta}) = \frac{\{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\alpha})}{\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\alpha}) d\boldsymbol{\alpha}},$$

where $\pi(\boldsymbol{\beta}) = N(\boldsymbol{\mu}_\beta, \boldsymbol{\Sigma}_\beta)$, $\pi(\boldsymbol{\alpha})$ is the prior density of $\boldsymbol{\alpha}$, and we assume that $\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\beta}) d\boldsymbol{\beta}$ and $\int \{\ell_j(\boldsymbol{\Omega})\}^{n/m_j} \pi(\boldsymbol{\alpha}) d\boldsymbol{\alpha}$ respectively are finite. The j th full conditional densities of $\mathbf{y}^*$ and $\mathbf{w}^*$ are calculated after modifying the likelihood of $\mathbf{w}_j$ using stochastic approximation. Given $\mathbf{y}_j$, $\mathbf{X}_j$, and $\boldsymbol{\Omega}$, straightforward calculation yields that the j th subset pseudo posterior predictive density of $\mathbf{w}^*$ is $\pi_{m_j}(\mathbf{w}^* \mid \mathbf{y}_j, \boldsymbol{\Omega}) = N(\mathbf{w}^* \mid \mathbf{m}_{j*}, \mathbf{V}_{j*})$, with

$$\mathbf{V}_{j*} = \mathbf{C}_{*,*}(\boldsymbol{\alpha}) - \mathbf{C}_{*j}(\boldsymbol{\alpha}) \mathbf{V}_j(\boldsymbol{\alpha})^{-1} \mathbf{C}_{*j}(\boldsymbol{\alpha})^T, \quad \mathbf{m}_{j*} = \mathbf{C}_{*j}(\boldsymbol{\alpha}) \mathbf{V}_j(\boldsymbol{\alpha})^{-1} (\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta}),$$

where $\mathbf{V}_j(\boldsymbol{\alpha}) = \mathbf{C}_j(\boldsymbol{\alpha}) + (n/m_j)^{-1} \mathbf{D}_j(\boldsymbol{\alpha})$ and $\tilde{\mathbf{V}}_j(\boldsymbol{\alpha}) = \tilde{\mathbf{C}}_j(\boldsymbol{\alpha}) + (n/m_j)^{-1} \mathbf{D}_j(\boldsymbol{\alpha})$ for full-rank and low-rank GP priors, respectively, and $\mathbf{C}_{*,*}(\boldsymbol{\alpha})$, $\mathbf{C}_{*j}(\boldsymbol{\alpha})$ are $l \times l$, $l \times m_j$ matrices obtained by extending the definition in (3) to subset j for full-rank and low-rank GP priors with covariance functions $C_\alpha(\cdot, \cdot)$ and $\tilde{C}_\alpha(\cdot, \cdot)$, respectively. We note that the stochastic approximation exponent, n/m_j , scales $\mathbf{D}_j(\boldsymbol{\alpha})$ in $\mathbf{V}_j(\boldsymbol{\alpha})$ so that the uncertainty in subset and full data posterior distributions are of the same order (as a function of n). The j th subset pseudo posterior predictive density of $\mathbf{y}^*$ given the MCMC samples of $\mathbf{w}^*$ and $\boldsymbol{\Omega}$ in the j th subset is $N\{\mathbf{y}^* \mid \mathbf{X}^* \boldsymbol{\beta} + \mathbf{w}^*, \mathbf{D}_j(\boldsymbol{\alpha})\}$.

We specialize the sampling algorithm (Steps 1–3) introduced in Section 2 to subset j ($j = 1, \dots, k$), sampling $\{\boldsymbol{\beta}, \boldsymbol{\alpha}, \mathbf{y}^*, \mathbf{w}^*\}$ in each subset across multiple MCMC iterations; see supplementary material for subset pseudo posterior sampling algorithms in the full-rank and low-rank GP priors. The computational complexity of j th subset pseudo posterior computations follows from their full

data counterparts if we replace n by m_j . Specifically, the computational complexities for sampling a subset pseudo posterior are $O(m^3)$ and $O(mr^2)$ flops per iteration if the model in (5) uses a full-rank or a low-rank GP prior, respectively, where $m = \max_j m_j$. Performing subset pseudo posterior computations in parallel across k servers also alleviates the need to store large covariance matrices. We hereon refer to subset pseudo posterior as subset posterior.

Our second step in the distributed framework resembles some existing methods based on the composite likelihood (Varin et al., 2011); for example, Chandler and Bate (2007) and Ribatet et al. (2012) construct pseudo likelihood to replace the full data likelihood, where the pseudo likelihood attempts to capture important features of the full data likelihood while offering computational efficiency. In the context of geostatistical modeling with GP or its variants, for computational efficiency, the pseudo likelihood will naturally be based on independence of data blocks at some level. To make up for the incorrect asymptotic distribution of the posterior distribution due to the incorrect independence assumption, they propose a number of adjustments in the composite log likelihood (e.g., the margin adjustment and the curvature adjustment). Similar to these approaches, the likelihood adjustment in each subset for the second step of our general distributed approach is also born out of consideration to scale the asymptotic variance of subset posteriors to the same order as the asymptotic variance of the full posterior; however, unlike these composite likelihood approaches, the distributed approaches we focus on do not assume any restrictive structure (e.g., block independence) in the data likelihood. In fact, there is no guarantee that the induced data likelihood that leads to the combined pseudo posterior for any distributed method assumes any block independence form. Moreover, Savitsky and Srivastava (2018) represents an example of embedding a composite likelihood in a distributed setup that computes the Wasserstein barycenter. Likewise, we believe that most of these “flexible” composite likelihoods can be used in extensions of the distributed framework for subset sampling in models where the true likelihood is unavailable or expensive to compute.

3.3 Third Step: Combination of Subset Posterior Distributions

We now discuss strategies for combining subset posteriors to construct a “combined pseudo posterior”, which is used as an computationally efficient alternative to the full data posterior. The combination strategies discussed here include representative approaches used for the distributed Bayesian inference in independent data, but they have not been studied empirically or theoretically for correlated spatial data setting. Specifically, we compare the following combination schemes with our approach: (i) Consensus Monte Carlo (CMC); (ii) Double Parallel Monte Carlo (DPMC); and (iii) combination through the Wasserstein barycenter.

Consensus Monte Carlo (CMC)

For a scalar or vector parameter of interest θ , Consensus Monte Carlo (CMC) (Scott et al., 2016) draws samples from an approximation of the full posterior. In our setting, θ can be taken as β , α , $\mathbf{w}^*$, $\mathbf{y}^*$, their individual components, or any functional of these parameters. Let $\{\theta_1^{(j)}, \dots, \theta_T^{(j)}\}$ denote the T posterior samples of θ generated from subset j post convergence. Based on the Bernstein-von Mises (BvM) theorem, Scott et al. (2016) proposed to use the weighted average $\sum_{j=1}^k w_j \theta_i^{(j)}$, $i = 1, \dots, T$ to approximate T samples from the full data

posterior, where the BvM theorem says that the full data posterior tends to a normal distribution centered around the true parameter value as n grows and w_j is the inverse of the empirical covariance matrix of $\{\theta_1^{(j)}, \dots, \theta_T^{(j)}\}$. This algorithm is exact when the samples are independent and each subset posterior is Gaussian, but this assumption is rarely satisfied in spatial applications.

Double Parallel Monte Carlo (DPMC)

Following the notation for CMC, let θ be the parameter of interest. Denote the average of θ draws on the subset j as $\bar{\theta}^{(j)} = (\theta_1^{(j)} + \dots + \theta_T^{(j)})/T$ ($j = 1, \dots, k$) and $\bar{\theta} = (\bar{\theta}^{(1)} + \dots + \bar{\theta}^{(k)})/k$ be their average. DPMC (Xue and Liang, 2019) re-centers each of the subset posteriors to $\bar{\theta}$ and then uses the mixture of re-centered subset posteriors, given by $\frac{1}{k} \sum_{j=1}^k \pi_{m_j}(\theta - \bar{\theta} + \bar{\theta}^{(j)} | \mathbf{y}_j)$, to approximate the full data posterior. Following the implementation of DPMC in the context of independent data, we simply transform $\theta_t^{(j)}$ to $\bar{\theta} + (\theta_t^{(j)} - \bar{\theta}^{(j)})$ ($t = 1, \dots, T; j = 1, \dots, k$) and treat them as draws from the combined posterior distribution.

Combining subset posteriors using Wasserstein barycenter

This combination algorithm relies on the notion of Wasserstein barycenter (Srivastava et al., 2015). If $\nu_1, \dots, \nu_k$ are the k subset posterior distributions of θ , then the combined pseudo posterior $\bar{\nu}$ is the Wasserstein barycenter that is defined as

$$(7) \quad \bar{\nu} = \operatorname{argmin}_{\nu \in \mathcal{P}_2(\Theta)} \frac{1}{k} \sum_{j=1}^k W_2^2(\nu, \nu_j), \quad W_2^2(\mu, \nu) = \inf_{\pi \in \Pi(\mu, \nu)} \int_{\Theta \times \Theta} \|x - y\|^2 d\pi(x, y),$$

where $\|\cdot\|$ is a metric on the parameter space Θ , $\mathcal{P}(\Theta)$ be the space of all probability measures on Θ , $\mathcal{P}_2(\Theta) = \{\mu \in \mathcal{P}(\Theta) : \int_{\Theta} \|\theta - \theta_0\|^2 \mu(d\theta) < \infty\}$, $W_2(\mu, \nu)$ is the Wasserstein distance between $\mu, \nu \in \mathcal{P}_2(\Theta)$, and $\Pi(\mu, \nu)$ is the space of all joint distributions of $\Theta \times \Theta$ with μ, ν as marginals. It is known that $\bar{\nu}$ exists and is unique (Agueh and Carlier, 2011).

In practice, ν_j is replaced by its empirical approximation obtained using the θ draws from subset j . A variety of efficient algorithms are available to provide an empirical approximation of $\bar{\nu}$ ($j = 1, \dots, k$) (Cuturi and Doucet, 2014). This approach for combining subset posteriors leads to the combined pseudo posterior referred to as the Wasserstein posterior (WASP), which is preferred over several other combination methods for independent data (Srivastava et al., 2018); for example, directly averaging over many subset posterior densities with different means can usually result in an undesirable multimodal pseudo posterior distribution, but the WASP does not have this problem; see Figure 1 in Srivastava et al. (2018). Besides, the WASP does not rely on the asymptotic normality of the subset posterior distributions as in other approaches, such as the CMC.

Computing the WASP with constraints

Computing the WASP is inefficient if k is large, so $\bar{\nu}$ is computed with additional constraints (Srivastava and Xu, 2021). One such approach constrains θ to be a one-dimensional functional of β , α , $\mathbf{w}^*$, or $\mathbf{y}^*$. For a scalar parameter, the Wasserstein barycenter of θ can be easily obtained by averaging empirical subset posterior quantiles (Li et al., 2017). We refer to this approach as distributed kriging (DISK) and the combined pseudo posterior is called as the DISK posterior. Let ν and ν_j be the full posterior and j th subset posterior distribution of θ , and $\bar{\nu}$ be the Wasserstein barycenter of $\nu_1, \dots, \nu_k$ as in (7). For any $q \in (0, 1)$, let $\hat{\nu}_j^q$

be the q th empirical quantile of ν_j based on the MCMC samples from ν_j , and $\hat{\bar{\nu}}^q$ be the q th quantile of the empirical version of $\bar{\nu}$. Then, $\hat{\bar{\nu}}^q$ can be computed as

$$(8) \quad \hat{\bar{\nu}}^q = \frac{1}{k} \sum_{j=1}^k \hat{\nu}_j^q, \quad q = \xi, 2\xi, \dots, 1 - \xi,$$

where ξ is the grid-size of the quantiles. If the ξ -grid is fine enough, then the θ draws from the marginal DISK distribution are obtained by inverting the empirical distribution function supported on the quantile estimates (Li et al., 2017). In practice, the primary interest often lies in the posterior distribution of some one-dimensional functional of θ ; therefore, the univariate WASP obtained by averaging quantiles in (8) accomplishes this with great generality and convenient implementation. Our simulation studies in Section 4 investigate if the multi-variate combination approaches in CMC, DPMC, or WASP lead to any notable improvement over the univariate quantile combination in (8).

The choice of the grid size is mainly determined by the Monte Carlo approximation error of each subset posterior. In general, the Monte Carlo approximation error to subset posteriors can be measured in terms of the size of MCMC samples (say T). This error is evaluated by taking T to infinity and differs from the statistical error, where n tends to infinity. In the context of distributed Bayesian inference for independent data, Theorem 3 in the supplementary material of Li et al. (2017) has shown that the Monte Carlo error is usually in some polynomial order of T such as $O(T^{-1/2})$ and $O(T^{-1/4})$ depending on the distance measure and is independent of the statistical error defined in terms of n . Following this intuition, in application, we usually draw at least 10^4 MCMC samples for each subset posterior and use all of them to construct the quantiles.

A key feature of the combination scheme for the four distributed approaches is that given the subset posterior MCMC samples, the combination step is agnostic to the choice of a model. Specifically, given MCMC samples from the k subset posterior distributions, (8) remains the same for models based on a full-rank GP prior, a low-rank GP prior, such as MPP, or any other model described in Section 1.1. Since the combination step over k subsets takes $O(k)$ flops for all four combination schemes and $k < n$, the total time for computing the empirical quantile estimates of the combined pseudo posterior in inference or prediction requires $O(k) + O(m^3)$ and $O(k) + O(rm^2)$ flops in models based on full-rank and low-rank GP priors, respectively. Assuming that we have abundant computational resources, k is chosen large enough so that $O(m^3)$ computations are feasible. This would enable applications of the proposed distributed framework in models based on both full-rank and low-rank GP priors in massive n settings.

3.4 Bayes L_2 -Risk: Bias-Variance Decomposition and Convergence Rates

In the distributed Bayesian setup, it is already known that when the data are independent and identically distributed (i.i.d.), the combined posterior distribution using the Wasserstein barycenter of subset posteriors approximates the full data posterior distribution at a near optimal parametric rate, under certain conditions as $n, k, m_1, \dots, m_k \rightarrow \infty$ (Li et al., 2017, Srivastava et al., 2018); however, in models based on spatial process, data are not i.i.d. and inference on the infinite dimensional true spatial surface is of primary importance. Few formal theoretical results are available in this nonparametric distributed Bayes setup. The recent

work (Szabo and van Zanten, 2019) has shown that combination using Wasserstein barycenter has optimal Bayes risk and adapts to the smoothness of $w_0(\cdot)$, the true but unknown $w(\cdot)$, in the Gaussian white noise model, which is a special case of (1) with additional smoothness assumptions on $w_0(\cdot)$.

We mainly focus on the theoretical properties of the DISK posterior of the mean surface $\mathbf{x}(\cdot)^T \boldsymbol{\beta} + w(\cdot)$, and our theoretical framework can be possibly extended to the other three combination schemes described in Section 3.3. For ease of presentation, we assume that $m_1 = \dots = m_k = m$ and $k = n/m$. Determining the appropriate order for k in terms of n is one of the key issues for all distributed statistical methods. Our theory below reveals that the number of subsets k cannot increase too fast with n , or equivalently, the subset size m cannot be too small, mainly because a small subset size m will result in larger random errors in the estimation from subset posterior distributions.

We formally explain the model setup for our theory development. Suppose that the data generation process follows the model (1) with the true parameter value $\boldsymbol{\Omega}_0 = (\boldsymbol{\alpha}_0, \boldsymbol{\beta}_0)$ and the true spatial surface $w_0(\cdot)$. We focus on the Bayes L_2 -risk of the DISK predictive posterior for the mean function in (1); that is, $\mathbf{x}(\mathbf{s}^*)^T \boldsymbol{\beta} + w(\mathbf{s}^*)$ for any testing location $\mathbf{s}^* \in \mathcal{S}$. To ease the complexity of our theory, we first present two theorems below for the simplified model

$$(9) \quad y(\mathbf{s}_i) = w(\mathbf{s}_i) + \epsilon(\mathbf{s}_i), \quad \epsilon(\mathbf{s}_i) \sim N(0, \tau^2), \quad w(\cdot) \sim \text{GP}\{0, \lambda_n^{-1} C_{\boldsymbol{\alpha}}(\cdot, \cdot)\},$$

for $i = 1, \dots, n$. Compared to the spatial model (1), the model (9) does not contain the regression term $\mathbf{x}(\mathbf{s})^T \boldsymbol{\beta}$; however, our theory includes this regression term later by modifying the covariance function; see Corollary 3.3 below. The tuning parameter λ_n is a user-chosen deterministic sequence that depends on n . In real applications, one can simply set $\lambda_n = 1$, but one can also choose λ_n such that the posterior convergence rate is minimax optimal; see Theorem 3.2 below and the discussions therein.

We introduce some theoretical definitions used in stating our results. Let $\boldsymbol{\alpha}_0$ be the true kernel parameter. Let $\mathbb{P}_{\mathbf{s}}$ be a design distribution of $\mathbf{s}$ over $\mathcal{D}$, $L_2(\mathbb{P}_{\mathbf{s}})$ be the L_2 space under $\mathbb{P}_{\mathbf{s}}$, the inner product in $L_2(\mathbb{P}_{\mathbf{s}})$ is defined as $\langle f, g \rangle_{L_2(\mathbb{P}_{\mathbf{s}})} = \mathbb{E}_{\mathbb{P}_{\mathbf{s}}}(fg)$ for any $f, g \in L_2(\mathbb{P}_{\mathbf{s}})$ where $\mathbb{E}_{\mathbb{P}_{\mathbf{s}}}(\cdot)$ represents an expectation taken with respect to the distribution, $\mathbb{P}_{\mathbf{s}}$. For any $f \in L_2(\mathbb{P}_{\mathbf{s}})$ and $\mathbf{s} \in \mathcal{D}$, define the linear operator $(T_{\boldsymbol{\alpha}_0} f)(\mathbf{s}) = \int_{\mathcal{D}} C_{\boldsymbol{\alpha}_0}(\mathbf{s}, \mathbf{s}') f(\mathbf{s}') d\mathbb{P}_{\mathbf{s}}(\mathbf{s}')$. According to the Mercer's theorem, there exists an orthonormal basis $\{\varphi_i(\mathbf{s})\}_{i=1}^{\infty}$ in $L_2(\mathbb{P}_{\mathbf{s}})$, such that $C_{\boldsymbol{\alpha}_0}(\mathbf{s}, \mathbf{s}') = \sum_{i=1}^{\infty} \mu_i \varphi_i(\mathbf{s}) \varphi_i(\mathbf{s}')$, where $\mu_1 \geq \mu_2 \geq \dots \geq 0$ are the eigenvalues and $\{\varphi_i(\mathbf{s})\}_{i=1}^{\infty}$ are the eigenfunctions of $T_{\boldsymbol{\alpha}_0}$. The trace of the kernel $C_{\boldsymbol{\alpha}_0}$ is defined as $\text{tr}(C_{\boldsymbol{\alpha}_0}) = \sum_{i=1}^{\infty} \mu_i$. Any $f \in L_2(\mathbb{P}_{\mathbf{s}})$ has the series expansion $f(\mathbf{s}) = \sum_{i=1}^{\infty} \theta_i \varphi_i(\mathbf{s})$, where $\theta_i = \langle f, \varphi_i \rangle_{L_2(\mathbb{P}_{\mathbf{s}})}$. The reproducing kernel Hilbert space (RKHS) $\mathbb{H}$ attached to $C_{\boldsymbol{\alpha}_0}$ is the space of all functions $f \in L_2(\mathbb{P}_{\mathbf{s}})$ such that the $\mathbb{H}$ -norm $\|f\|_{\mathbb{H}}^2 = \sum_{i=1}^{\infty} \theta_i^2 / \mu_i < \infty$. The RKHS $\mathbb{H}$ is the completion of the linear space of functions defined as $\sum_{i=1}^I a_i C_{\boldsymbol{\alpha}_0}(\mathbf{s}_i, \cdot)$, where I is a positive integer, $\mathbf{s}_i \in \mathcal{D}$, and $a_i \in \mathbb{R}$ ($i = 1, \dots, I$); see van der Vaart and van Zanten (2008) for more details on RKHS.

We impose the following assumptions.

- A.1 (Sampling) The locations $\mathcal{S} = \{\mathbf{s}_1, \dots, \mathbf{s}_n\}$ and $\mathbf{s}^*$ are independently drawn from the same sampling distribution $\mathbb{P}_{\mathbf{s}}$. $\mathcal{S}_1, \dots, \mathcal{S}_k$ is a random disjoint partition of $\mathcal{S}$, each with size $m = n/k$.

- A.2 (True model) The true function w_0 is an element of the RKHS $\mathbb{H}$ attached to the kernel C_{α_0} . At a location $\mathbf{s}$, the observation is $y(\mathbf{s}) = w_0(\mathbf{s}) + \epsilon(\mathbf{s})$, where $\epsilon(\mathbf{s})$ is a white noise process with the true variance $\tau_0^2 < \infty$.
- A.3 (Trace class kernel) $\text{tr}(C_{\alpha_0}) < \infty$.
- A.4 (Moment condition) There are positive constants ρ and $q > 4$ such that $\mathbb{E}_{\mathbb{P}_{\mathbf{s}}} \{\varphi_i^{2q}(\mathbf{s})\} \leq \rho^{2q}$ for every $i \in \mathbb{N}$.

The random partition in A.1 guarantees that each individual subset $\mathcal{S}_j$ ($j = 1, \dots, k$) is a random sample from $\mathbb{P}_{\mathbf{s}}$. In general, the RKHS $\mathbb{H}$ in A.2 can be a smaller space relative to the support of the GP prior. While we use $w_0 \in \mathbb{H}$ in A.2 mainly for technical simplicity, this assumption can be possibly relaxed by considering sieves with increasing $\mathbb{H}$ -norms, similar to Assumption B' and Theorem 2 in Zhang et al. (2015). Furthermore, A.2 only requires that the true unknown error distribution to have a finite variance. Although we fit the data using the normal error in model (9), our theory below allows this error distribution to be misspecified and not normal; therefore, our posterior convergence rate results also hold for heavy-tailed error distributions such as t_4 , which are more general than van der Vaart and van Zanten (2011) whose techniques fully depend on the normal error assumption. In A.3, $\text{tr}(C_{\alpha})$ measures the size of the covariance function and imposes conditions on the regularity of functions that DISK can learn. A.4 on the eigenfunctions controls the error in approximating $C_{\alpha_0}(\mathbf{s}, \mathbf{s}')$ by a finite sum, similar to Assumption A in Zhang et al. (2015).

We first consider the case where both the error variance τ^2 and the kernel parameter α are fixed and known, similar to van der Vaart and van Zanten (2011). We extend our results to a special case where τ^2 is assigned a prior with bounded support in Corollary 1.1 of the supplementary material.

- A.5 (Fixed parameters) α and τ^2 are fixed at their true values $\alpha = \alpha_0$, $\tau^2 = \tau_0^2$.

We begin by examining the Bayes L_2 -risk of the DISK posterior for estimating w_0 in (9). Let $\bar{w}(\mathbf{s}^*)$ be a random variable that follows the DISK posterior for estimating $w_0(\mathbf{s}^*)$. Let $\mathbb{E}_{\mathbf{s}^*}$, $\mathbb{E}_{\mathcal{S}}$, and $\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*}$ respectively be the expectations with respect to the distributions of $\mathbf{s}^*$, $\mathcal{S}$, and $\{\mathbf{y}, \bar{w}(\mathbf{s}^*)\}$ given $\mathcal{S}, \mathbf{s}^*$. Given the random partition assumption in A.1, each individual subset $\mathcal{S}_j$ ($j = 1, \dots, k$) is a random sample from $\mathbb{P}_{\mathbf{s}}$. By A.5, we can drop the subscript “0” in α_0 and τ_0^2 . Then, $\bar{w}(\mathbf{s}^*)$ given $\mathbf{y}, \mathcal{S}, \mathbf{s}^*$ has the density $N(\bar{m}, \bar{v})$, where

$$\begin{aligned} \bar{m} &= \frac{1}{k} \sum_{j=1}^k \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{y}_j, \\ (10) \quad \bar{v}^{1/2} &= \frac{1}{k} \sum_{j=1}^k v_j^{1/2}, \quad v_j = \lambda_n^{-1} \left\{ c_{*,*} - \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{c}_{j,*} \right\}, \end{aligned}$$

$c_{*,*} = \text{cov}\{w(\mathbf{s}^*), w(\mathbf{s}^*)\}$, and $\mathbf{c}_{j,*}^T = [\text{cov}\{w(\mathbf{s}_{j1}), w(\mathbf{s}^*)\}, \dots, \text{cov}\{w(\mathbf{s}_{jm}), w(\mathbf{s}^*)\}]$. The Bayes L_2 -risk of DISK in estimating w_0 is $\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2$. This risk can be used to quantify how quickly the DISK posterior concentrates around the unknown true surface $w_0(\cdot)$ as the total sample size n increases to infinity. The convergence rate of this Bayes L_2 -risk towards zero also gives the posterior contraction rate of the DISK posterior defined in the same way as in Bayesian nonparametrics, such as van der Vaart and van Zanten (2011, Theorem

2). When the parameters τ^2 and α are fixed and known, it is straightforward to show (see the proof of Theorem 3.1 in the supplementary material) that this Bayes L_2 -risk can be decomposed into the squared bias, the variance of subset posterior means, and the variance of DISK posterior terms as

$$\begin{aligned} \text{bias}^2 &= \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{w}_0 - w_0(\mathbf{s}^*) \}^2, \\ \text{var}_{\text{mean}} &= \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-2} \mathbf{c}_* \}, \\ (11) \quad \text{var}_{\text{DISK}} &= \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \bar{v}(\mathbf{s}^*) \}, \end{aligned}$$

where $\bar{v}(\mathbf{s}^*) = \mathbb{E}_{\mathbf{y}|\mathcal{S}} [\text{var}\{\bar{w}(\mathbf{s}^*) \mid \mathbf{y}\}]$, $\mathbf{c}_*^T = (\mathbf{c}_{1,*}^T, \dots, \mathbf{c}_{k,*}^T)$, $\mathbf{w}_{0j} = \{w_0(\mathbf{s}_{j1}), \dots, w_0(\mathbf{s}_{jk})\}$ for $j = 1, \dots, k$, $\mathbf{w}_0^T = (\mathbf{w}_{01}, \dots, \mathbf{w}_{0k})$, and $\mathbf{L}$ is a block-diagonal matrix with $\mathbf{C}_{1,1}, \dots, \mathbf{C}_{k,k}$ along the diagonal. The next theorem provides theoretical upper bounds for each of the three terms in (11).

Theorem 3.1 *If Assumptions A.1–A.5 hold, then*

$$\begin{aligned} \text{Bayes } L_2 \text{ risk} &= \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|\mathcal{S}, \mathbf{s}^*} \{ \bar{w}(\mathbf{s}_*) - w_0(\mathbf{s}_*) \}^2 \\ &= \text{bias}^2 + \text{var}_{\text{mean}} + \text{var}_{\text{DISK}}, \\ \text{bias}^2 &\leq \frac{8\tau^2 \lambda_n}{n} \|w_0\|_{\mathbb{H}}^2 \\ &\quad + \|w_0\|_{\mathbb{H}}^2 \inf_{d \in \mathbb{N}} \left[\frac{8n}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \mu_1 R(m, n, d, q) \right], \\ \text{var}_{\text{mean}} &\leq \left(\frac{2n}{k \lambda_n} + \frac{4\|w_0\|_{\mathbb{H}}^2}{k} \right) \inf_{d \in \mathbb{N}} \left[\mu_{d+1} + \frac{12n}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \right. \\ &\quad \left. + R(m, n, d, q) \right] + \frac{12\tau^2 \lambda_n}{kn} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right), \\ \text{var}_{\text{DISK}} &\leq 3 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) + \inf_{d \in \mathbb{N}} \left[\left\{ \frac{4n}{\tau^2 \lambda_n^2} \text{tr}(C_{\alpha}) + \frac{1}{\lambda_n} \right\} \text{tr}(C_{\alpha}^d) \right. \\ (12) \quad &\quad \left. + \lambda_n^{-1} \text{tr}(C_{\alpha}) R(m, n, d, q) \right], \end{aligned}$$

where $\mathbb{N}$ is the set of all positive integers, A is a global positive constant that does not depend on any of the quantities here, and

$$\begin{aligned} b(m, d, q) &= \max \left(\sqrt{\max(q, \log d)}, \frac{\max(q, \log d)}{m^{1/2-1/q}} \right), \\ R(m, n, d, q) &= \left\{ \frac{A \rho^2 b(m, d, q) \gamma(\tau^2 \lambda_n / n)}{\sqrt{m}} \right\}^q, \\ \gamma(a) &= \sum_{i=1}^{\infty} \frac{\mu_i}{\mu_i + a} \text{ for any } a > 0, \quad \text{tr}(C_{\alpha}^d) = \sum_{i=d+1}^{\infty} \mu_i. \end{aligned}$$

These upper bounds are similar to the bounds obtained in Theorem 1 of Zhang et al. (2015) for the frequentist distributed estimator in kernel ridge regression. Although the upper bounds in (12) appear very complicated and involve many terms, the dominant term among them is $\frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right)$, where the function $\gamma(\cdot)$ is

related to the “effective dimensionality” of the covariance function C_α (Zhang, 2005). This term determines how fast the Bayes L_2 -risk converges to zero, as long as k is chosen to be some proper order of n such that all the other terms in the upper bounds of (12) can be made negligible compared to $\frac{\tau^2}{n}\gamma\left(\frac{\tau^2\lambda_n}{n}\right)$. In particular, the term $R(m, n, d, q)$ that quantifies the random error and appears in the infimums in all three upper bounds of (12) generally decreases with m and increases with k ; therefore, to ensure the dominance of $\frac{\tau^2}{n}\gamma\left(\frac{\tau^2\lambda_n}{n}\right)$, k cannot increase too fast with n ; see Theorem 3.2 below.

In contrast to the frequentist literature such as Zhang et al. (2015), a significant difference in our Theorem 3.1 is that our risk bounds involve two different variance terms. Our analysis naturally introduces the variance term var_{DISK} that corresponds to the variance of the DISK posterior distribution, while frequentist kernel ridge regression only finds a point estimate of w_0 and thus does not include this variance term. Each of the three upper bounds in Theorem 3.1 can be made close to zero as n increases to ∞ and k is chosen to grow at an appropriate rate depending on n . The next theorem finds the appropriate order for k in terms of n , such that the DISK posterior achieves nearly minimax optimal rates in its Bayes L_2 -risk (12), for three types of commonly used covariance functions/kernels, (i) degenerate kernels, (ii) kernels with exponentially decaying eigenvalues, and (iii) kernels with polynomially decaying eigenvalues. The kernel C_α is a degenerate kernel of rank d^* if there is some constant positive integer d^* such that $\mu_1 \geq \mu_2 \geq \dots \geq \mu_{d^*} > 0$ and $\mu_{d^*+1} = \mu_{d^*+2} = \dots = \mu_\infty = 0$.

Theorem 3.2 *If Assumptions A.1–A.5 hold, then as $n \rightarrow \infty$,*

- (i) *if C_α is a degenerate kernel of rank d^* , $\lambda_n = 1$, and $k \leq cn^{\frac{q-4}{q-2}}/(\log n)^{\frac{2q}{q-2}}$ for some constant $c > 0$, then the Bayes L_2 -risk of DISK posterior satisfies $\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 = O(n^{-1})$;*
- (ii) *if $\mu_i \leq c_{1\mu} \exp(-c_{2\mu} i^\kappa)$ for some constants $c_{1\mu} > 0, c_{2\mu} > 0, \kappa > 0$ and all $i \in \mathbb{N}$, $\lambda_n = 1$, and for some constant $c > 0$, $k \leq cn^{\frac{q-4}{q-2}}/(\log n)^{\frac{2(q\kappa+q-1)}{\kappa(q-2)}}$, then the Bayes L_2 -risk of DISK posterior satisfies $\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 = O\{(\log n)^{1/\kappa}/n\}$;*
- (iii) *if $\mu_i \leq c_\mu i^{-2\eta}$ for some constants $c_\mu > 0, \eta > \frac{q-1}{q-4}$ and all $i \in \mathbb{N}$, $\lambda_n = 1$, and for some constant $c > 0$, $k \leq cn^{\frac{(q-4)\eta-(q-1)}{(q-2)\eta}}/(\log n)^{\frac{2q}{q-2}}$, then the Bayes L_2 -risk of DISK posterior satisfies $\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 = O\left(n^{-\frac{2\eta-1}{2\eta}}\right)$; and*
- (iv) *if $\mu_i \leq c_\mu i^{-2\eta}$ for some constants $c_\mu > 0, \eta > \frac{q-1}{q-4}$ and all $i \in \mathbb{N}$, $\lambda_n = c_1 n^{1/(2\eta+1)}$, and $k \leq c_2 n^{\frac{(2\eta-1)q-8\eta}{(q-2)(2\eta+1)}}/(\log n)^{\frac{2q}{q-2}}$ for some positive constants c_1, c_2 , then the Bayes L_2 -risk of DISK posterior satisfies $\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 = O\left(n^{-\frac{2\eta}{2\eta+1}}\right)$.*

In Theorem 3.2, the space of w_0 is the RKHS $\mathbb{H}$ attached to C_α by Assumption A.2. In Case (i), the RKHS of C_α is a d^* -dimensional space of functions. For example, the covariance functions in subset of regressors approximation (Quiñonero-Candela and Rasmussen, 2005) and predictive process (Banerjee et al., 2008) are both degenerate with their ranks equaling the number of inducing variables and

knots, respectively. One example of Case (ii) is the squared exponential kernel, which is popular in machine learning. The squared exponential kernel defined on $\mathbb{R}$ with $\mathbb{P}_s$ being a Gaussian measure has exponentially decaying eigenvalues (Zhu et al., 1998), and its RKHS only contains functions with infinite smoothness. The rate of decay of the L_2 -risks in Case (i) and Case (ii) with $\kappa = 2$ are known to be minimax optimal (Raskutti et al., 2012, Yang et al., 2017).

Cases (iii) and (iv) apply to the class of kernels with polynomially decaying eigenvalues. For example, consider the Matérn covariance function $C_{\sigma^2, \phi, \nu}(\mathbf{s}, \mathbf{s}') = \sigma^2 \frac{2^{1-\nu}}{\Gamma(\nu)} (\phi \|\mathbf{s} - \mathbf{s}'\|)^\nu \mathcal{K}_\nu(\phi \|\mathbf{s} - \mathbf{s}'\|)$, where $\mathbf{s}, \mathbf{s}' \in \mathcal{D} \subset \mathbb{R}^d$, $\sigma^2 > 0$, $\phi > 0$, $\boldsymbol{\alpha} = (\sigma^2, \phi)$, $\nu \geq d/2$ is known, $\Gamma(\cdot)$ is the gamma function, and $\mathcal{K}_\nu(\cdot)$ is the modified Bessel function of the second kind. Then the RKHS of $C_{\sigma^2, \phi, \nu}(\mathbf{s}, \mathbf{s}')$ defined on a compact domain $\mathcal{D}$ with Lipschitz boundary is norm equivalent to the Sobolev space with order $\nu + d/2$ (Wendland 2005, Corollary 10.48). Furthermore, when $\mathbb{P}_s$ is the uniform distribution on $\mathcal{D}$, the eigenvalues of Matérn kernels decay as $\mu_i \leq c_\mu i^{-2\nu/d}$ for all $i \in \mathbb{N}$, such that $\eta = \nu/d$ in Cases (iii) and (iv) (Santin and Schaback 2016, Theorem 6). In the special case of $\nu = 1/2$ and $d = 1$, $C_{\sigma^2, \phi, 1/2}(\mathbf{s}, \mathbf{s}') = \sigma^2 \exp(-\phi \|\mathbf{s} - \mathbf{s}'\|)$ is the exponential kernel, whose eigenfunctions are bounded sine and cosine functions, so (A.4) is also satisfied with $q = +\infty$ (Van Trees 2001, Section 3.4.1). It is unknown whether the eigenfunctions of Matérn kernels can be uniformly bounded for general ν and d .

When $\eta = \nu/d$ in Cases (iii) and (iv), the rate $O(n^{-\frac{2\nu-d}{2\nu}})$ for the Bayes L_2 -risk in Case (iii) is not minimax optimal for estimating functions in the Sobolev space of order $\nu + d/2$, whereas the faster rate $O(n^{-\frac{2\nu}{2\nu+d}})$ in Case (iv) is minimax optimal. This is because (iv) has used the additional optimal tuning parameter $\lambda_n = c_1 n^{\nu/(2\nu+d)}$, while setting $\lambda_n = 1$ is sub-optimal in this case. The use of a tuning parameter to achieve optimal convergence is common in Gaussian process regression and kernel ridge regression (Zhang et al., 2015, Yang et al., 2017). Although van der Vaart and van Zanten (2011) have shown the minimax optimal posterior convergence rates for the Matérn kernel without using tuning parameters, their proof only works when the true error distribution of $\epsilon(\mathbf{s})$ is sub-Gaussian. In comparison, our Assumption A.1 only requires that $\epsilon(\mathbf{s})$ has a finite variance without the normality assumption, which is more general and allows the model (9) to be misspecified in the error distribution.

For the conditions on k , in the case when $q = +\infty$, the upper bounds on k in (i), (ii), (iii), and (iv) reduce to $k = O\{n/(\log n)^2\}$, $k = O\{n/(\log n)^{2/\kappa}\}$, $k = O\{n^{\frac{\eta-1}{\eta}}/(\log n)^2\}$, and $k = O\{n^{\frac{2\eta-1}{2\eta+1}}/(\log n)^2\}$, respectively. The convergence rate results in Theorem 3.2 hold as long as k does not grow too fast with n .

We can generalize the results in Theorems 3.1 and 3.2 to the model (1). Besides A.1–A.4, we further make the following assumption on $\mathbf{x}(\cdot)$ and the prior on $\boldsymbol{\beta}$:

B.1 All p components of $\mathbf{x}(\cdot)$ are non-random functions in $\mathcal{S}$. The prior on $\boldsymbol{\beta}$ is $N(\boldsymbol{\mu}_\beta, \Sigma_\beta)$ and it is independent of the prior on $w(\cdot)$, which is $\text{GP}\{0, C_\alpha(\cdot, \cdot)\}$.

By the normality and joint independence in Assumption B.1, it is straightforward to show that the mean function $\mathbf{x}(\mathbf{s})^T \boldsymbol{\beta} + w(\mathbf{s})$ has a GP prior $\text{GP}\{\mathbf{x}(\cdot)^T \boldsymbol{\mu}_\beta, \check{C}_\alpha(\cdot, \cdot)\}$, where the modified covariance function $\check{C}_\alpha$ is given by

$$\begin{aligned} \check{C}_\alpha(\mathbf{s}_1, \mathbf{s}_2) &= \text{cov}\{\mathbf{x}(\mathbf{s}_1)^T \boldsymbol{\beta} + w(\mathbf{s}_1), \mathbf{x}(\mathbf{s}_2)^T \boldsymbol{\beta} + w(\mathbf{s}_2)\} \\ (13) \quad &= \mathbf{x}(\mathbf{s}_1)^T \Sigma_\beta \mathbf{x}(\mathbf{s}_2) + C_\alpha(\mathbf{s}_1, \mathbf{s}_2), \text{ for any } \mathbf{s}_1, \mathbf{s}_2 \in \mathcal{S}. \end{aligned}$$

With this modified covariance function, we have the following corollary:

Corollary 3.3 *If Assumption B.1 holds, Assumptions A.1–A.5 hold with all C_α replaced by $\tilde{C}_\alpha$ in (13), and $\mu_\beta = \mathbf{0}$, the conclusions of Theorems 3.1 and 3.2 hold for the Bayes L_2 -risk of the mean surface $\mathbf{x}(\cdot)^T \beta + w(\cdot)$ in the model (1).*

4. EXPERIMENTS

4.1 Simulation Setup

This section presents a comparative study of important non-distributed and distributed approaches on large spatial data based on the performance in learning the process parameters, interpolating the unobserved spatial surface, and predicting the response at new locations. Two simulation studies and a real data analysis are presented. The first simulation (*Simulation 1*) generates the data from a spatial linear model, where the spatial process is simulated from a GP with an exponential covariance function, leading to a fairly rough (nowhere differentiable) spatial surface. Following Gramacy and Apley (2015), we use an analytic function with local features to simulate the data in the second simulation (*Simulation 2*). The number of locations in the two simulations is moderately large with $n = 10,000$. Our real data analysis is based on a large data subset of sea surface temperature data with $n = 1,000,000$ locations. For the two simulations and in the real data analysis, the response at $(n + l)$ locations is modeled as

$$(14) \quad y(\mathbf{s}_i) = \beta_0 + x(\mathbf{s}_i)\beta_1 + w(\mathbf{s}_i) + \epsilon_i, \quad \epsilon_i \sim N(0, \tau^2), \quad \mathbf{s}_i \in \mathcal{D} \subset \mathbb{R}^2,$$

for $i = 1, \dots, n + l$, where $\mathcal{D}$ is the spatial domain, $y(\mathbf{s}_i)$, $x(\mathbf{s}_i)$, $w(\mathbf{s}_i)$, and ϵ_i are the response, covariate, spatial process, and idiosyncratic error values at the location $\mathbf{s}_i$, β_0 is the intercept, β_1 models the covariate effect, and l is the number of new locations where surface interpolation and prediction are sought.

A number of popular and state-of-the-art non-distributed Bayesian and non-Bayesian spatial models are compared with a few important distributed Bayesian approaches in the two simulations and in the real data analysis. Among non-distributed Bayesian and non-Bayesian methods, we fit: (i) Integrated nested Laplace approximation (INLA) using the INLA package in R (Illian et al., 2012); (ii) LatticeKrig (Nychka et al., 2015) using the LatticeKrig package in R with 3 resolutions (Nychka et al., 2016); (iii) modified predictive process (MPP) using the spBayes package in R with the full data; (iv) nearest neighbor Gaussian process (NNGP) using the spNNGP package in R with the number of nearest neighbors m set to be 10, 20, and 30 (Datta et al., 2016); (v) locally approximated Gaussian process (laGP) using the laGP package in R (Gramacy and Apley, 2015); (vi) Vecchia’s approximation using the GPvecchia package in R with the number of nearest neighbors m set to be 10, 20, and 30 (Katzfuss and Guinness, 2021); (vii) Fisher Scoring of Vecchia’s Approximation using the GpGp (Guinness, 2019).

In fitting (i), (ii), (iv), (v), (vi), (vii), we assume an exponential correlation in the random field given by $\text{cov}\{w(\mathbf{s}), w(\mathbf{s}')\} = \sigma^2 e^{-\phi \|\mathbf{s} - \mathbf{s}'\|}$, $\mathbf{s}, \mathbf{s}' \in \mathcal{D}$. To fit MPP for (iii), the MPP prior on $w(\cdot)$ is fitted with rank $r = 200, 400$ in Simulations 1, 2 and with $r = 400, 600$ in the real data analysis, where r knots are selected randomly from $\mathcal{D}$. For Bayesian model fitting, we apply a flat prior on (β_0, β_1) , a $\text{IG}(2, 0.1)$ prior on τ^2 , an $\text{IG}(2, 2)$ prior on σ^2 and a uniform prior on ϕ , where $\text{IG}(a, b)$ is the Inverse-Gamma distribution with mean $b/(a - 1)$.

The non-distributed approaches are compared with distributed Bayesian methods for model-free subset posterior aggregation discussed in Section 3 of this article. They are (viii) CMC (Scott et al. (2016)); (ix) DPMC (Xue and Liang (2019)); (x) WASP (Srivastava et al. (2015)); (xi) DISK (with $\xi = 10^{-4}$), for our exposition. Identical priors, covariance functions, ranks, and knots are used for the non-distributed process models and their distributed counterparts for a fair comparison. We emphasize that the distributed methods do not compete with the non-distributed methods in (i)-(vii). Instead, each of them can be potentially embedded in the second step of any of the distributed methods for improved performance because the distributed approaches are not model-specific. More importantly, MPP is not considered to be the state-of-the-art, so it is instructive to investigate the competitiveness of (viii)-(xi) with MPP fitted on each subset.

In the interest of space, we present the performance comparison between distributed and non-distributed approaches only, and similar comparisons between CMC, DISK, DPMC and WASP are presented in the supplementary material. Because DISK shows better or similar performance as its distributed competitors in all simulations, we only present results from DISK with the non-distributed methods in the main text. Notably, DISK combines one-dimensional marginals of subset posteriors, but DPMC and WASP aggregate subset posteriors of multivariate parameters; therefore, similar performances of DISK, DPMC, and WASP in the supplementary material shows that combining subset posteriors of univariate parameters does not lead to any significant loss in inference or predictions.

Any distributed method has two important choices: (A) the value of k and (B) the construction of subsets. We choose k in our experiments based on two broad guidelines: (a) available computational resources and (b) the subset size to draw reliable inference on the spatial surface with data subsets. To assess (b), we plot the histograms or density estimates of subset posterior draws of representative parameters and see if they are very far from each other. If so, the subset posteriors fail to provide a noisy approximation of the full data posterior, resulting in inaccuracy of the combined pseudo posterior for a distributed approach. Empirically, we also propose computing the pairwise Wasserstein or total variation distance between the subset posteriors of representative parameters. If the average of these distances is much larger than the average distance between the combined and subset posterior distributions, then the combined pseudo posterior provides a poor approximation of the full data posterior. Assuming that the fitted model can reasonably capture variation of the data, these checks would imply that one has to fit a distributed approach with a smaller value of k .

Regarding (B), we present performance of the distributed approaches when data subsets are constructed (a) under a random partitioning scheme and (b) under a grid partitioning scheme. Random partitioning scheme randomly partitions the data into subsets. In contrast, grid partitioning scheme partitions the domain into a number of sub-domains and creates each subset with representative samples from each sub-domain. All tables in the main article and in supplementary material show results from both partitioning schemes.

We run all the experiments on an Oracle Grid Engine cluster with 2.6GHz 16 core compute nodes. The non-distributed methods (INLA, LatticeKrig, MPP, NNGP, laGP, GPvecchia, and GpGp) and the distributed methods (DISK, DPMC, CMC, and WASP) are allotted memory resources of 64GB and 16GB, respec-

tively. Every MCMC algorithm runs for 10,000 iterations, out of which the first 5,000 MCMC samples are discarded as burn-ins and the rest of the chain is thinned by collecting every fifth MCMC sample. We also refer to Section 5 of the supplementary material that presents comparison between effective sample size of model parameters averaged over all subsets to the effective sample size of model parameters from the full data posterior in simulations. We compare the quality of prediction and estimation of spatial surface at predictive locations $\mathcal{S}^* = \{\mathbf{s}_1^*, \dots, \mathbf{s}_l^*\}$. If $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$ are the value of the spatial surface and response at $\mathbf{s}_{i'}^* \in \mathcal{S}^*$, then the estimation and prediction errors are defined as

$$(15) \quad \text{Est Err}^2 = \frac{1}{l} \sum_{i'=1}^l \{\hat{w}(\mathbf{s}_{i'}^*) - w(\mathbf{s}_{i'}^*)\}^2, \quad \text{Pred Err}^2 = \frac{1}{l} \sum_{i'=1}^l \{\hat{y}(\mathbf{s}_{i'}^*) - y(\mathbf{s}_{i'}^*)\}^2,$$

where $\hat{w}(\mathbf{s}_{i'}^*)$ and $\hat{y}(\mathbf{s}_{i'}^*)$ denote the point estimates of $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$ obtained using any distributed or non-distributed methods. For sampling-based methods, we set $\hat{w}(\mathbf{s}_{i'}^*)$ and $\hat{y}(\mathbf{s}_{i'}^*)$ to be the medians of posterior MCMC samples for $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$, respectively, for $i' = 1, \dots, l$. We also estimate the point-wise 95% credible or confidence intervals (CIs) of $w(\mathbf{s}_{i'}^*)$ and predictive intervals (PIs) of $y(\mathbf{s}_{i'}^*)$ for every $\mathbf{s}_{i'}^* \in \mathcal{S}^*$ and compare the CI and PI coverages and lengths for every method. Finally, we compare the performance of all the methods for parameter estimation using the posterior medians and the 95% CIs. Posterior medians are reported instead of posterior means as point estimators since they are easily estimated for the DISK combined posterior following equation (8).

4.2 Simulation 1: Spatial Linear Model Based On GP

TABLE 1

The errors in estimating the parameters $\beta = (\beta_0, \beta_1), \sigma^2, \phi, \tau^2$ in Simulation 1. The parameter estimates for the Bayesian methods $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1)$, $\hat{\sigma}^2, \hat{\phi}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples and their true values are $\beta_0 = (1, 2)$, $\sigma_0^2 = 1$, $\phi_0 = 4$ and $\tau_0^2 = 0.1$. The entries in the table are averaged across 10 simulation replications.

	$\ \hat{\beta} - \beta_0\ $	$ \hat{\sigma}^2 - \sigma_0^2 $	$ \hat{\phi} - \phi_0 $	$ \hat{\tau}^2 - \tau_0^2 $
INLA	0.21	-	-	-
LaGP	0.08	-	-	-
NNGP ($m = 10$)	0.11	0.07	0.37	0.00
NNGP ($m = 20$)	0.12	0.09	0.51	0.00
NNGP ($m = 30$)	0.11	0.11	0.58	0.00
LatticeKrig	0.11	0.09	1.59	0.06
GpGp	0.08	0.11	0.64	0.01
Vecchia ($m = 10$)	0.10	0.11	0.51	0.01
Vecchia ($m = 20$)	0.10	0.10	0.55	0.01
Vecchia ($m = 30$)	0.10	0.38	1.13	0.01
MPP ($r = 200$)	0.35	0.23	1.98	0.17
MPP ($r = 400$)	0.19	0.09	1.88	0.07
Random Partitioning				
DISK ($r = 200, k = 10$)	0.09	0.11	0.64	0.01
DISK ($r = 400, k = 10$)	0.09	0.11	0.64	0.01
DISK ($r = 200, k = 20$)	0.10	0.12	0.66	0.02
DISK ($r = 400, k = 20$)	0.10	0.12	0.66	0.02
Grid-Based Partitioning				
DISK ($r = 200, k = 10$)	0.09	0.12	0.62	0.01
DISK ($r = 400, k = 10$)	0.09	0.12	0.62	0.01
DISK ($r = 200, k = 20$)	0.10	0.12	0.63	0.01
DISK ($r = 400, k = 20$)	0.10	0.12	0.64	0.01

Our first simulation generates data using the spatial linear model in (14). We set $\mathcal{D} = [-2, 2] \times [-2, 2] \subset \mathbb{R}^2$, $n = 10,000$, $l = 500$ and uniformly draw

TABLE 2

The estimates of parameters $\beta = (\beta_0, \beta_1)$, σ^2 , ϕ , τ^2 and their 95% marginal credible intervals (CIs) in Simulation 1. The parameter estimates for the Bayesian methods $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1)$, $\hat{\sigma}^2$, $\hat{\phi}$, $\hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications; '-' indicates that the uncertainty estimates are not provided by the software or the competitor.

	β_0	β_1	σ^2	ϕ	τ^2
Truth	1.00	2.00	1.00	4.00	0.10
Parameter Estimates					
INLA	1.00	2.00	-	-	-
laGP	1.01	2.00	-	-	-
NNGP ($m = 10$)	1.02	2.00	0.99	4.00	0.10
NNGP ($m = 20$)	0.98	2.00	0.94	4.30	0.10
NNGP ($m = 30$)	0.99	2.00	0.94	4.34	0.10
LatticeKrig	1.01	2.00	0.93	2.42	0.16
GpGp	0.99	2.00	0.92	4.43	0.11
Vecchia ($m = 10$)	0.99	2.00	0.94	3.93	0.09
Vecchia ($m = 20$)	0.99	2.00	0.95	3.93	0.09
Vecchia ($m = 30$)	1.00	2.00	1.10	3.68	0.09
MPP ($r = 200$)	1.26	2.00	0.77	2.02	0.27
MPP ($r = 400$)	1.08	2.00	0.99	2.14	0.17
DISK ($r = 200, k = 10$)	1.00	2.00	0.92	4.35	0.11
DISK ($r = 400, k = 10$)	1.00	2.00	0.92	4.35	0.11
DISK ($r = 200, k = 20$)	1.00	2.00	0.91	4.38	0.11
DISK ($r = 400, k = 20$)	1.00	2.00	0.91	4.38	0.11
95% Credible Intervals					
INLA	(0.26, 1.73)	(1.98, 2.02)	-	-	-
laGP	(0.99, 1.03)	(1.98, 2.02)	-	-	-
NNGP ($m = 10$)	(0.87, 1.15)	(1.99, 2.01)	(0.86, 1.24)	(3.15, 4.70)	(0.09, 0.11)
NNGP ($m = 20$)	(0.85, 1.13)	(1.99, 2.01)	(0.82, 1.14)	(3.46, 4.95)	(0.09, 0.11)
NNGP ($m = 30$)	(0.86, 1.12)	(1.99, 2.01)	(0.81, 1.11)	(3.62, 5.03)	(0.09, 0.11)
LatticeKrig	-	-	-	-	-
GpGp	(0.75, 1.23)	(1.99, 2.01)	-	-	-
Vecchia ($m = 10$)	-	-	-	-	-
Vecchia ($m = 20$)	-	-	-	-	-
Vecchia ($m = 30$)	-	-	-	-	-
MPP ($r = 200$)	(1.06, 1.26)	(1.98, 2.00)	(0.70, 0.85)	(2.01, 2.07)	(0.24, 0.30)
MPP ($r = 400$)	(0.76, 1.08)	(1.99, 2.00)	(0.91, 1.08)	(2.07, 2.26)	(0.15, 0.19)
DISK ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
DISK ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
DISK ($r = 200, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.07, 4.67)	(0.09, 0.13)
DISK ($r = 400, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.07, 4.68)	(0.09, 0.13)

$(n + l)$ spatial locations $\mathbf{s}_i = (s_{i1}, s_{i2})$ in $\mathcal{D}$ ($i = 1, \dots, n + l$). The spatial surface $w(\cdot)$ at the $(n + l)$ locations, $\{w(\mathbf{s}_1), \dots, w(\mathbf{s}_{n+l})\}$, is simulated from $\text{GP}(0, \sigma^2 \exp\{-\phi \|\mathbf{s} - \mathbf{s}'\|\})$, where $\mathbf{s}, \mathbf{s}' \in \mathcal{D}$, $\phi = 4$, and $\sigma^2 = 1$. The covariance function ensures the generated spatial surface is continuous everywhere but differentiable nowhere, which is a more familiar simulation scenario in the spatial context. Setting $\beta_0 = 1$, $\beta_1 = 2$, and $\tau^2 = 0.1$, we simulate the responses at $(n + l)$ locations using (14). The three-step distributed frameworks are applied using the low-rank MPP priors with $k = 10$ and $k = 20$, having average subset sizes 1000 and 500, respectively. We replicate this simulation ten times.

DISK with MPP prior, NNGP, and GPvecchia have similar performance in parameter estimation (Tables 1 and 2). The parameter estimates obtained using DISK are very close to their true values and the estimation errors are very similar to those of NNGP and non-Bayesian methods based on the Vecchia approximation, including GpGp and GPvecchia. The 95% credible intervals of β_0, β_1, τ^2 in DISK cover the true values and their lower and upper quantiles are very similar

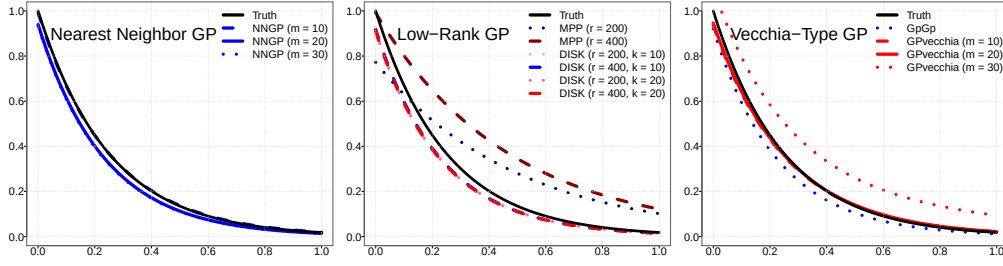


Fig 1: Estimated covariance function using three types of GP priors on the spatial surface. The true covariance function is $\text{cov}\{w(\mathbf{s}_i), w(\mathbf{s}_j)\} = \exp(-4\|\mathbf{s}_i - \mathbf{s}_j\|_2)$.

TABLE 3

Inference on the values of spatial surface and response at the locations in $\mathcal{S}_$ in Simulation 1. The estimation and prediction errors are defined in (15) and coverage and credible intervals are calculated pointwise for the locations in $\mathcal{S}_*$. The entries in the table are averaged over 10 simulation replications; ‘-’ indicates that the estimates are not provided by the software or the competitor.*

	Est Err	Pred Err	95% CI Coverage		95% CI Length	
	GP	Y	GP	Y	GP	Y
INLA	-	0.90	-	0.80	-	0.17
laGP	0.20	0.28	0.98	0.95	2.06	1.04
NNGP ($m = 10$)	0.38	0.47	0.93	0.95	1.39	1.84
NNGP ($m = 20$)	0.38	0.47	0.93	0.95	1.38	1.81
NNGP ($m = 30$)	0.38	0.47	0.92	0.95	1.37	1.82
LatticeKrig	0.38	0.47	-	0.73	-	1.08
GpGp	-	0.47	-	-	-	-
Vecchia ($m = 10$)	-	0.47	-	0.87	-	1.43
Vecchia ($m = 20$)	-	0.47	-	0.86	-	1.41
Vecchia ($m = 30$)	-	0.47	-	0.86	-	1.41
MPP ($r = 200$)	0.73	0.59	0.93	0.95	3.05	3.02
MPP ($r = 400$)	0.43	0.47	0.96	0.95	2.76	2.67
Random Partitioning						
DISK ($r = 200, k = 10$)	0.55	0.64	0.97	0.97	3.20	3.45
DISK ($r = 400, k = 10$)	0.42	0.51	0.97	0.97	2.88	3.15
DISK ($r = 200, k = 20$)	0.58	0.67	0.97	0.97	3.25	3.51
DISK ($r = 400, k = 20$)	0.46	0.55	0.97	0.97	2.98	3.25
Grid-Based Partitioning						
DISK ($r = 200, k = 10$)	0.75	0.80	0.97	0.97	3.45	3.45
DISK ($r = 400, k = 10$)	0.65	0.72	0.97	0.97	3.15	3.15
DISK ($r = 200, k = 20$)	0.76	0.82	0.97	0.97	3.51	3.51
DISK ($r = 400, k = 20$)	0.68	0.74	0.97	0.97	3.26	3.26

to those of NNGP. DISK underestimates σ^2 and overestimates ϕ slightly. Both results are the impacts of parent MPP prior, which also shows less accurate estimation of the posterior distribution of σ^2 and ϕ for the two choices of the number of knots r . More importantly, the impacts the choice of r on parameter estimation are less severe in the distributed methods compared to that in its parent MPP prior. The CIs are not available from GPvecchia, LatticeKrig and laGP, so that the cells corresponding these methods are kept blank in Table 2.

Despite the discrepancy in parameter estimates, the correlation function estimates obtained using the combined posteriors from distributed competitors (DISK pseudo posterior being a representative) are very close to those obtained using NNGP and GPvecchia (Figure 1). Similar to the observations of Sang and Huang (2012), there is considerable discrepancy between the estimated and true correlation functions when the MPP prior is used. In contrast, for the same

choices of r as its parent MPP prior, DISK’s estimate of the correlation function is much closer to the truth and is insensitive to the choice of $k = 10, 20$. DISK estimates are similar to those obtained using Vecchia-type approximation, except when the number of nearest neighbor is 30 and the GPvecchia-based estimate of the correlation function has a significant positive bias.

The predictive performance of the representative distributed competitor DISK is little inferior to that of NNGP. NNGP, MPP, and DISK have close to nominal predictive coverage, but the PIs of NNGP have smaller lengths for every choice of nearest neighbor. The PI coverage values and lengths of MPP and DISK are similar and stable for the different choices of r and k . PIs in GPvecchia have the smallest length and their coverage values are smaller than the nominal value for all the three choices of nearest neighbor. Focusing on spatial surface interpolation, the estimation error of DISK is smaller than that of MPP for both choices of r when $k = 10$ and is slightly larger when $k = 20$ and $r = 400$. Similarly, MPP’s coverage of the spatial surface is smaller than the nominal value when $r = 200$, but DISK shows better coverage than its parent MPP prior for both choices of k . Consequently, the lengths of DISK’s credible intervals are slightly larger than those obtained using its parent MPP prior.

In summary, the distributed methods are competitive with state-of-the-art non-distributed methods NNGP and GPvecchia in inference on the spatial surface and predictions, respectively. laGP is the only non-distributed competing method that yields comprehensively better inferential and predictive performance than all distributed methods, but it is not designed to provide estimates for the σ^2 , ϕ , and τ^2 . LatticeKrig has a very similar point estimation, but inferior uncertainty quantification compared to GpGp and GPvecchia. INLA underperforms in surface interpolation and prediction. Supplementary material shows comparative performance of distributed competitors and also ensures that stochastic approximation does not impact the mixing of the Markov chains on the subsets. The model free nature of the distributed methods also allows us to fit a nearest neighbor approach, including NNGP, on each subset to improve inference and expedite computations by multiple folds. Finally, the results show that the random partitioning scheme yields little better point estimation with similar uncertainty quantification compared to a more sophisticated grid partitioning scheme.

4.3 Simulation 2: Spatial Linear Model Based On Analytic Spatial Surface

Our second simulation generates data by setting $w(\cdot)$ in (14) to be an analytic function. For any $s \in [-2, 2]$, define the function $f_0(s) = e^{-(s-1)^2} + e^{-0.8(s+1)^2} - 0.05 \sin\{8(s + 0.1)\}$ and set $w(\mathbf{s}_i) = -f_0(s_{i1})f_0(s_{i2})$. Although the function $w(\cdot)$ simulated in this way is theoretically infinitely smooth, the response surface simulated from (14) exhibits complex local behavior, which is challenging to capture using spatial process-based models as we demonstrate later. We set $\beta_0 = 1$, $\beta_1 = 0$, and $\tau^2 = 0.01$, use the same values of the spatial domain, k , and r as used in the previous simulation, and replicate this simulation 10 times.

The parameter estimation results in this simulation are similar to those in Simulation 1 with one important exception in inference on β_0 (Tables 4 and 5). All the methods except GpGp show excellent performance in estimating τ^2 ; however, NNGP, GPvecchia, and MPP estimate β_0 with a significant bias. 95% credible intervals of β_0 computed from DISK has better coverage properties than

TABLE 4

The errors in estimating the parameters β, τ^2 in Simulation 2. The parameter estimates for the Bayesian methods $\hat{\beta}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples and $\beta_0 = 1$ and $\tau_0^2 = 0.01$. The entries in the table are averaged across 10 simulation replications.

	$\ \hat{\beta} - \beta_0\ $	$ \hat{\tau}^2 - \tau_0^2 $
INLA	0.18	-
LaGP	-	-
NNGP ($m = 10$)	0.84	0.03
NNGP ($m = 20$)	0.84	0.03
NNGP ($m = 30$)	0.84	0.03
LatticeKrig	-	0.01
GpGp	0.31	0.39
Vecchia ($m = 10$)	0.85	0.01
Vecchia ($m = 20$)	0.85	0.01
Vecchia ($m = 30$)	0.85	0.01
MPP ($r = 200$)	0.75	0.05
MPP ($r = 400$)	0.48	0.04
Random Partitioning		
DISK ($r = 200, k = 10$)	0.18	0.04
DISK ($r = 400, k = 10$)	0.13	0.04
DISK ($r = 200, k = 20$)	0.18	0.04
DISK ($r = 400, k = 20$)	0.13	0.04
Grid-Based Partitioning		
DISK ($r = 200, k = 10$)	0.03	0.09
DISK ($r = 400, k = 10$)	0.03	0.09
DISK ($r = 200, k = 20$)	0.02	0.09
DISK ($r = 400, k = 20$)	0.02	0.09

those of NNGP. Unlike our observation in the previous section, all the methods underestimate τ^2 slightly, and the 95% credible intervals of NNGP, MPP prior, and DISK fail to cover the true value. Similar to the previous simulation results, DISK performs better than its parent MPP prior for both choices of r .

The predictive and inferential performance of distributed methods in this simulation are also very similar to those in Simulation 1. The prediction error, PI coverage, and PI length of all the methods except GPvecchia are fairly similar and are close to the nominal value. The PI length of GPvecchia is the smallest, but its coverage values are critically low for all choices of nearest neighbor; that is, GPvecchia has a relatively inferior performance for estimating spatial surfaces that are not simulated from a GP. The PI coverage values of distributed method DISK is a little higher than those of NNGP and MPP priors while the PI lengths of DISK are very close to those of MPP and NNGP priors. A noticeable feature of our comparison is that the distributed methods improve the performance of their parent MPP prior when $r = 200$. In this case, the CI coverage values of distributed methods for both choices of k are greater the nominal value, whereas the parent MPP prior has fails to cover the spatial surface. Intuitively, for most competitors in this simulation the estimation of fixed and random effects are mixed up, whereas the overall mean effect is estimated correctly by all competitors.

As in Simulation 1, INLA still underperforms in surface interpolation and prediction, and laGP maintains its superior predictive and inferential performance, especially because it is tuned for inference in such analytic surfaces with many local features (Gramacy and Apley, 2015). LatticeKrig also offers excellent performance and it outperforms the distributed methods in terms of surface estimation and prediction. Simulation 2 shows that grid based partitioning yields better point estimation for β_0 , but inferior point estimation for τ_0^2 (Table 4). This leads

TABLE 5

The estimates of parameters $\beta, \sigma^2, \phi, \tau^2$ and their 95% marginal credible intervals (CIs) in Simulation 2. The parameter estimates for the Bayesian methods $\hat{\beta}, \hat{\sigma}^2, \hat{\phi}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications; '-' indicates that the uncertainty estimates are not provided by the software or the competitor.

	β	σ^2	ϕ	τ^2
Truth	1.00	-	-	0.01
Parameter Estimates				
INLA	0.8161	-	-	-
laGP	-	-	-	-
NNGP ($m = 10$)	0.2897	0.1933	0.1075	0.0091
NNGP ($m = 20$)	0.3002	0.1660	0.1059	0.0092
NNGP ($m = 30$)	0.2892	0.1557	0.1058	0.0093
LatticeKrig	-	-	0.0842	0.0099
GpGp	1.0346	0.0669	0.2643	0.1620
Vecchia ($m = 10$)	0.2792	0.4063	0.7796	0.0099
Vecchia ($m = 20$)	0.2792	0.2904	0.9479	0.0099
Vecchia ($m = 30$)	0.2792	0.2746	0.9587	0.0099
MPP ($r = 200$)	1.5634	0.1535	0.1185	0.0077
MPP ($r = 400$)	1.2333	0.1586	0.1200	0.0080
DISK ($r = 200, k = 10$)	1.0322	0.2133	0.1196	0.0087
DISK ($r = 400, k = 10$)	0.9830	0.2185	0.1402	0.0082
DISK ($r = 200, k = 20$)	1.0328	0.2133	0.1194	0.0087
DISK ($r = 400, k = 20$)	0.9822	0.2185	0.1402	0.0082
95% Credible Intervals				
INLA	(0.5320, 1.2108)	-	-	-
laGP	-	-	-	-
NNGP ($m = 10$)	(0.2678, 0.3143)	(0.1568, 0.2223)	(0.1010, 0.1339)	(0.0088, 0.0094)
NNGP ($m = 20$)	(0.2801, 0.3226)	(0.1361, 0.1906)	(0.1009, 0.1279)	(0.0089, 0.0095)
NNGP ($m = 30$)	(0.2660, 0.3103)	(0.1293, 0.1794)	(0.1009, 0.1284)	(0.0090, 0.0095)
LatticeKrig	-	-	-	-
GpGp	(0.7090, 1.3601)	-	-	-
Vecchia ($m = 10$)	-	-	-	-
Vecchia ($m = 20$)	-	-	-	-
Vecchia ($m = 30$)	-	-	-	-
MPP ($r = 200$)	(0.9931, 2.1464)	(0.1307, 0.1760)	(0.1104, 0.1327)	(0.0073, 0.0081)
MPP ($r = 400$)	(0.6130, 1.8412)	(0.1269, 0.1876)	(0.1096, 0.1480)	(0.0076, 0.0084)
DISK ($r = 200, k = 10$)	(0.7961, 1.2722)	(0.1783, 0.2418)	(0.1088, 0.1439)	(0.0084, 0.0091)
DISK ($r = 400, k = 10$)	(0.8180, 1.1582)	(0.1743, 0.2589)	(0.1192, 0.1773)	(0.0079, 0.0086)
DISK ($r = 200, k = 20$)	(0.7987, 1.2719)	(0.1781, 0.2417)	(0.1087, 0.1434)	(0.0084, 0.0091)
DISK ($r = 400, k = 20$)	(0.8172, 1.1568)	(0.1721, 0.2588)	(0.1190, 0.1806)	(0.0079, 0.0086)

to little better surface estimation for random partitioning scheme than grid-based partitioning scheme, but practically indistinguishable predictive performance as demonstrated in Table 6. We conclude that the distributed methods are promising tools even when the spatial surface is not simulated from a GP.

4.4 Real Data Analysis: Sea Surface Temperature Data

A description of the evolution and dynamics of the SST is a key component of the study of the Earth's climate. SST data (in centigrade) from ocean samples have been collected by voluntary observing ships, buoys, and military and scientific cruises for decades. During the last 20 years or so, the SST database has been complemented by regular streams of remotely sensed observations from satellite orbiting the earth. A careful quantification of variability of SST data is important for climatological research, which includes determining the formation of sea breezes and sea fog and calibrating measurements from weather satellites (Di Lorenzo et al., 2008). A number of articles have appeared to address this issue in recent years; see Berliner et al. (2000), Lemos and Sansó (2009), Wikle

TABLE 6

Inference on the values of spatial surface and response at the locations in S_ in Simulation 2. The estimation and prediction errors are defined in (15) and coverage and credible intervals are calculated pointwise for the locations in S_* . The entries in the table are averaged over 10 simulation replications; ‘-’ indicates that the estimates are not provided by the software or the competitor.*

	Est Err	Pred Err	95% CI Coverage		95% CI Length	
	GP	Y	GP	Y	GP	Y
INLA	-	0.1552	-	0.0755	-	0.0268
laGP	0.0004	0.0103	1.0000	0.9456	0.3890	0.3902
NNGP ($m = 10$)	0.5058	0.0104	0.0000	0.9439	0.1496	0.3949
NNGP ($m = 20$)	0.4908	0.0103	0.0000	0.9456	0.1392	0.3938
NNGP ($m = 30$)	0.5103	0.0103	0.0000	0.9479	0.1388	0.3969
LatticeKrig	0.0002	0.0101	0.9867	0.9463	-	0.3901
GpGp	-	0.0103	-	-	-	-
Vecchia ($m = 10$)	-	0.0106	-	0.3559	-	0.0951
Vecchia ($m = 20$)	-	0.0103	-	0.2815	-	0.0728
Vecchia ($m = 30$)	-	0.0102	-	0.2612	-	0.0674
MPP ($r = 200$)	0.3732	0.0105	0.0000	0.9498	0.4061	0.4061
MPP ($r = 400$)	0.0623	0.0102	0.2946	0.9477	0.3976	0.3976
Random Partitioning						
DISK ($r = 200, k = 10$)	0.0017	0.1035	1.0000	0.9696	0.5388	0.4449
DISK ($r = 400, k = 10$)	0.0009	0.1026	1.0000	0.9724	0.4477	0.4578
DISK ($r = 200, k = 20$)	0.0015	0.1041	1.0000	0.9646	0.5211	0.4248
DISK ($r = 400, k = 20$)	0.0007	0.1031	1.0000	0.9672	0.4253	0.4359
Grid-Based Partitioning						
DISK ($r = 200, k = 10$)	0.0394	0.1036	1.0000	0.9660	0.4452	0.4452
DISK ($r = 400, k = 10$)	0.0368	0.1028	1.0000	0.9594	0.4249	0.4249
DISK ($r = 200, k = 20$)	0.0304	0.1040	1.0000	0.9700	0.4590	0.4590
DISK ($r = 400, k = 20$)	0.0268	0.1030	1.0000	0.9642	0.4371	0.4371

and Holan (2011), Hazra and Huser (2021).

We consider the problem of capturing the spatial trend and characterizing the uncertainties in the SST in the west coast of mainland U.S.A., Canada, and Alaska between 40° – 65° north latitudes and 100° – 180° west longitudes. The data is obtained from NODC World Ocean Database (https://www.nodc.noaa.gov/OC5/WOD/pr_wod.html) and the entire data corresponds to sea surface temperature measured by remote sensing satellites on 16th August 2016. All data locations are distinct and there is no time replicate; therefore, we can practically ignore the temporal variation of sea surface temperature for our analysis. After screening the data for quality control, we choose a random subset of 1,000,800 spatial observations over the selected domain. From these observations, we randomly select 10^6 observations as training data and the remaining observations are used to compare the performance of distributed and non-distributed competitors. We replicate this setup ten times. The selected domain is large enough to allow considerable spatial variation in SST from north to south and provides an important first step in extending these models for analyzing global-scale SST database.

The SST data in the selected domain shows a clear decreasing trend in SST with increasing latitude (Figure 2). Based on this observation, we add latitude as a linear predictor in the univariate spatial regression model (14) to explain the long-range directional variability in the SST. Similar to Simulation 1 and 2, Section 4.4 in the supplementary material shows that among distributed competitors DISK shows identical or little better performance than the other distributed approaches for the sea surface data. Thus, we only present results from DISK in this section due to space constraint considering it as a representative

distributed competitor. The detailed performance comparison of all distributed competitors in the real data can be found in Section 4.4 of the supplementary material. To fit distributed competitors, we set $k = 300$, which results in subsets of approximately 3300 locations. Since each subset has larger sample size than the simulation studies, we increase the number of knots in each subset for model fitting and use MPP priors with 400 and 600 knots, respectively, in each subset. All the non-distributed competitors except laGP fail to produce results due to numerical issues. Specifically, GPvecchia and GpGp fail after 8 and 21 iterations with an error in `vecchia.Linv` function, INLA fails with an error in `GMRFLib_factorise_sparse_matrix.TAUCS` function, spNNGP fails an error in the `dpotrf` function, and MPP fails from memory bottlenecks. Due to the lack of ground truth for estimating $w(\mathbf{s}^*)$, we compare the DISK and laGP in terms of their inference on $\mathbf{\Omega}$ and prediction of $y(\mathbf{s}^*)$ for $\mathbf{s}^* \in \mathcal{S}^*$ in terms of MSPE and the length and coverage of 95% posterior PIs.

DISK provides inference on the covariance function, including credible intervals for σ^2 , ϕ , and τ^2 , which are unavailable in laGP. The 50%, 2.5%, and 97.5% quantiles of the posterior distributions for $\mathbf{\Omega}$, $w(\mathbf{s}^*)$ and $y(\mathbf{s}^*)$ for every $\mathbf{s}^* \in \mathcal{S}^*$ are used for estimation and uncertainty quantification. We observe significantly higher estimation of spatial variability than non-spatial variability from DISK indicating local spatial variation in SST. Importantly, the point estimate of β_1 is negative and its 95% CI does not include zero, which confirms that SST decreases as latitude increases. For every $\mathbf{s}^* \in \mathcal{S}^*$, laGP's and DISK's estimates of $w(\mathbf{s}^*)$ and $y(\mathbf{s}^*)$ agree closely (Figures 2 and 3 and Table 7). The pointwise predictive coverages of laGP and DISK match their nominal levels; however, the 95% posterior PIs of DISK are wider than those of laGP because DISK accounts for uncertainty due to the error term and unknown parameters (Figure 2 and Table 7). As a whole, SST data analysis reinforces our findings on the importance of distributed Bayesian methods as computationally efficient and flexible tools for full Bayesian inference.

TABLE 7

Parametric inference and prediction in SST data. DISK uses MPP-based modeling with $r = 400, 600$ on $k = 300$ subsets. For parametric inference posterior medians are provided along with The 95% credible intervals (CIs) in the parentheses, where available. Similarly mean squared prediction errors (MSPEs) along with length and coverage of 95% predictive intervals (PIs) are presented, where available. The upper and lower quantiles of 95% CIs and PIs are averaged over 10 simulation replications; '-' indicates that the parameter estimate or prediction is not provided by the software or the competitor

	β_0	β_1	σ^2	ϕ	τ^2
Parameter Estimate					
laGP	32.98	-0.37	-	-	-
DISK ($r = 400$)	32.33	-0.32	11.82	0.04	0.18
DISK ($r = 600$)	32.33	-0.32	11.85	0.04	0.18
95% Credible Interval					
laGP	-	-	-	-	-
DISK ($r = 400$)	(31.72, 32.93)	(-0.33, -0.31)	(11.24, 12.43)	(0.0373, 0.0412)	(0.18, 0.19)
DISK ($r = 600$)	(31.72, 32.93)	(-0.33, -0.31)	(11.25, 12.45)	(0.0372, 0.0413)	(0.18, 0.19)
Predictions					
	MSPE	95% PI Coverage	95% PI Length		
laGP	0.24	0.95	1.35		
DISK ($r = 400$)	0.43	0.95	2.65		
DISK ($r = 600$)	0.36	0.95	2.34		

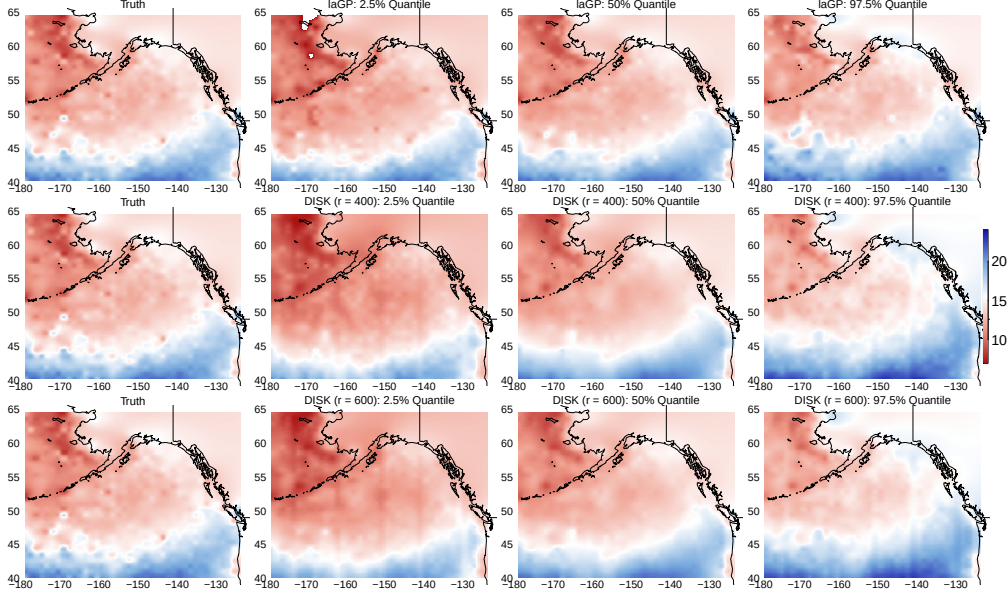


Fig 2: Prediction of sea surface temperatures at the locations in $\mathcal{S}^*$. Negative longitude means degree west from Greenwich. DISK uses MPP-based modeling with $r = 400, 600$ on $k = 300$ subsets and laGP uses the ‘nn’ method. The 2.5%, 50%, and 97.5% quantile surfaces, respectively, represent pointwise quantiles of the posterior distribution for $y(\mathbf{s}^*)$ for every $\mathbf{s}^* \in \mathcal{S}^*$.

5. DISCUSSION

This article presents a comparative study of a class of distributed Bayesian and popular non-distributed methods that are tuned for spatial GP regression in massive data settings. As part of our exposition, we have demonstrated through simulated and real data analyses that distributed Bayesian methods compare well with state-of-the-art non-distributed methods. Motivated by the promising empirical performance, we provide theoretical support for our numerical results. In particular, under commonly assumed regularity conditions, we have provided explicit upper bound on the number of subsets k depending on the analytic properties of the spatial surface so that the Bayes L_2 -risk of the combined pseudo posterior for a subclass of distributed methods is nearly minimax optimal. Additional empirical and theoretical results in the supplementary material shed light on the relative empirical performances of different distributed Bayesian methods in simulations and in the real data analyses.

The simplicity and generality of distributed frameworks enable scaling of any spatial model. For example, recent applications have confirmed that the NNGP prior requires modifications if scalability is desired for even a few millions of locations (Finley et al., 2019b). While computing subset posteriors with MCMC algorithm, we have tacitly assumed that the MCMC chain converges to the subset posterior. While there is no theoretical result to support this, we find enough empirical evidence regarding convergence of the MCMC chain for each subset posterior. We plan to explore this issue theoretically in a future work. In future, we also aim to scale ordinary NNGP and other multiscale approaches to tens of

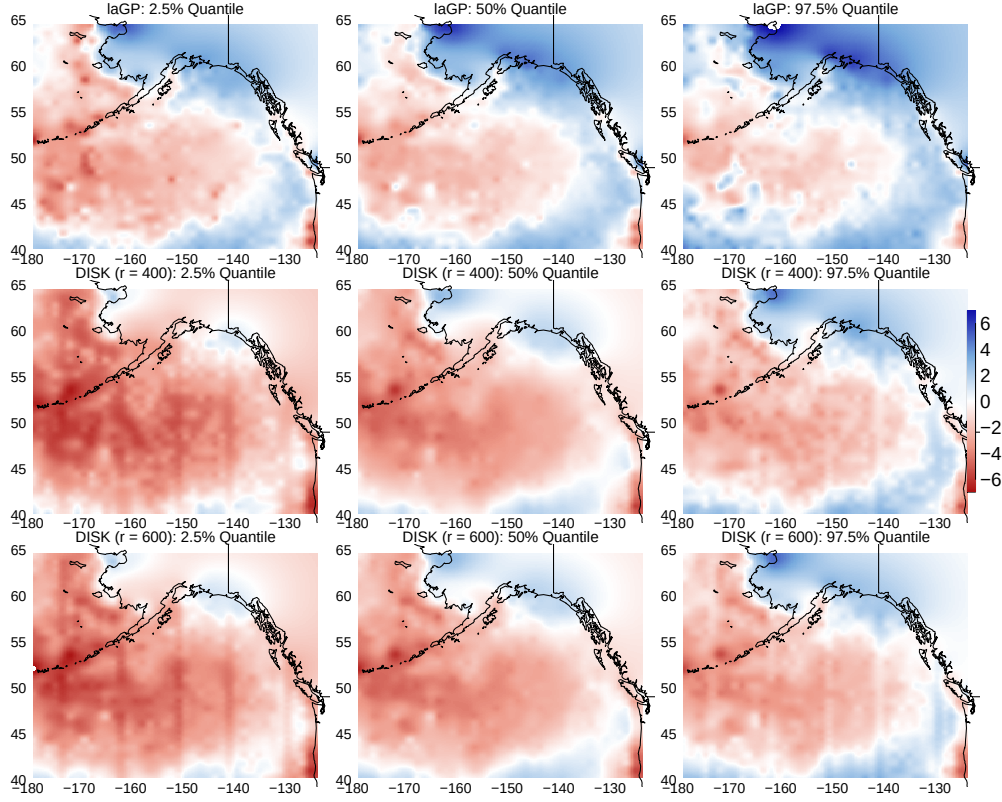


Fig 3: Interpolated spatial surface w at the locations in $\mathcal{S}^*$. Negative longitude means degree west from Greenwich. DISK uses MPP-based modeling with $r = 400, 600$ on $k = 300$ subsets and laGP uses the ‘nn’ method. The 2.5%, 50%, and 97.5% quantile surfaces, respectively, represent pointwise quantiles of the posterior distribution for $w(\mathbf{s}^*)$ for every $\mathbf{s}^* \in \mathcal{S}^*$.

millions of locations with distributed frameworks.

We have focused on developing the distributed framework for spatial modeling due to the motivating applications from massive geostatistical data. The distributed frameworks, however, are applicable to any mixed effects model where the random effects are assigned a GP prior, which includes Bayesian nonparametric regression using GP prior. We plan to explore more general applications in the future with high dimensional covariates.

ACKNOWLEDGEMENTS

The four authors contributed equally to this work. Rajarshi Guhaniyogi’s and Sanvesh Srivastava’s research are partially supported by from Office of Naval Research award no. N00014-18-1-2741 and National Science Foundation DMS-1854662/1854667. Cheng Li’s research is supported by Singapore Ministry of Education Academic Research Funds Tier 1 Grants R155000172133 and R155000201114.

REFERENCES

Agueh, M. and G. Carlier (2011). Barycenters in the Wasserstein space. *SIAM Journal on Mathematical Analysis* 43(2), 904–924.

- Anderes, E., R. Huser, D. Nychka, and M. Coram (2013). Nonstationary positive definite tapering on the plane. *Journal of Computational and Graphical Statistics* 22(4), 848–865.
- Anderson, C., D. Lee, and N. Dean (2014). Identifying clusters in Bayesian disease mapping. *Biostatistics* 15(3), 457–469.
- Bai, Y., P. X.-K. Song, and T. Raghunathan (2012). Joint composite estimating functions in spatiotemporal models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 74(5), 799–824.
- Banerjee, S., B. P. Carlin, and A. E. Gelfand (2014). *Hierarchical Modeling and Analysis for Spatial Data*. CRC Press.
- Banerjee, S., A. E. Gelfand, A. O. Finley, and H. Sang (2008). Gaussian predictive process models for large spatial data sets. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 70(4), 825–848.
- Barbian, M. H. and R. M. Assunção (2017). Spatial subsemble estimator for large geostatistical data. *Spatial Statistics* 22, 68–88.
- Berliner, L. M., C. K. Wikle, and N. Cressie (2000). Long-lead prediction of pacific ssts via bayesian dynamic modeling. *Journal of Climate* 13(22), 3953–3968.
- Bevilacqua, M., C. Caamano-Carrillo, and E. Porcu (2020). Unifying compactly supported and matern covariance functions in spatial statistics. *arXiv preprint arXiv:2008.02904*.
- Bevilacqua, M. and C. Gaetan (2015). Comparing composite likelihood methods based on pairs for spatial gaussian random fields. *Statistics and Computing* 25(5), 877–892.
- Bolin, D. and F. Lindgren (2013). A comparison between markov approximations and other methods for large spatial data sets. *Computational Statistics & Data Analysis* 61, 7–21.
- Bolin, D. and J. Wallin (2020). Multivariate type g matern stochastic partial differential equation random fields. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 82(1), 215–239.
- Chandler, R. E. and S. Bate (2007). Inference for clustered data using the independence log-likelihood. *Biometrika* 94(1), 167–183.
- Cheng, G. and Z. Shang (2017). Computational limits of divide-and-conquer method. *Journal of Machine Learning Research* 18, 1–37.
- Cressie, N. and G. Johannesson (2008). Fixed rank kriging for very large spatial data sets. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 70(1), 209–226.
- Cressie, N. and C. Wikle (2011). *Statistics for Spatio-Temporal Data*. Wiley, Hoboken, NJ.
- Cuturi, M. and A. Doucet (2014). Fast computation of Wasserstein barycenters. In *Proceedings of the 31st International Conference on Machine Learning, JMLR W&CP*, Volume 32.
- Daley, D. J., E. Porcu, and M. Bevilacqua (2015). Classes of compactly supported covariance functions for multivariate random fields. *Stochastic Environmental Research and Risk Assessment* 29(4), 1249–1263.
- Datta, A., S. Banerjee, A. O. Finley, and A. E. Gelfand (2016). Hierarchical nearest-neighbor Gaussian process models for large geostatistical datasets. *Journal of the American Statistical Association* 111(514), 800–812.
- Di Lorenzo, E., N. Schneider, K. Cobb, P. Franks, K. Chhak, A. Miller, J. McWilliams, S. Bograd, H. Arango, E. Curchitser, et al. (2008). North Pacific gyre oscillation links ocean climate and ecosystem change. *Geophysical Research Letters* 35(8).
- Eidsvik, J., B. A. Shaby, B. J. Reich, M. Wheeler, and J. Niemi (2014). Estimation and prediction in spatial models with block composite likelihoods. *Journal of Computational and Graphical Statistics* 23(2), 295–315.
- Finley, A. O., A. Datta, B. D. Cook, D. C. Morton, H. E. Andersen, and S. Banerjee (2019a). Efficient algorithms for bayesian nearest neighbor gaussian processes. *Journal of Computational and Graphical Statistics* 28(2), 401–414.
- Finley, A. O., A. Datta, B. D. Cook, D. C. Morton, H. E. Andersen, and S. Banerjee (2019b). Efficient algorithms for bayesian nearest neighbor gaussian processes. *Journal of Computational and Graphical Statistics* 28(2), 401–414.
- Finley, A. O., H. Sang, S. Banerjee, and A. E. Gelfand (2009). Improving the performance of predictive process modeling for large datasets. *Computational Statistics & Data Analysis* 53(8), 2873–2884.
- Furrer, R., M. G. Genton, and D. Nychka (2006). Covariance tapering for interpolation of large spatial datasets. *Journal of Computational and Graphical Statistics* 15(3), 502–523.
- Gramacy, R. B. and D. W. Apley (2015). Local Gaussian process approximation for large computer experiments. *Journal of Computational and Graphical Statistics* 24(2), 561–578.

- Gramacy, R. B. and B. Haaland (2016). Speeding up neighborhood search in local gaussian process prediction. *Technometrics* 58(3), 294–303.
- Guhaniyogi, R. and S. Banerjee (2018). Meta-kriging: Scalable Bayesian modeling and inference for massive spatial datasets. *Technometrics* 60(4), 430–444.
- Guhaniyogi, R. and S. Banerjee (2019). Multivariate spatial meta kriging. *Statistics & probability letters* 144, 3–8.
- Guhaniyogi, R., A. O. Finley, S. Banerjee, and A. E. Gelfand (2011). Adaptive Gaussian predictive process models for large spatial datasets. *Environmetrics* 22(8), 997–1007.
- Guhaniyogi, R. and B. Sanso (2017). Large multiscale spatial modeling using tree shrinkage priors. *UCSC Technical Report*.
- Guinness, J. (2018). Permutation methods for sharpening Gaussian process approximations. *Technometrics* 60(4), 415–429.
- Guinness, J. (2019). Gaussian process learning via fisher scoring of vecchia’s approximation. *arXiv preprint arXiv:1905.08374*.
- Harville, D. A. (1997). *Matrix algebra from a statistician’s perspective*, Volume 1. Springer.
- Hazra, A. and R. Huser (2021). Estimating high-resolution red sea surface temperature hotspots, using a low-rank semiparametric spatial model. *The Annals of Applied Statistics* 15(2), 572–596.
- Heaton, M. J., W. F. Christensen, and M. A. Terres (2017). Nonstationary gaussian process models using spatial hierarchical clustering from finite differences. *Technometrics* 59(1), 93–101.
- Heaton, M. J., A. Datta, A. Finley, R. Furrer, R. Guhaniyogi, F. Gerber, R. B. Gramacy, D. Hammerling, M. Katzfuss, F. Lindgren, D. W. Nychka, F. Sun, and A. Zammit-Mangion (2019). A case study competition among methods for analyzing large spatial data. *Journal of Agricultural, Biological and Environmental Statistics* 24, 398–425.
- Illian, J. B., S. H. Sørbye, and H. Rue (2012). A toolbox for fitting complex spatial point process models using integrated nested Laplace approximation (INLA). *The Annals of Applied Statistics*, 1499–1530.
- Katzfuss, M. (2017). A multi-resolution approximation for massive spatial datasets. *Journal of the American Statistical Association* 112(517).
- Katzfuss, M. and J. Guinness (2021). A general framework for vecchia approximations of gaussian processes. *Statistical Science* 36(1), 124–141.
- Kaufman, C. G., M. J. Schervish, and D. W. Nychka (2008). Covariance tapering for likelihood-based estimation in large spatial data sets. *Journal of the American Statistical Association* 103(484), 1545–1555.
- Knorr-Held, L. and G. Raßer (2000). Bayesian detection of clusters and discontinuities in disease maps. *Biometrics* 56(1), 13–21.
- Lemos, R. T. and B. Sansó (2009). A spatio-temporal model for mean, anomaly, and trend fields of north Atlantic sea surface temperature. *Journal of the American Statistical Association* 104(485), 5–18.
- Li, C., S. Srivastava, and D. B. Dunson (2017). Simple, scalable and accurate posterior interval estimation. *Biometrika* 104(3), 665–680.
- Lindgren, F., H. Rue, and J. Lindström (2011). An explicit link between gaussian fields and gaussian markov random fields: the stochastic partial differential equation approach. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 73(4), 423–498.
- Lindsten, F., A. M. Johansen, C. A. Naesseth, B. Kirkpatrick, T. B. Schön, J. Aston, and A. Bouchard-Côté (2017). Divide-and-conquer with sequential monte carlo. *Journal of Computational and Graphical Statistics* 26(2), 445–458.
- Mehrotra, S., H. Brantley, J. Westman, L. Bangerter, and A. Maity (2021). Divide-and-conquer mcmc for multivariate binary data. *arXiv preprint arXiv:2102.09008*.
- Minsker, S. (2019). Distributed statistical estimation and rates of convergence in normal approximation. *Electronic Journal of Statistics* 13(2), 5213–5252.
- Minsker, S., S. Srivastava, L. Lin, and D. Dunson (2014). Scalable and robust Bayesian inference via the median posterior. In *Proceedings of the 31st International Conference on Machine Learning (ICML-14)*, pp. 1656–1664.
- Minsker, S., S. Srivastava, L. Lin, and D. B. Dunson (2017). Robust and scalable bayes via a median of subset posterior measures. *The Journal of Machine Learning Research* 18(1), 4488–4527.
- Neiswanger, W., C. Wang, and E. Xing (2014). Asymptotically exact, embarrassingly parallel

- MCMC. In *Proceedings of the 30th International Conference on Uncertainty in Artificial Intelligence*, pp. 623–632.
- Nychka, D., S. Bandyopadhyay, D. Hammerling, F. Lindgren, and S. Sain (2015). A multiresolution Gaussian process model for the analysis of large spatial datasets. *Journal of Computational and Graphical Statistics* 24(2), 579–599.
- Nychka, D., D. Hammerling, S. Sain, and N. Lenssen (2016). Latticekrig: Multiresolution kriging based on markov random fields. R package version 6.4.
- Quiñonero-Candela, J. and C. E. Rasmussen (2005). A unifying view of sparse approximate gaussian process regression. *Journal of Machine Learning Research* 6(Dec), 1939–1959.
- Raskutti, G., M. J. Wainwright, and B. Yu (2012). Minimax-optimal rates for sparse additive models over kernel classes via convex programming. *Journal of Machine Learning Research* 13(Feb), 389–427.
- Ribaret, M., D. Cooley, and A. C. Davison (2012). Bayesian inference from composite likelihoods, with an application to spatial extremes. *Statistica Sinica*, 813–845.
- Rue, H., S. Martino, and N. Chopin (2009). Approximate Bayesian inference for latent Gaussian models by using integrated nested Laplace approximations. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 71(2), 319–392.
- Sang, H. and J. Z. Huang (2012). A full scale approximation of covariance functions for large spatial data sets. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 74(1), 111–132.
- Santin, G. and R. Schaback (2016). Approximation of eigenfunctions in kernel-based spaces. *Advances in Computational Mathematics* 42(4), 973–993.
- Savitsky, T. D. and S. Srivastava (2018). Scalable Bayes under informative sampling. *Scandinavian Journal of Statistics* 45(3), 534–556.
- Scott, S. L., A. W. Blocker, F. V. Bonassi, H. A. Chipman, E. I. George, and R. E. McCulloch (2016). Bayes and big data: the consensus Monte Carlo algorithm. *International Journal of Management Science and Engineering Management* 11(2), 78–88.
- Shang, Z., B. Hao, and G. Cheng (2019). Nonparametric Bayesian aggregation for massive data. *Journal of Machine Learning Research* 20, 1–81.
- Simpson, D., F. Lindgren, and H. Rue (2012). In order to make spatial statistics computationally feasible, we need to forget about the covariance function. *Environmetrics* 23(1), 65–74.
- Srivastava, S., V. Cevher, Q. Dinh, and D. Dunson (2015). WASP: Scalable Bayes via barycenters of subset posteriors. In *Proceedings of the 18th International Conference on Artificial Intelligence and Statistics*, pp. 912–920.
- Srivastava, S., C. Li, and D. B. Dunson (2018). Scalable Bayes via barycenter in Wasserstein space. *Journal of Machine Learning Research* 19, 1–35.
- Srivastava, S. and Y. Xu (2021). Distributed bayesian inference in linear mixed-effects models. *Journal of Computational and Graphical Statistics* 30(3), 594–611.
- Stein, M. L. (2014). Limitations on low rank approximations for covariance matrices of spatial data. *Spatial Statistics* 8, 1–19.
- Stein, M. L., Z. Chi, and L. J. Welty (2004). Approximating likelihoods for large spatial data sets. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 66(2), 275–296.
- Su, Y. (2020). A divide and conquer algorithm of bayesian density estimation. *arXiv preprint arXiv:2002.07094*.
- Szabo, B. and H. van Zanten (2019). An asymptotic analysis of distributed nonparametric methods. *Journal of Machine Learning Research* 20, 1–30.
- van der Vaart, A. and H. van Zanten (2011). Information rates of nonparametric Gaussian process methods. *Journal of Machine Learning Research* 12(Jun), 2095–2119.
- van der Vaart, A. W. and J. H. van Zanten (2008). Reproducing kernel Hilbert spaces of Gaussian priors. In *Pushing the limits of contemporary statistics: contributions in honor of Jayanta K. Ghosh*, pp. 200–222. Institute of Mathematical Statistics.
- Van Trees, H. L. (2001). *Detection, Estimation, and Modulation Theory*. John Wiley & Sons.
- Varin, C., N. Reid, and D. Firth (2011). An overview of composite likelihood methods. *Statistica Sinica*, 5–42.
- Vecchia, A. V. (1988). Estimation and model identification for continuous spatial processes. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 50(2), 297–312.
- Wang, X. and D. B. Dunson (2013). Parallelizing MCMC via Weierstrass sampler. *arXiv preprint arXiv:1312.4605*.
- Wang, X., F. Guo, K. A. Heller, and D. B. Dunson (2015). Parallelizing MCMC with random

- partition trees. *arXiv preprint arXiv:1506.03164*.
- Wendland, H. (2005). *Scattered Data Approximation*. Cambridge University Press.
- Wikle, C. K. (2010). Low-rank representations for spatial processes. *Handbook of Spatial Statistics*, 107–118.
- Wikle, C. K. and S. H. Holan (2011). Polynomial nonlinear spatio-temporal integro-difference equation models. *Journal of Time Series Analysis* 32(4), 339–350.
- Xue, J. and F. Liang (2019). Double-parallel Monte Carlo for Bayesian analysis of big data. *Statistics and Computing* 29(1), 23–32.
- Yang, Y., A. Bhattacharya, and D. Pati (2017). Frequentist coverage and sup-norm convergence rate in gaussian process regression. *arXiv preprint arXiv:1708.04753*.
- Yang, Y., M. Pilanci, M. J. Wainwright, et al. (2017). Randomized sketches for kernels: Fast and optimal nonparametric regression. *The Annals of Statistics* 45(3), 991–1023.
- Zhang, M. M. and S. A. Williamson (2019). Embarrassingly parallel inference for gaussian processes. *Journal of Machine Learning Research* 20, 1–26.
- Zhang, T. (2005). Learning bounds for kernel regression using effective data dimensionality. *Neural Computation* 17(9), 2077–2098.
- Zhang, Y., J. C. Duchi, and M. J. Wainwright (2015). Divide and conquer kernel ridge regression: a distributed algorithm with minimax optimal rates. *Journal of Machine Learning Research* 16, 3299–3340.
- Zhu, H., C. K. I. Williams, R. J. Rohwer, and M. Morciniec (1998). Gaussian regression and optimal finite dimensional linear models. In C. M. Bishop (Ed.), *Neural Networks and Machine Learning*.

Distributed Bayesian Inference in Massive Spatial Data

Rajarshi Guhaniyogi*

Department of Statistics, University of California, Santa Cruz

Cheng Li†

Department of Statistics and Applied Probability, National University of Singapore

Terrance Savitsky‡

U.S. Bureau of Labor Statistics

Sanvesh Srivastava§

Department of Statistics and Actuarial Science, The University of Iowa

Abstract. We provide detailed technical proofs of the theorems and derive the posterior sampling algorithms described in Section 3 of the main manuscript, as well as an extension to the case when τ^2 is unknown and assigned a prior in Section 1.3. We include more detailed information on parameter estimation and convergence of Markov chains for the DISK posterior for the two simulation examples and real data analysis in Section 4 of the main manuscript.

Key words and phrases: Distributed Bayesian inference, Gaussian process, low-rank Gaussian process, modified predictive process, massive spatial data, Wasserstein distance, Wasserstein barycenter.

1. PROOF OF THEOREMS IN SECTION 3.4

Recall that the spatial regression model with a GP prior considered in Section 3.4 is

$$(1) \quad \begin{aligned} y(\mathbf{s}_i) &= w(\mathbf{s}_i) + \epsilon(\mathbf{s}_i), & \epsilon(\mathbf{s}_i) &\sim N(0, \tau^2), \\ w(\cdot) &\sim \text{GP}\{0, \lambda_n^{-1} C_\alpha(\cdot, \cdot)\}, & i &= 1, \dots, n. \end{aligned}$$

Writing this model for the n locations in $\mathcal{S}$ gives

$$(2) \quad \mathbf{y} = \mathbf{w}_0 + \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \mid \mathcal{S} \sim N(\mathbf{0}, \tau^2 \mathbf{I}), \quad \mathbf{y} \mid \mathcal{S} \sim N(\mathbf{w}_0, \tau^2 \mathbf{I}),$$

where $\mathbf{w}_0 = \{w_0(\mathbf{s}_1), \dots, w_0(\mathbf{s}_n)\}$ and $\boldsymbol{\epsilon} = \{\epsilon(\mathbf{s}_1), \dots, \epsilon(\mathbf{s}_n)\}$ are the true value of the residual spatial surface and white noise realized at the locations in $\mathcal{S}$. We can write the model in a similar format for each data subset. Let $\mathbf{s} \in \mathcal{D}$ be a location, $w_0(\mathbf{s})$ be the true value of the residual spatial surface, $\mathbb{E}_{\mathbf{s}^*}$, $\mathbb{E}_0$, $\mathbb{E}_{\mathcal{S}}$, $\mathbb{E}_{\mathbf{y} \mid \mathcal{S}}$, and

*rguhaniy@ucsc.edu

†stalic@nus.edu.sg

‡savitsky.terrance@bls.gov

§sanvesh-srivastava@uiowa.edu Corresponding author.

$\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}}$ respectively be the expectations with respect to the distributions of $\mathbf{s}^*$, $(\mathcal{S}, \mathbf{y})$, $\mathcal{S}$, $\mathbf{y}$ given $\mathcal{S}$, and $(\mathbf{y}, \bar{w}(\mathbf{s}^*))$ given $\mathcal{S}, \mathbf{s}^*$.

In this section, we assume the Assumption A.5, so that the parameters $(\tau^2, \boldsymbol{\alpha})$ are fixed at their truth and the same across all subsets. In this case, if $\bar{w}(\mathbf{s}^*)$ is a random variable that follows the DISK posterior for estimating $w_0(\mathbf{s}^*)$, then conditional on τ^2 and $\boldsymbol{\alpha}$, $\bar{w}(\mathbf{s}^*)$ has the density $N(\bar{m}, \bar{v})$, where

$$\begin{aligned} \bar{m} &= \frac{1}{k} \sum_{j=1}^k \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{y}_j, \\ (3) \quad \bar{v}^{1/2} &= \frac{1}{k} \sum_{j=1}^k v_j^{1/2}, \quad v_j = \lambda_n^{-1} \left\{ c_{*,*} - \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{c}_{j,*} \right\}, \end{aligned}$$

where $c_{*,*} = C_{\boldsymbol{\alpha}}(\mathbf{s}^*, \mathbf{s}^*)$, and $\mathbf{c}_{j,*}^T = \mathbf{c}_j^T(\mathbf{s}^*) = [C_{\boldsymbol{\alpha}}(\mathbf{s}_{j1}, \mathbf{s}^*), \dots, C_{\boldsymbol{\alpha}}(\mathbf{s}_{jm}, \mathbf{s}^*)]$. In the proofs below, without confusion, we use the notation $\mathbf{c}_{j,*}$ and $\mathbf{c}_j(\mathbf{s}^*)$ interchangeably.

The Bayes L_2 -risk in estimating w_0 using the DISK posterior is defined as

$$\begin{aligned} &\mathbb{E}_0 \mathbb{E}_{\mathbf{s}^*} [\{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2] \\ (4) \quad &\stackrel{(i)}{=} \mathbb{E}_{\mathcal{S}} \int_{\mathcal{D}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} [\{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2] \mathbb{P}_{\mathbf{s}}(d\mathbf{s}^*), \end{aligned}$$

where (i) follows from Fubini's theorem. Using bias-variance decomposition,

$$\begin{aligned} &\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} [\{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2] \\ &= \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} [\bar{w}(\mathbf{s}^*) - \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} + \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} - w_0(\mathbf{s}^*)]^2 \\ &= [\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} - w_0(\mathbf{s}^*)]^2 + \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} [\bar{w}(\mathbf{s}^*) - \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\}]^2 \\ &\equiv \text{bias}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}}^2 \{\bar{w}(\mathbf{s}^*)\} + \text{var}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\}. \end{aligned}$$

If $\mathbf{c}_j^T(\cdot) = [\text{cov}\{w(\cdot), w(\mathbf{s}_{j1})\}, \dots, \text{cov}\{w(\cdot), w(\mathbf{s}_{jm})\}] = \{C_{\boldsymbol{\alpha}}(\mathbf{s}_{j1}, \cdot), \dots, C_{\boldsymbol{\alpha}}(\mathbf{s}_{jm}, \cdot)\}$, $\mathbf{c}^T(\cdot) = \{\mathbf{c}_1^T(\cdot), \dots, \mathbf{c}_k^T(\cdot)\}$, $\mathbf{w}_{0j}^T = \{w_0(\mathbf{s}_{j1}), \dots, w_0(\mathbf{s}_{jm})\}$, and $\mathbf{w}_0^T = \{\mathbf{w}_{01}^T, \dots, \mathbf{w}_{0k}^T\}$, then the distribution of $\bar{w}(\mathbf{s}^*)$ in (3) implies that

$$\begin{aligned} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} &= \frac{1}{k} \sum_{j=1}^k \mathbf{c}_j^T(\mathbf{s}^*) \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \mathbf{w}_{0j} \\ &= \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{w}_0, \\ \text{var}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} &= \text{var}_{\mathbf{y} | \mathcal{S}} [\mathbb{E} \{\bar{w}(\mathbf{s}^*) \mid \mathbf{y}\}] + \mathbb{E}_{\mathbf{y} | \mathcal{S}} [\text{var} \{\bar{w}(\mathbf{s}^*) \mid \mathbf{y}\}] \\ &\stackrel{(i)}{=} \text{var}_{\mathbf{y} | \mathcal{S}} \left[\frac{1}{k} \sum_{j=1}^k \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{y}_j \right] + \mathbb{E}_{\mathbf{y} | \mathcal{S}} [\bar{v}(\mathbf{s}^*)] \\ (5) \quad &= \tau^2 \mathbf{c}^T(\mathbf{s}^*) (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-2} \mathbf{c}(\mathbf{s}^*) + \bar{v}(\mathbf{s}^*), \end{aligned}$$

where $\mathbf{L}$ is a block-diagonal matrix with $\mathbf{C}_{1,1}, \dots, \mathbf{C}_{k,k}$ along the diagonal. The equality (i) holds due to the following reasons: (i) The true data follows $y(\mathbf{s}_{ji}) = w_0(\mathbf{s}_{ji}) + \epsilon(\mathbf{s}_{ji})$ ($j = 1, \dots, k$ and $i = 1, \dots, m$), where $\epsilon(\mathbf{s}_{ji})$'s are all independent with variance $\tau_0^2 = \tau^2$ by Assumption A.5; (ii) $\{\mathbf{y}_1, \dots, \mathbf{y}_k\}$ conditional on $\mathcal{S}$ are

jointly independent since they are a disjoint (random) partition of the full dataset by Assumption A.1, which implies that $\text{var}_{\mathbf{y}|\mathcal{S}} \left(\sum_{j=1}^k \mathbf{a}_j^T \mathbf{y}_j \right) = \tau^2 \sum_{j=1}^k \mathbf{a}_j^T \mathbf{a}_j$ for any vectors $\mathbf{a}_1, \dots, \mathbf{a}_k \in \mathbb{R}^m$.

Therefore, the Bayes L_2 -risk in (4) can be decomposed into three parts:

$$(6) \quad \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{w}_0 - w_0(\mathbf{s}^*) \}^2 + \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \mathbf{I})^{-2} \mathbf{c}_* \} \\ + \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \bar{v}(\mathbf{s}^*) \},$$

which correspond to bias^2 , var_{mean} and var_{DISK} in Theorem 3.1.

1.1 Proof of Theorem 3.1

The next three sections find upper bounds for each of the three terms in (6). The conclusion of Theorem 3.1 follows directly by combining the three upper bounds.

1.1.1 An upper bound for the squared bias Consider the squared-bias term in (6). For ease of presentation, assume that $\{\mathbf{s}_1, \dots, \mathbf{s}_n\}$ are relabeled to

$$\{\mathbf{s}_{11}, \dots, \mathbf{s}_{1m}, \dots, \mathbf{s}_{k1}, \dots, \mathbf{s}_{km}\}$$

corresponding to the k subsets. Define $\xi_{\mathbf{s}_{ji}}(\cdot) = C_{\alpha}(\mathbf{s}_{ji}, \cdot)$,

$$(7) \quad \mathbf{w}_0^T = (\langle w_0, \xi_{\mathbf{s}_{11}} \rangle_{\mathbb{H}}, \dots, \langle w_0, \xi_{\mathbf{s}_{1m}} \rangle_{\mathbb{H}}, \dots, \langle w_0, \xi_{\mathbf{s}_{k1}} \rangle_{\mathbb{H}}, \dots, \langle w_0, \xi_{\mathbf{s}_{km}} \rangle_{\mathbb{H}}) \\ \equiv (\mathbf{w}_{01}^T, \dots, \mathbf{w}_{0k}^T), \\ \mathbf{c}^T(\cdot) = (\xi_{\mathbf{s}_{11}}, \dots, \xi_{\mathbf{s}_{1m}}, \dots, \xi_{\mathbf{s}_{k1}}, \dots, \xi_{\mathbf{s}_{km}}) \\ = \{\mathbf{c}_1^T(\cdot), \dots, \mathbf{c}_k^T(\cdot)\} \equiv (\mathbf{c}_1^T, \dots, \mathbf{c}_k^T).$$

The following lemma provides an upper bound on the squared bias of the DISK posterior.

Lemma 1.1 *If Assumptions A.1–A.5 in the main paper hold, then for some global constant $A > 0$,*

$$\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{w}_0 - w_0(\mathbf{s}^*) \}^2 \leq \frac{8\tau^2 \lambda_n}{n} \|w_0\|_{\mathbb{H}}^2 \\ + \|w_0\|_{\mathbb{H}}^2 \inf_{d \in \mathbb{N}} \left[\frac{8n}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \mu_1 \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2 \lambda_n}{n})}{\sqrt{m}} \right\}^q \right].$$

Proof Based on the term $\mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{w}_0$ in (6), we define Δ_j ($j = 1, \dots, k$) and Δ as

$$(8) \quad \Delta_j(\cdot) = \mathbf{y}_j^T (\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I})^{-1} \mathbf{c}_j(\cdot) - w_0(\cdot) \equiv \tilde{w}_j(\cdot) - w_0(\cdot),$$

$$\Delta(\cdot) = \mathbf{y}^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{c}(\cdot) - w_0(\cdot) = \frac{1}{k} \sum_{j=1}^k \{ \tilde{w}_j(\cdot) - w_0(\cdot) \} = \frac{1}{k} \sum_{j=1}^k \Delta_j(\cdot),$$

so that $\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta) = \mathbf{w}_0^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-1} \mathbf{c}(\cdot) - w_0(\cdot) = k^{-1} \sum_{j=1}^k \mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j)$ and $\mathbb{E}_{\mathcal{S}} \|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta)\|_2^2$ yields the bias^2 term in (6). Jensen's inequality implies that

$\|\mathbb{E}_{\mathbf{y}|S}(\Delta)\|_2^2 \leq k^{-1} \sum_{j=1}^k \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2$, so we only need to find upper bounds for $\|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2$ ($j = 1, \dots, k$).

We can recognize that the optimization problem below has $\tilde{w}_j(\cdot)$ defined in (8) as its solution,

$$(9) \quad \operatorname{argmin}_{w \in \mathcal{H}} \sum_{i=1}^m \frac{\{w(\mathbf{s}_{ji}) - y(\mathbf{s}_{ji})\}^2}{2\tau^2/k} + \frac{1}{2}\lambda_n \|w\|_{\mathbb{H}}^2, \quad j = 1, \dots, k.$$

Differentiating (9) and taking expectations with respect to $\mathbb{E}_{\mathbf{y}|S}$ implies that

$$(10) \quad \begin{aligned} & \sum_{i=1}^m \mathbb{E}_{\mathbf{y}|S} \{ \tilde{w}_j(\mathbf{s}_{ji}) - y(\mathbf{s}_{ji}) \} \xi_{\mathbf{s}_{ji}} + \frac{\tau^2 \lambda_n}{k} \mathbb{E}_{\mathbf{y}|S}(\tilde{w}_j) \\ &= \sum_{i=1}^m \langle \mathbb{E}_{\mathbf{y}|S}(\Delta_j), \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} \xi_{\mathbf{s}_{ji}} + \frac{\tau^2 \lambda_n}{k} \mathbb{E}_{\mathbf{y}|S}(\tilde{w}_j) = 0, \end{aligned}$$

where the last inequality follows because $y(\mathbf{s}_{ji}) = \langle w_0, \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} + \epsilon(\mathbf{s}_{ji})$ and $\langle \mathbb{E}_{\mathbf{y}|S}(\epsilon), \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} = \langle 0, \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} = 0$. Using (8), $\Delta_j = \tilde{w}_j - w_0$, $\mathbb{E}_{\mathbf{y}|S}(\tilde{w}_j) = \mathbb{E}_{\mathbf{y}|S}(\Delta_j) + w_0$, and dividing by m in (10), we obtain that

$$(11) \quad \frac{1}{m} \sum_{i=1}^m \langle \mathbb{E}_{\mathbf{y}|S}(\Delta_j), \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} \xi_{\mathbf{s}_{ji}} + \frac{\tau^2 \lambda_n}{km} \mathbb{E}_{\mathbf{y}|S}(\Delta_j) = -\frac{\tau^2 \lambda_n}{km} w_0.$$

If we define the j th sample covariance operator as $\hat{\Sigma}_j = \frac{1}{m} \sum_{i=1}^m \xi_{\mathbf{s}_{ji}} \otimes \xi_{\mathbf{s}_{ji}}$, then (11) reduces to

$$(12) \quad \begin{aligned} & \left(\hat{\Sigma}_j + \frac{\tau^2 \lambda_n}{km} \mathbf{I} \right) \mathbb{E}_{\mathbf{y}|S}(\Delta_j) = -\frac{\tau^2 \lambda_n}{km} w_0 \\ & \implies \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_{\mathbb{H}} \leq \|w_0\|_{\mathbb{H}}, \quad j = 1, \dots, k, \end{aligned}$$

where the last inequality follows because $\hat{\Sigma}_j$ is a positive semi-definite matrix.

The rest of the proof finds an upper bound for $\|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2$. We now reduce this problem to a finite dimensional one indexed by a chosen $d \in \mathbb{N}$. Let $\boldsymbol{\delta}_j = (\delta_{j1}, \dots, \delta_{jd}, \delta_{j(d+1)}, \dots, \delta_{j\infty}) \in L_2(\mathbb{N})$ such that

$$(13) \quad \begin{aligned} \mathbb{E}_{\mathbf{y}|S}(\Delta_j) &= \sum_{i=1}^{\infty} \delta_{ji} \varphi_i, \quad \delta_{ji} = \langle \mathbb{E}_{\mathbf{y}|S}(\Delta_j), \varphi_i \rangle_{L^2(\mathbb{P})}, \\ \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2 &= \sum_{i=1}^{\infty} \delta_{ji}^2, \quad j = 1, \dots, k. \end{aligned}$$

Define the vectors $\boldsymbol{\delta}_j^{\downarrow} = (\delta_{j1}, \dots, \delta_{jd})$ and $\boldsymbol{\delta}_j^{\uparrow} = (\delta_{j(d+1)}, \dots, \delta_{j\infty})$, so

$\|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2 = \|\boldsymbol{\delta}_j^{\downarrow}\|_2^2 + \|\boldsymbol{\delta}_j^{\uparrow}\|_2^2$ and we upper bound $\|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2$ by separately upper bounding $\|\boldsymbol{\delta}_j^{\downarrow}\|_2^2$ and $\|\boldsymbol{\delta}_j^{\uparrow}\|_2^2$. Using the expansion $C_{\alpha}(\mathbf{s}, \mathbf{s}') = \sum_{j=1}^{\infty} \mu_j \varphi_j(\mathbf{s}) \varphi_j(\mathbf{s}')$ for any $\mathbf{s}, \mathbf{s}' \in \mathcal{D}$, we have the following upper bound for $\|\boldsymbol{\delta}_j^{\uparrow}\|_2^2$:

$$(14) \quad \|\boldsymbol{\delta}_j^{\uparrow}\|_2^2 = \frac{\mu_{d+1}}{\mu_{d+1}} \sum_{i=d+1}^{\infty} \delta_{ji}^2 \leq \mu_{d+1} \sum_{i=d+1}^{\infty} \frac{\delta_{ji}^2}{\mu_i} \stackrel{(i)}{\leq} \mu_{d+1} \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_{\mathbb{H}}^2 \stackrel{(ii)}{\leq} \mu_{d+1} \|w_0\|_{\mathbb{H}}^2,$$

where (i) follows because $\|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j)\|_{\mathbb{H}}^2 = \sum_{i=1}^{\infty} \delta_{ji}^2/\mu_i$ and (ii) follows from (12).

We then derive an upper bound for $\|\delta_j^\downarrow\|_2^2$. Let $\mathbf{M} = \text{diag}(\mu_1, \dots, \mu_d) \in \mathbb{R}^{d \times d}$, $\Phi^j \in \mathbb{R}^{m \times d}$ be a matrix such that

$$(15) \quad \Phi_{ih}^j = \varphi_h(\mathbf{s}_{ji}), \quad i = 1, \dots, m, \quad h = 1, \dots, d, \quad j = 1, \dots, k,$$

$w_0 = \sum_{i=1}^{\infty} \theta_i \varphi_i$, and the tail error vector $\mathbf{v}_j = (v_{j1}, \dots, v_{jm})^T \in \mathbb{R}^m$ ($j = 1, \dots, k$) such that

$$v_{ji} = \sum_{h=d+1}^{\infty} \delta_{jh} \varphi_h(\mathbf{s}_{ji}), \quad i = 1, \dots, m.$$

For any $g \in \{1, \dots, d\}$, taking the $\mathbb{H}$ -inner product with respect φ_g in (12) yields

$$(16) \quad \begin{aligned} & \left\langle \left(\frac{1}{m} \sum_{i=1}^m \xi_{\mathbf{s}_{ji}} \otimes \xi_{\mathbf{s}_{ji}} + \frac{\tau^2 \lambda_n}{km} \mathbf{I} \right) \mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j), \varphi_g \right\rangle_{\mathbb{H}} \\ &= -\frac{\tau^2 \lambda_n}{km} \langle w_0, \varphi_g \rangle_{\mathbb{H}} = -\frac{\tau^2 \lambda_n}{km} \frac{\theta_g}{\mu_g}, \quad j = 1, \dots, k. \end{aligned}$$

Expanding the left hand side in (16), we obtain that

$$\begin{aligned} & \frac{1}{m} \sum_{i=1}^m \langle \varphi_g, \xi_{\mathbf{s}_{ji}} \rangle_{\mathbb{H}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \{ \Delta_j(\mathbf{s}_{ji}) \} + \frac{\tau^2 \lambda_n}{km} \langle \varphi_g, \mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j) \rangle_{\mathbb{H}} \\ &= \frac{1}{m} \sum_{i=1}^m \varphi_g(\mathbf{s}_{ji}) \mathbb{E}_{\mathbf{y}|\mathcal{S}} \{ \Delta_j(\mathbf{s}_{ji}) \} + \frac{\tau^2 \lambda_n}{km} \frac{\delta_{jg}}{\mu_g}. \end{aligned}$$

The term $\frac{1}{m} \sum_{i=1}^m \varphi_g(\mathbf{s}_{ji}) \mathbb{E}_{\mathbf{y}|\mathcal{S}} \{ \Delta_j(\mathbf{s}_{ji}) \}$ on the right hand side is

$$(17) \quad \begin{aligned} &= \frac{1}{m} \sum_{i=1}^m \Phi_{ig}^j \sum_{h=1}^d \delta_{jh} \varphi_h(\mathbf{s}_{ji}) + \frac{1}{m} \sum_{i=1}^m \Phi_{ig}^j \sum_{h=d+1}^{\infty} \delta_{jh} \varphi_h(\mathbf{s}_{ji}) \\ &= \frac{1}{m} \sum_{h=1}^d \delta_{jh} \sum_{i=1}^m \Phi_{ig}^j \Phi_{ih}^j + \frac{1}{m} \sum_{i=1}^m \Phi_{ig}^j v_{ji} \\ &= \frac{1}{m} \sum_{h=1}^d \delta_{jh} \left(\Phi^{jT} \Phi^j \right)_{gh} + \frac{1}{m} \sum_{i=1}^m \left(\Phi^{jT} \mathbf{v}_j \right)_g \\ &= \frac{1}{m} \left(\Phi^{jT} \Phi^j \delta_j^\downarrow \right)_g + \frac{1}{m} \left(\Phi^{jT} \mathbf{v}_j \right)_g. \end{aligned}$$

Substitute (17) in (16) for $g = 1, \dots, d$ to obtain that

$$(18) \quad \begin{aligned} & \frac{1}{m} \Phi^{jT} \Phi^j \delta_j^\downarrow + \frac{1}{m} \Phi^{jT} \mathbf{v}_j + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \delta_j^\downarrow = -\frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \\ & \left(\frac{1}{m} \Phi^{jT} \Phi^j + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \right) \delta_j^\downarrow = -\frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow - \frac{1}{m} \Phi^{jT} \mathbf{v}_j. \end{aligned}$$

The proof is completed by showing that the right hand side expression in (18) gives an upper bound for $\|\delta_j^\downarrow\|_2^2$. Define $\mathbf{Q} = \left(\mathbf{I} + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \right)^{1/2}$, then

$$\frac{1}{m} \Phi^{jT} \Phi^j + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} = \mathbf{I} + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} + \frac{1}{m} \Phi^{jT} \Phi^j - \mathbf{I}$$

$$= \mathbf{Q} \left\{ \mathbf{I} + \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{j^T} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \right\} \mathbf{Q}$$

and using this in (18) gives

$$(19) \quad \left\{ \mathbf{I} + \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{j^T} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \right\} \mathbf{Q} \delta_j^\downarrow = -\frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow - \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j.$$

Now we define the $\mathbb{P}$ -measureable event

$$(20) \quad \mathcal{E}_1 = \left\{ \left\| \left\| \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{j^T} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \right\| \right\| \leq 1/2 \right\},$$

where $\|\cdot\|$ is the matrix operator norm. We have that $\mathbf{I} + \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{j^T} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \succeq (1/2) \mathbf{I}$ whenever $\mathcal{E}_1$ occurs. Furthermore, when $\mathcal{E}_1$ occurs, (19) implies that

$$\begin{aligned} \|\delta_j^\downarrow\|_2^2 &\leq \|\mathbf{Q} \delta_j^\downarrow\|_2^2 \leq 4 \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow + \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2 \\ &\leq 8 \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \right\|_2^2 + 8 \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2, \end{aligned}$$

where the last inequality follows because $(a+b)^2 \leq 2a^2 + 2b^2$ for any $a, b \in \mathbb{R}$.

Since $\mathcal{E}_1$ is $\mathbb{P}$ -measureable, $\mathbb{E}_{\mathcal{S}} \left(\|\delta_j^\downarrow\|_2^2 \right) = \mathbb{E}_{\mathcal{S}} \left\{ \|\delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1) \right\} + \mathbb{E}_{\mathcal{S}} \left\{ \|\delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1^c) \right\}$ and the previous display gives

$$(21) \quad \mathbb{E}_{\mathcal{S}} \left\{ \|\delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1) \right\} \leq 8 \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \right\|_2^2 + 8 \mathbb{E}_{\mathcal{S}} \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2.$$

From Lemma 10 in Zhang et al. (2015), we have that under our assumptions A.1-A.5, there exists a universal constant $A > 0$ that does not depend on λ_n, n, τ^2 , such that

$$\begin{aligned} \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \right\|_2^2 &\leq \frac{\tau^2 \lambda_n}{km} \|w_0\|_{\mathbb{H}}^2, \\ \mathbb{E}_{\mathcal{S}} \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2 &\leq \frac{km}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\boldsymbol{\alpha}}) \text{tr}(C_{\boldsymbol{\alpha}}^d) \|w_0\|_{\mathbb{H}}^2, \\ \mathbb{P}(\mathcal{E}_1^c) &\leq \left\{ A \max \left(\sqrt{\max(q, \log d)}, \frac{\max(q, \log d)}{m^{1/2-1/q}} \right) \frac{\rho^2 \gamma(\frac{\tau^2 \lambda_n}{km})}{\sqrt{m}} \right\}^q \\ (22) \quad &= \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2 \lambda_n}{km})}{\sqrt{m}} \right\}^q. \end{aligned}$$

Since $\mu_1 \geq \mu_2 \geq \dots \geq 0$, the optimality condition in (12) implies that

$$(23) \quad \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_2^2 = \frac{\mu_1}{\mu_1} \sum_{i=1}^{\infty} \delta_{ji} \varphi_i \leq \mu_1 \sum_{i=1}^{\infty} \frac{\delta_{ji}}{\mu_i} \varphi_i = \mu_1 \|\mathbb{E}_{\mathbf{y}|S}(\Delta_j)\|_{\mathbb{H}}^2 \leq \mu_1 \|w_0\|_{\mathbb{H}}^2.$$

Using the shorthand (22) and (23), we obtain that

$$(24) \quad \mathbb{E}_{\mathcal{S}} \left\{ \|\boldsymbol{\delta}_j^\dagger\|_2^2 \mathbf{1}(\mathcal{E}_1^c) \right\} \leq \mathbb{E}_{\mathcal{S}} \left\{ \|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j)\|_2^2 \mathbf{1}(\mathcal{E}_1^c) \right\} \leq \mathbb{P}(\mathcal{E}_1^c) \mu_1 \|w_0\|_{\mathbb{H}}^2.$$

Combining (21) and (24) gives

$$(25) \quad \begin{aligned} \mathbb{E}_{\mathcal{S}}(\|\boldsymbol{\delta}_j\|_2^2) &\leq \frac{8\tau^2\lambda_n}{km} \|w_0\|_{\mathbb{H}}^2 + \frac{8km}{\tau^2\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \|w_0\|_{\mathbb{H}}^2 \\ &\quad + \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2\lambda_n}{km})}{\sqrt{m}} \right\}^q \mu_1 \|w_0\|_{\mathbb{H}}^2. \end{aligned}$$

Finally, we use that $\|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta)\|_2^2 \leq k^{-1} \sum_{j=1}^k \|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta_j)\|_2^2 = k^{-1} \sum_{j=1}^k \|\boldsymbol{\delta}_j\|_2^2$ to obtain that

$$(26) \quad \begin{aligned} \mathbb{E}_{\mathcal{S}}(\|\mathbb{E}_{\mathbf{y}|\mathcal{S}}(\Delta)\|_2^2) &\leq \frac{8\tau^2\lambda_n}{km} \|w_0\|_{\mathbb{H}}^2 + \frac{8km}{\tau^2\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \|w_0\|_{\mathbb{H}}^2 \\ &\quad + \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2\lambda_n}{km})}{\sqrt{m}} \right\}^q \mu_1 \|w_0\|_{\mathbb{H}}^2 \\ &= \frac{8\tau^2\lambda_n}{n} \|w_0\|_{\mathbb{H}}^2 + \|w_0\|_{\mathbb{H}}^2 \left[\frac{8n}{\tau^2\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \mu_1 \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2\lambda_n}{n})}{\sqrt{m}} \right\}^q \right], \end{aligned}$$

where we have replaced km by n in the last equality. Taking the infimum over $d \in \mathbb{N}$ leads to the proof. $\blacksquare$

1.1.2 An upper bound for the first variance term The following lemma provides an upper bound the first part of the variance term in (6).

Lemma 1.2 *If Assumptions A.1–A.5 in the main paper hold, then*

$$\begin{aligned} \tau^2 \mathbb{E}_{\mathcal{S}^*} \mathbb{E}_{\mathcal{S}} \left\{ \mathbf{c}_*^T (k\mathbf{L} + \tau^2\lambda_n \mathbf{I})^{-2} \mathbf{c}_* \right\} &\leq \\ &\left(\frac{2n}{k\lambda_n} + \frac{4\|w_0\|_{\mathbb{H}}^2}{k} \right) \inf_{d \in \mathbb{N}} \left[\mu_{d+1} + 12 \frac{n}{\tau^2\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \right. \\ &\quad \left. + \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2\lambda_n}{n})}{\sqrt{m}} \right\}^q \right] + \frac{12\tau^2\lambda_n}{kn} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau^2\lambda_n}{n} \gamma \left(\frac{\tau^2\lambda_n}{n} \right). \end{aligned}$$

Proof Continuing from (8), we start by finding an upper bound for $\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|\mathcal{S}} \|\Delta_j\|_{\mathbb{H}}^2$, which is required later to upper bound $\mathbb{E}_0 \|\Delta_j\|_{\mathbb{H}}^2$. From (8) we have

$$(27) \quad \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|\mathcal{S}} \|\Delta_j\|_{\mathbb{H}}^2 \leq 2 \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|\mathcal{S}} \|\tilde{w}_j\|_{\mathbb{H}}^2 + 2\|w_0\|_{\mathbb{H}}^2.$$

An upper bound for $\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|\mathcal{S}} \|\tilde{w}_j\|_{\mathbb{H}}^2$ gives the desired bound. Using the objective in (9),

$$\frac{1}{2} \|\tilde{w}_j\|_{\mathbb{H}}^2 \stackrel{(i)}{\leq} \sum_{i=1}^m \frac{\{\tilde{w}_j(\mathbf{s}_{ji}) - y(\mathbf{s}_{ji})\}^2}{2\tau^2\lambda_n/k} + \frac{1}{2} \|\tilde{w}_j\|_{\mathbb{H}}^2$$

$$(28) \quad \stackrel{(ii)}{\leq} \sum_{i=1}^m \frac{\{w_0(\mathbf{s}_{ji}) - y(\mathbf{s}_{ji})\}^2}{2\tau^2\lambda_n/k} + \frac{1}{2}\|w_0\|_{\mathbb{H}}^2,$$

where (i) follows because the term inside the summation is non-negative and (ii) follows because $\tilde{w}_j$ minimizes the objective. Since $w(\mathbf{s}_{ji}) - y(\mathbf{s}_{ji}) = -\epsilon(\mathbf{s}_{ji})$ and $\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \{\epsilon^2(\mathbf{s}_{ji})\} \leq \tau^2$ by Assumption A.2, (28) reduces to

$$(29) \quad \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\tilde{w}_j\|_{\mathbb{H}}^2 \leq \frac{k}{\tau^2\lambda_n} \sum_{i=1}^m \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \{\epsilon(\mathbf{s}_{ji})\}^2 + \|w_0\|_{\mathbb{H}}^2 \leq \frac{km}{\lambda_n} + \|w_0\|_{\mathbb{H}}^2.$$

Substituting (29) in (27) gives

$$(30) \quad \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j\|_{\mathbb{H}}^2 \leq \frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2.$$

First notice that

$$(31) \quad \begin{aligned} & \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k\mathbf{L} + \tau^2\lambda_n \mathbf{I})^{-2} \mathbf{c}_* \} \\ &= \frac{1}{k^2} \sum_{j=1}^k \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \left\{ \mathbf{c}_{j*}^T \left(\mathbf{C}_{j,j} + \frac{\tau^2\lambda_n}{k} \mathbf{I} \right)^{-2} \mathbf{c}_{j*} \right\}. \end{aligned}$$

and from (6) we have

$$(32) \quad \begin{aligned} & \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \left\{ \mathbf{c}_{j*}^T \left(\mathbf{C}_{j,j} + \frac{\tau^2\lambda_n}{k} \mathbf{I} \right)^{-2} \mathbf{c}_{j*} \right\} \\ &= \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \text{var}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \left\{ \mathbf{c}_{j*}^T \left(\mathbf{C}_{j,j} + \frac{\tau^2\lambda_n}{k} \mathbf{I} \right)^{-1} \mathbf{y}_j \right\} \\ &\leq \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \left\{ \mathbf{c}_{j*}^T \left(\mathbf{C}_{j,j} + \frac{\tau^2\lambda_n}{k} \mathbf{I} \right)^{-1} \mathbf{y}_j - w_0(\mathbf{s}^*) \right\}^2 \\ &= \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j\|_2^2. \end{aligned}$$

Substituting (32) to (31) leads to

$$(33) \quad \tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k\mathbf{L} + \tau^2\lambda_n \mathbf{I})^{-2} \mathbf{c}_* \} \leq \mathbb{E}_{\mathbf{s}^*} \left\{ \frac{1}{k^2} \sum_{j=1}^k \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j\|_2^2 \right\}.$$

We then find an upper bound for $\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j\|_2^2$ by following similar steps to the proof of Lemma 1.1. Let $\boldsymbol{\delta}_j \in L_2(\mathbb{N})$ be the expansion of Δ_j in the basis $\{\varphi_i\}_{i=1}^\infty$, so that $\Delta_j = \sum_{i=1}^\infty \delta_{ji} \varphi_i$ (the $\boldsymbol{\delta}_j$ sequence here is different from the one in the previous section). Similar to Section 1.1.1, choose a fixed $d \in \mathbb{N}$ and truncate Δ_j by defining $\Delta_j^\downarrow$, $\Delta_j^\uparrow$, $\boldsymbol{\delta}_j^\downarrow$, and $\boldsymbol{\delta}_j^\uparrow$ as

$$\begin{aligned} \Delta_j^\downarrow &= \sum_{i=1}^d \delta_{ji} \varphi_i, & \Delta_j^\uparrow &= \sum_{i=d+1}^\infty \delta_{ji} \varphi_i = \Delta_j - \Delta_j^\downarrow, \\ \boldsymbol{\delta}_j^\downarrow &= (\delta_{j1}, \dots, \delta_{jd}), & \boldsymbol{\delta}_j^\uparrow &= (\delta_{j(d+1)}, \dots, \delta_{j\infty}). \end{aligned}$$

The orthonormality of $\{\varphi_i\}_{i=1}^\infty$ implies that

$$\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j\|_2^2 = \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*)|S} \|\Delta_j^\downarrow\|_2^2 + \mathbb{E}_{\mathcal{S}} \mathbb{E}_{0|S} \|\Delta_j^\uparrow\|_2^2$$

$$(34) \quad = \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\delta_j^\downarrow\|_2^2 + \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\delta_j^\uparrow\|_2^2.$$

First, the upper bound for $\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\delta_j^\uparrow\|_2^2$ follows from (14),

$$\begin{aligned} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j^\uparrow\|_2^2 &= \sum_{i=d+1}^{\infty} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} (\delta_{ji}^2) \\ &= \mu_{d+1} \sum_{i=d+1}^{\infty} \frac{\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} (\delta_{ji}^2)}{\mu_{d+1}} \leq \mu_{d+1} \sum_{i=d+1}^{\infty} \frac{\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} (\delta_{ji}^2)}{\mu_i} \\ &= \mu_{d+1} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j^\uparrow\|_{\mathbb{H}}^2 \leq \mu_{d+1} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j\|_{\mathbb{H}}^2, \end{aligned}$$

and using (30),

$$(35) \quad \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j^\uparrow\|_2^2 \leq \mu_{d+1} \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right).$$

We now find an upper bound for $\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j^\downarrow\|_2^2$. Following Section 1.1.1, define the error vector $\mathbf{v}_j = (v_{j1}, \dots, v_{jm})^T \in \mathbb{R}^m$ with $v_{ji} = \sum_{h=d+1}^{\infty} \delta_{ji} \varphi_h(\mathbf{s}_{ji})$ ($i = 1, \dots, m$), and $\mathbf{M} = \text{diag}(\mu_1, \dots, \mu_d)$. From (9) and (10), $\tilde{w}_j(\cdot)$ in (8) satisfies

$$(36) \quad \frac{1}{m} \sum_{i=1}^m \langle \xi_{\mathbf{s}_{ji}}, \tilde{w}_j - w_0 - \epsilon \rangle_{\mathbb{H}} \xi_{\mathbf{s}_{ji}} + \frac{\tau^2 \lambda_n}{km} \tilde{w}_j = 0.$$

For any $g \in \{1, \dots, d\}$, taking the $\mathbb{H}$ -inner product with respect φ_g in (36) to obtain that

$$\begin{aligned} &\frac{1}{m} \sum_{i=1}^m \langle \xi_{\mathbf{s}_{ji}}, \Delta_j - \epsilon \rangle_{\mathbb{H}} \langle \xi_{\mathbf{s}_{ji}}, \varphi_g \rangle_{\mathbb{H}} + \frac{\tau^2 \lambda_n}{km} \langle \Delta_j + w_0, \varphi_g \rangle_{\mathbb{H}} = \\ &\frac{1}{m} \sum_{i=1}^m \{ \Delta_j(\mathbf{s}_{ji}) - \epsilon(\mathbf{s}_{ji}) \} \varphi_g(\mathbf{s}_{ji}) + \frac{\tau^2 \lambda_n}{km} \frac{\delta_{jg}}{\mu_g} + \frac{\tau^2 \lambda_n}{km} \frac{\theta_g}{\mu_g} = 0, \\ &\frac{1}{m} \sum_{i=1}^m \left\{ \sum_{h=1}^d \delta_{jh} \varphi_h(\mathbf{s}_{ji}) + \sum_{h=d+1}^{\infty} \delta_{jh} \varphi_h(\mathbf{s}_{ji}) - \epsilon(\mathbf{s}_{ji}) \right\} \varphi_g(\mathbf{s}_{ji}) + \frac{\tau^2 \lambda_n}{km} \frac{\delta_{jg}}{\mu_g} = -\frac{\tau^2 \lambda_n}{km} \frac{\theta_g}{\mu_g}, \\ &\frac{1}{m} \sum_{h=1}^d \left\{ \sum_{i=1}^m \varphi_h(\mathbf{s}_{ji}) \varphi_g(\mathbf{s}_{ji}) \right\} \delta_{jh} + \frac{1}{m} \sum_{i=1}^m \{ v_{ji} - \epsilon(\mathbf{s}_{ji}) \} \varphi_g(\mathbf{s}_{ji}) + \frac{\tau^2 \lambda_n}{km} \frac{\delta_{jg}}{\mu_g} = -\frac{\tau^2 \lambda_n}{km} \frac{\theta_g}{\mu_g}, \\ &\frac{1}{m} \left(\Phi^{jT} \Phi^j \delta_j^\downarrow \right)_g + \frac{1}{m} \left\{ \Phi^{jT} (\mathbf{v}_j - \epsilon_j) \right\}_g + \frac{\tau^2 \lambda_n}{km} (\mathbf{M}^{-1} \delta_j^\downarrow)_g = -\frac{\tau^2 \lambda_n}{km} (\mathbf{M}^{-1} \theta^\downarrow)_g. \end{aligned}$$

Writing this equation in the matrix form yields,

$$(37) \quad \left(\frac{1}{m} \Phi^{jT} \Phi^j + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \right) \delta_j^\downarrow = -\frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \theta^\downarrow - \frac{1}{m} \Phi^{jT} \mathbf{v}_j + \frac{1}{m} \Phi^{jT} \epsilon_j.$$

Following Section 1.1.1, by defining $\mathbf{Q} = (\mathbf{I} + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1})^{1/2}$, (37) reduces to

$$\left\{ \mathbf{I} + \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{jT} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \right\} \mathbf{Q} \delta_j^\downarrow$$

$$(38) \quad = -\frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow - \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j + \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \boldsymbol{\epsilon}_j.$$

On the event $\mathcal{E}_1$ defined as in (20), we have that $\mathbf{I} + \mathbf{Q}^{-1} \left(\frac{1}{m} \Phi^{j^T} \Phi^j - \mathbf{I} \right) \mathbf{Q}^{-1} \succeq (1/2) \mathbf{I}$. Furthermore, when $\mathcal{E}_1$ occurs, (38) implies that

$$\begin{aligned} \|\Delta_j^\downarrow\|_2^2 &\leq \|\mathbf{Q} \delta_j^\downarrow\|_2^2 \leq 4 \left\| -\frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow - \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j + \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \boldsymbol{\epsilon}_j \right\|_2^2 \\ &\leq 12 \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \right\|_2^2 + 12 \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2 + 12 \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \boldsymbol{\epsilon}_j \right\|_2^2, \end{aligned}$$

where the last inequality follows because $(a + b + c)^2 \leq 3a^2 + 3b^2 + 3c^2$ for any $a, b, c \in \mathbb{R}$. Since $\mathcal{E}_1$ is $\mathbb{P}$ -measureable,

$$\mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left(\|\Delta_j^\downarrow\|_2^2 \right) = \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \|\Delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1) \right\} + \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \|\Delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1^c) \right\}.$$

If the event $\mathcal{E}_1$ occurs, then the upper bounds for the first term and the last two terms in the last inequality are given by Lemmas 10 and 7 of Zhang et al. (2015), respectively, and we have that

$$\begin{aligned} \left\| \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-1} \mathbf{M}^{-1} \boldsymbol{\theta}^\downarrow \right\|_2^2 &\leq \frac{\tau^2 \lambda_n}{km} \|w_0\|_{\mathbb{H}}^2, \\ \mathbb{E}_{\mathcal{S}} \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \mathbf{v}_j \right\|_2^2 &\leq \frac{km}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\boldsymbol{\alpha}}) \text{tr}(C_{\boldsymbol{\alpha}}^d) \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right), \\ \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \boldsymbol{\epsilon}_j \right\|_2^2 \\ (39) \quad &\leq \frac{1}{m^2} \sum_{h=1}^d \sum_{i=1}^m \frac{1}{1 + \frac{\tau^2 \lambda_n}{km} \frac{1}{\mu_h}} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \varphi_h^2(\mathbf{s}_{ji}) \epsilon^2(\mathbf{s}_{ji}) \right\}. \end{aligned}$$

Since the error $\epsilon(\cdot)$ and $w(\cdot)$ are independent, by Assumption A.4,

$$\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \varphi_h^2(\mathbf{s}_{ji}) \epsilon^2(\mathbf{s}_{ji}) \right\} = \mathbb{E}_{\mathcal{S}} \left\{ \varphi_h^2(\mathbf{s}_{ji}) \right\} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \epsilon^2(\mathbf{s}_{ji}) \right\} \leq \tau^2,$$

and the last inequality in (39) simplifies to

$$\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\| \frac{1}{m} \mathbf{Q}^{-1} \Phi^{j^T} \boldsymbol{\epsilon}_j \right\|_2^2 \leq \frac{\tau^2}{m} \sum_{h=1}^d \frac{1}{1 + \frac{\tau^2 \lambda_n}{km} \frac{1}{\mu_h}} \leq \frac{\tau^2}{m} \gamma \left(\frac{\tau^2 \lambda_n}{km} \right).$$

Hence when the event $\mathcal{E}_1$ occurs,

$$\begin{aligned} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \|\Delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1) \right\} &\leq \\ (40) \quad &12 \frac{\tau^2 \lambda_n}{km} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{km}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\boldsymbol{\alpha}}) \text{tr}(C_{\boldsymbol{\alpha}}^d) \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right) + 12 \frac{\tau^2}{m} \gamma \left(\frac{\tau^2 \lambda_n}{km} \right). \end{aligned}$$

If the event $\mathcal{E}_1$ does not occur, then

$$\mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \left\{ \|\Delta_j^\downarrow\|_2^2 \mathbf{1}(\mathcal{E}_1^c) \right\}$$

$$\begin{aligned}
&\leq \mathbb{E}_{\mathcal{S}} \left\{ \mathbf{1}(\mathcal{E}_1^c) \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \|\Delta_j^\dagger\|_2^2 \right\} \stackrel{(i)}{\leq} \mathbb{P}(\mathcal{E}_1^c) \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right) \\
(41) \quad &\stackrel{(ii)}{=} \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2 \lambda_n}{km})}{\sqrt{m}} \right\}^q \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right),
\end{aligned}$$

where (i) follows from (30) and (ii) follows from (22). Substituting (40), (41), and (35) in (34) implies that

$$\begin{aligned}
(42) \quad \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{ \|\Delta_j\|_2^2 \} &\leq 12 \frac{\tau^2 \lambda_n}{km} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau^2}{m} \gamma \left(\frac{\tau^2 \lambda_n}{km} \right) + \\
&\quad \left[\mu_{d+1} + 12 \frac{km}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \right. \\
&\quad \left. \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2 \lambda_n}{km})}{\sqrt{m}} \right\}^q \right] \left(\frac{2km}{\lambda_n} + 4\|w_0\|_{\mathbb{H}}^2 \right).
\end{aligned}$$

Therefore, substituting (42) in (33) implies that

$$\begin{aligned}
(43) \quad &\tau^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-2} \mathbf{c}_* \} \leq \\
&\quad \left(\frac{2n}{k \lambda_n} + \frac{4\|w_0\|_{\mathbb{H}}^2}{k} \right) \left[\mu_{d+1} + 12 \frac{n}{\tau^2 \lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \right. \\
&\quad \left. \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2 \lambda_n}{n})}{\sqrt{m}} \right\}^q \right] + \frac{12 \tau^2 \lambda_n}{kn} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right).
\end{aligned}$$

where we have replace km by n . Taking the infimum over $d \in \mathbb{N}$ leads to the proof. $\blacksquare$

1.1.3 An upper bound for the second variance term The following lemma provides an upper bound the second part of the variance term in (6).

Lemma 1.3 *If Assumptions A.1–A.5 in the main paper hold, then*

$$\begin{aligned}
\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \bar{v}(\mathbf{s}^*) &\leq 3 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) \\
&\quad + \inf_{d \in \mathbb{N}} \left[\left\{ \frac{4n}{\tau^2 \lambda_n^2} \text{tr}(C_{\alpha}) + \frac{1}{\lambda_n} \right\} \text{tr}(C_{\alpha}^d) + \lambda_n^{-1} \text{tr}(C_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\tau^2}{n})}{\sqrt{m}} \right\}^q \right].
\end{aligned}$$

Proof First we have the following relation between $\bar{v}$ and the subset variance v_j :

$$\begin{aligned}
(44) \quad \bar{v}(\mathbf{s}^*) &= \left\{ k^{-1} \sum_{j=1}^k v_j^{1/2}(\mathbf{s}^*) \right\}^2 \leq \frac{1}{k} \sum_{j=1}^k v_j(\mathbf{s}^*) \\
&= \frac{1}{k} \sum_{j=1}^k \lambda_n^{-1} \left\{ C_{\alpha}(\mathbf{s}^*, \mathbf{s}^*) - \mathbf{c}_j^T(\mathbf{s}^*) \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \mathbf{c}_j(\mathbf{s}^*) \right\}.
\end{aligned}$$

Since $C_{\alpha}(\mathbf{s}, \mathbf{s}') = \sum_{i=1}^{\infty} \mu_i \varphi_i(\mathbf{s}) \varphi_i(\mathbf{s}')$ for $\mathbf{s}, \mathbf{s}' \in \mathcal{D}$, we have

$$C_{\alpha}(\mathbf{s}^*, \mathbf{s}^*) = \sum_{a=1}^{\infty} \mu_a \varphi_a^2(\mathbf{s}^*), \quad \{\mathbf{c}_j(\mathbf{s}^*)\}_i = \sum_{a=1}^{\infty} \mu_a \varphi_a(\mathbf{s}_{ji}) \varphi_a(\mathbf{s}^*), \quad i = 1, \dots, m.$$

These together with the orthogonality property of $\{\varphi_i\}_{i=1}^{\infty}$ imply that

$$\begin{aligned} \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{v_j(\mathbf{s}^*)\} &= \lambda_n^{-1} \sum_{a=1}^{\infty} \mu_a \mathbb{E}_{\mathbf{s}^*} \varphi_a^2(\mathbf{s}^*) \\ &\quad - \lambda_n^{-1} \sum_{i=1}^m \sum_{i'=1}^m \sum_{a=1}^{\infty} \sum_{b=1}^{\infty} \mu_a \mu_b \left\{ \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \right\}_{i'i''} \\ &\quad \times \mathbb{E}_{\mathcal{S}} [\varphi_a(\mathbf{s}_{ji}) \varphi_b(\mathbf{s}_{ji'}) \mathbb{E}_{\mathbf{s}^*} \{\varphi_a(\mathbf{s}^*) \varphi_b(\mathbf{s}^*)\}] \\ &= \lambda_n^{-1} \text{tr}(C_{\alpha}) - \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \sum_{i=1}^m \sum_{i'=1}^m \sum_{a=1}^{\infty} \mu_a^2 \left\{ \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \right\}_{ii'} \varphi_a(\mathbf{s}_{ji}) \varphi_a(\mathbf{s}_{ji'}) \\ &= \lambda_n^{-1} \sum_{a=1}^d \mu_a - \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \sum_{a=1}^d \mu_a^2 \left[\sum_{i=1}^m \sum_{i'=1}^m \left\{ \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \right\}_{ii'} \varphi_a(\mathbf{s}_{ji}) \varphi_a(\mathbf{s}_{ji'}) \right] + \\ &\quad \lambda_n^{-1} \text{tr}(C_{\alpha}^d) - \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \sum_{a=d+1}^{\infty} \mu_a^2 \left[\sum_{i=1}^m \sum_{i'=1}^m \left\{ \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \right\}_{i'i''} \varphi_a(\mathbf{s}_{ji}) \varphi_a(\mathbf{s}_{ji'}) \right] \\ (45) \quad &\stackrel{(i)}{\leq} \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \sum_{a=1}^d \left\{ \mu_a - \mu_a^2 \varphi_a^{jT} \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \varphi_a^j \right\} + \lambda_n^{-1} \text{tr}(C_{\alpha}^d), \end{aligned}$$

where i ath element of the matrix Φ^j (defined in the proof of Lemma 1.1) is $\varphi_a(\mathbf{s}_{ji})$, φ_a^j is the a th column of Φ^j , and (i) follows because $\left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)$ is a positive definite matrix and $\varphi_a^{jT} \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \varphi_a^j \geq 0$.

Let $\mathbf{M} = \text{diag}(\mu_1, \dots, \mu_d)$ and $\mathbf{Q} = \left(\mathbf{I} + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} \right)^{1/2}$ as defined in the proofs of Lemmas 1.1 and 1.2. Define a $d \times d$ matrix $\mathbf{B} \equiv \mathbf{M} - \mathbf{M} \Phi^{jT} \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \Phi^j \mathbf{M}$, so that from (45),

$$\begin{aligned} \text{tr}(\mathbf{B}) &= \sum_{a=1}^d \left\{ \mu_a - \mu_a^2 \varphi_a^{jT} \left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \varphi_a^j \right\}, \\ (46) \quad \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{v_j(\mathbf{s}^*)\} &\leq \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \text{tr}(\mathbf{B}) + \lambda_n^{-1} \text{tr}(C_{\alpha}^d). \end{aligned}$$

Let

$$\begin{aligned} \mathbf{C}_{j,j} &= \Phi^j \mathbf{M} \Phi^{jT} + \Phi^{j\uparrow} \mathbf{M}^{\uparrow} \Phi^{j\uparrow T} \equiv \Phi^j \mathbf{M} \Phi^{jT} + \mathbf{C}_{j,j}^{\uparrow}, \\ \mathbf{M}^{\uparrow} &= \text{diag}(\mu_{d+1}, \dots, \mu_{\infty}), \quad \Phi^{j\uparrow} = [\varphi_{d+1}^j, \dots, \varphi_{\infty}^j], \end{aligned}$$

then the Woodbury formula (Harville, 1997) and the definition of $\mathbf{Q}$ imply that

$$\mathbf{B} = \left\{ \mathbf{M}^{-1} + \Phi^{jT} \left(\mathbf{C}_{j,j}^{\uparrow} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1} \Phi^j \right\}^{-1}$$

$$\begin{aligned}
&= \frac{\tau^2 \lambda_n}{km} \left\{ \mathbf{I} + \frac{\tau^2 \lambda_n}{km} \mathbf{M}^{-1} + \frac{1}{m} \Phi^{jT} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow + \mathbf{I} \right)^{-1} \Phi^j - \mathbf{I} \right\}^{-1} \\
(47) \quad &= \frac{\tau^2 \lambda_n}{km} \mathbf{Q}^{-2} \left[\mathbf{I} + \mathbf{Q}^{-1} \left\{ \frac{1}{m} \Phi^{jT} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow + \mathbf{I} \right)^{-1} \Phi^j - \mathbf{I} \right\} \mathbf{Q}^{-1} \right]^{-1}.
\end{aligned}$$

Define the event $\mathcal{E}_2 = \left\{ \frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow \preceq \frac{1}{4} \mathbf{I} \right\}$. Since the matrix $\mathbf{C}_{j,j}^\uparrow$ is nonnegative definite, we have the relation that

$$\left\{ \text{tr} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow \right) \leq \frac{1}{4} \right\} \subseteq \left\{ s_{\max} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow \right) \leq \frac{1}{4} \right\} \subseteq \mathcal{E}_2,$$

$s_{\max}(\mathbf{A})$ is the maximum eigenvalue of the square matrix $\mathbf{A}$. Therefore, by Markov's inequality, we have that

$$\begin{aligned}
\mathbb{P}(\mathcal{E}_2^c) &\leq \mathbb{P} \left\{ \text{tr} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow \right) > \frac{1}{4} \right\} \leq 4 \mathbb{E}_{\mathcal{S}} \text{tr} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow \right) \\
(48) \quad &= \frac{4k}{\tau^2 \lambda_n} \sum_{i=1}^m \sum_{a=d+1}^{\infty} \mu_a \mathbb{E}_{\mathcal{S}} \varphi_a^2(\mathbf{s}_{ji}) = \frac{4km}{\tau^2 \lambda_n} \text{tr} (C_{\alpha}^d).
\end{aligned}$$

Now on the event $\mathcal{E}_1 \cap \mathcal{E}_2$ (with $\mathcal{E}_1$ defined in (20)), we have that

$$\begin{aligned}
&\mathbf{I} + \mathbf{Q}^{-1} \left\{ \frac{1}{m} \Phi^{jT} \left(\frac{k}{\tau^2 \lambda_n} \mathbf{C}_{j,j}^\uparrow + \mathbf{I} \right)^{-1} \Phi^j - \mathbf{I} \right\} \mathbf{Q}^{-1} \\
&\stackrel{(i)}{\succeq} \mathbf{I} + \mathbf{Q}^{-1} \left\{ \frac{1}{m} \Phi^{jT} \left(\frac{1}{4} \mathbf{I} + \mathbf{I} \right)^{-1} \Phi^j - \mathbf{I} \right\} \mathbf{Q}^{-1} \\
&= \mathbf{I} - \frac{1}{5} \mathbf{Q}^{-2} + \frac{4}{5} \mathbf{Q}^{-1} \left\{ \frac{1}{m} \Phi^{jT} \Phi^j - \mathbf{I} \right\} \mathbf{Q}^{-1} \\
(49) \quad &\stackrel{(ii)}{\succeq} \mathbf{I} - \frac{1}{5} \mathbf{I} - \frac{4}{5} \cdot \frac{1}{2} \mathbf{I} = \frac{2}{5} \mathbf{I},
\end{aligned}$$

where (i) follows on the event $\mathcal{E}_2$, and (ii) holds on the event $\mathcal{E}_1$ and from the fact $\mathbf{Q}^{-2} \preceq \mathbf{I}$.

Therefore, by combining (48), (49), and the upper bound for $\mathbb{P}(\mathcal{E}_1^c)$ given in (22) under our assumptions, we obtain that

$$\begin{aligned}
&\mathbb{E}_{\mathcal{S}} \text{tr}(\mathbf{B}) \\
&\leq \mathbb{E}_{\mathcal{S}} \{ \text{tr}(\mathbf{B}) \mathbf{1}(\mathcal{E}_1 \cap \mathcal{E}_2) \} + \mathbb{E}_{\mathcal{S}} [\text{tr}(\mathbf{B}) \{ \mathbf{1}(\mathcal{E}_1^c) + \mathbf{1}(\mathcal{E}_2^c) \}] \\
&\stackrel{(i)}{\leq} \frac{5}{2} \frac{\tau^2 \lambda_n}{km} \text{tr}(\mathbf{Q}^{-2}) + \text{tr}(C_{\alpha}) \{ \mathbb{P}(\mathcal{E}_1^c) + \mathbb{P}(\mathcal{E}_2^c) \} \\
&\stackrel{(ii)}{\leq} 3 \frac{\tau^2 \lambda_n}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) + \frac{4n}{\tau^2 \lambda_n} \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \\
(50) \quad &+ \text{tr}(C_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma \left(\frac{\tau^2 \lambda_n}{n} \right)}{\sqrt{m}} \right\}^q,
\end{aligned}$$

where (i) follows from (49), and (ii) follows from (48), (22), and by replacing km with n .

(45), (47), and (50) together yield

$$\begin{aligned}
& \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \bar{v}_j(\mathbf{s}^*) \} \\
& \leq \lambda_n^{-1} \mathbb{E}_{\mathcal{S}} \operatorname{tr}(\mathbf{B}) + \lambda_n^{-1} \operatorname{tr} \left(C_{\alpha}^d \right) \\
& \leq 3 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) + \left\{ \frac{4n}{\tau^2 \lambda_n^2} \operatorname{tr}(C_{\alpha}) + \frac{1}{\lambda_n} \right\} \operatorname{tr}(C_{\alpha}^d) \\
(51) \quad & + \lambda_n^{-1} \operatorname{tr}(C_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma \left(\frac{\tau^2}{n} \right)}{\sqrt{m}} \right\}^q.
\end{aligned}$$

Since the righthand side of (51) does not depend on j , a further upper bound for (44) is given by

$$\begin{aligned}
& \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \bar{v}(\mathbf{s}^*) \} \leq \frac{1}{k} \sum_{j=1}^k \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \bar{v}_j(\mathbf{s}^*) \} \\
& \leq 3 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) + \left\{ \frac{4n}{\tau^2 \lambda_n^2} \operatorname{tr}(C_{\alpha}) + \frac{1}{\lambda_n} \right\} \operatorname{tr}(C_{\alpha}^d) \\
(52) \quad & + \lambda_n^{-1} \operatorname{tr}(C_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma \left(\frac{\tau^2}{n} \right)}{\sqrt{m}} \right\}^q.
\end{aligned}$$

Taking the infimum over $d \in \mathbb{N}$ leads to the proof. $\blacksquare$

1.2 Proof of Theorem 3.2

The proof of parts (i)–(iv) are as follows.

(i) Since d^* is a constant integer and $k = o(n)$, we can take m sufficiently large such that $n \geq m > \max(d^*, e^q)$. In the upper bounds of Theorem 3.1, we choose $d = n$ in every infimum to make the upper bounds larger. This implies that $\operatorname{tr}(C_{\alpha}^d) = 0$, $\mu_{d+1} = 0$, and $b(m, d, q) \leq \log n$. Also notice that in this case, $\gamma(a) \leq d^*$ for any $a > 0$. Then, with $\lambda_n = 1$, Theorem 3.1 implies that

$$\begin{aligned}
& \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{ \bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*) \}^2 \\
& \leq (8 \|w_0\|_{\mathbb{H}}^2 + 12k^{-1} \|w_0\|_{\mathbb{H}}^2 + 15d^*) \frac{\tau^2}{n} \\
& + \left\{ \mu_1 \|w_0\|_{\mathbb{H}}^2 + \frac{2n}{k} + \frac{4 \|w_0\|_{\mathbb{H}}^2}{k} + \operatorname{tr}(C_{\alpha}) \right\} \left(\frac{A \rho^2 d^* \log n}{\sqrt{n/k}} \right)^q \\
& \leq O(n^{-1}) + \{1 + o(1)\} \frac{2 (A \rho^2 d^* \log n)^q k^{r/2-1}}{n^{r/2-1}} \\
& = O(n^{-1}),
\end{aligned}$$

where the last equality follows from the condition on k .

(ii) In the upper bounds of Theorem 3.1, we choose $d = n^2$ in every infimum for sufficiently large n such that $\log d = 2 \log n > q$. Then

$$\mu_{d+1} \leq c_{1\mu} \exp(-c_{2\mu} n^{2\kappa}) = O(n^{-4}),$$

$$\begin{aligned}
b(m, d, q) &\leq \max \left(\sqrt{\log d}, \frac{\log d}{m^{1/2-1/q}} \right) \leq \log d \leq 2 \log n, \\
\text{tr} \left(C_{\alpha}^d \right) &= \sum_{i=n^2+1}^{\infty} \mu_i \leq \sum_{i=n^2+1}^{\infty} c_{1\mu} \exp(-c_{2\mu} i^{\kappa}) \leq c_{1\mu} \int_{n^2}^{\infty} \exp(-c_{2\mu} z^{\kappa}) dz \\
(53) \quad &= c_{1\mu} \int_{n^{2\kappa}}^{\infty} \frac{1}{\kappa} t^{\frac{1}{\kappa}-1} \exp(-c_{2\mu} t) dt,
\end{aligned}$$

where in the last step, we use the change of variable $t = z^{\kappa}$. If $\kappa \geq 1$, then since $t \geq n^{2\kappa} \geq 1$, we have $t^{\frac{1}{\kappa}-1} \leq 1$. If $0 < \kappa < 1$, then there exists a large $n_0 \in \mathbb{N}$ that depends on only $c_{2\mu}$ and κ , such that for all $n \geq n_0$ and $t \geq n^{2\kappa}$, we have $t^{\frac{1}{\kappa}-1} \leq \exp(c_{2\mu} t/2)$. Therefore, in all cases,

$$(54) \quad \text{tr} \left(C_{\alpha}^d \right) \leq \frac{c_{1\mu}}{\kappa} \int_{n^{2\kappa}}^{\infty} \exp(-c_{2\mu} t/2) dt = \frac{2c_{1\mu}}{c_{2\mu}\kappa} \exp(-c_{2\mu} n^{2\kappa}/2) = O(n^{-4}).$$

Let $d_1 = \left(\frac{2}{c_{2\mu}} \log n \right)^{1/\kappa}$. For sufficiently large n , with $\lambda_n \equiv 1$, $\gamma(\tau^2 \lambda_n/n)$ can be bounded as

$$\begin{aligned}
\gamma(\tau^2 \lambda_n/n) &= \gamma(\tau^2/n) = \sum_{i=1}^{\infty} \frac{\mu_i}{\mu_i + \frac{\tau^2}{n}} = \sum_{i=1}^{\lfloor d_1 \rfloor + 1} \frac{\mu_i}{\mu_i + \frac{\tau^2}{n}} + \sum_{i=\lfloor d_1 \rfloor + 2}^{\infty} \frac{\mu_i}{\mu_i + \frac{\tau^2}{n}} \\
&\leq d_1 + 1 + \frac{n}{\tau^2} \sum_{i=\lfloor d_1 \rfloor + 1}^{\infty} c_{1\mu} \exp(-c_{2\mu} i^{\kappa}) \\
&\leq d_1 + 1 + \frac{n}{\tau^2} \int_{d_1}^{\infty} c_{1\mu} \exp(-c_{2\mu} z^{\kappa}) dz \\
&= d_1 + 1 + \frac{nc_{1\mu}}{\tau^2 \kappa} \int_{d_1^{\kappa}}^{\infty} t^{\frac{1}{\kappa}-1} \exp(-c_{2\mu} t) dt \\
&\leq d_1 + 1 + \frac{nc_{1\mu}}{\tau^2 \kappa} \int_{d_1^{\kappa}}^{\infty} \exp(-c_{2\mu} t/2) dt \\
&= d_1 + 1 + \frac{nc_{1\mu}}{c_{2\mu} \tau^2 \kappa} \exp(-c_{2\mu} d_1^{\kappa}/2) \\
(55) \quad &= \left(\frac{2}{c_{2\mu}} \log n \right)^{1/\kappa} + 1 + \frac{c_{1\mu}}{c_{2\mu} \tau^2 \kappa} = O\left((\log n)^{1/\kappa}\right).
\end{aligned}$$

Therefore, from (53), (54), (55), and the bounds in Theorem 3.1, we obtain that

$$\begin{aligned}
&\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 \\
&\leq O(n^{-1}) + 15 \frac{\tau^2}{n} \gamma\left(\frac{\tau^2}{n}\right) + \{1 + o(1)\} \frac{2n}{k} \left\{ \frac{Ab(m, d, q) \rho^2 \gamma\left(\frac{\tau^2}{n}\right)}{\sqrt{m}} \right\}^q \\
&\leq O(n^{-1}) + O\left((\log n)^{1/\kappa}/n\right) + O(1) \cdot \frac{n}{k} \left\{ \frac{(\log n)^{1/\kappa} \cdot \log n}{\sqrt{n/k}} \right\}^q \\
&\leq O\left((\log n)^{1/\kappa}/n\right) + O(1) \cdot \frac{k^{\frac{q}{2}-1} (\log n)^{\frac{q(1+\kappa)}{\kappa}}}{n^{\frac{q}{2}-1}} \\
&= O\left((\log n)^{1/\kappa}/n\right),
\end{aligned}$$

where the last equality follows from the condition on k .

(iii) Let $\lambda_n = 1$. In the upper bounds of Theorem 3.1, we choose $d = \lfloor n^{3/(2\eta-1)} \rfloor$ in every infimum for sufficiently large n such that $\log d \geq \log \left(n^{\frac{3}{2\eta-1}} - 1 \right) > q$. Then

$$\begin{aligned}
 \mu_{d+1} &\leq c_\mu n^{-6\eta/(2\eta-1)} \leq c_\mu n^{-3}, \\
 \text{tr} \left(C_\alpha^d \right) &= \sum_{i=d+1}^{\infty} \mu_i \leq \sum_{i=d+1}^{\infty} c_\mu i^{-2\eta} \leq c_\mu \int_d^{\infty} \frac{1}{z^{2\eta}} dz \\
 &= \frac{c_\mu}{2\eta-1} d^{-(2\eta-1)} \leq \frac{c_\mu}{2\eta-1} n^{-3}, \\
 (56) \quad b(m, d, q) &\leq \max \left(\sqrt{\log d}, \frac{\log d}{m^{1/2-1/q}} \right) \leq \log d \leq \frac{3}{2\eta-1} \log n.
 \end{aligned}$$

$\gamma(\tau^2 \lambda_n / n) = \gamma(\tau^2 / n)$ can be bounded as

$$\begin{aligned}
 \gamma(\tau^2 / n) &= \sum_{i=1}^{\infty} \frac{1}{1 + \frac{\tau^2}{n\mu_i}} \leq \sum_{i=1}^{\infty} \frac{1}{1 + \frac{\tau^2 i^{2\eta}}{c_\mu n}} \\
 &\leq n^{1/(2\eta)} + 1 + \frac{c_\mu n}{\tau^2} \sum_{i=\lfloor n^{1/(2\eta)} \rfloor + 2}^{\infty} \frac{1}{i^{2\eta}} \\
 &\leq n^{1/(2\eta)} + 1 + \frac{c_\mu n}{\tau^2} \int_{n^{1/(2\eta)}}^{\infty} \frac{1}{z^{2\eta}} dz \\
 (57) \quad &= n^{1/(2\eta)} + 1 + \frac{c_\mu n}{\tau^2(2\eta-1)n^{(2\eta-1)/(2\eta)}} \leq \left(\frac{c_\mu}{\tau^2(2\eta-1)} + 1 \right) n^{1/(2\eta)}.
 \end{aligned}$$

From (56), (57), and the bounds in Theorem 3.1, we obtain that

$$\begin{aligned}
 &\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{ \bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*) \}^2 \\
 &\leq O(n^{-1}) + 15 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2}{n} \right) + \{1 + o(1)\} \frac{2n}{k} \left\{ \frac{Ab(m, d, q) \rho^2 \gamma \left(\frac{\tau^2}{n} \right)}{\sqrt{m}} \right\}^q \\
 &\leq O(n^{-1}) + \frac{15\tau^2 \left(2 + \frac{c_\mu}{\tau^2(2\eta-1)} \right) n^{1/(2\eta)}}{n} \\
 &\quad + \{1 + o(1)\} \frac{2n}{k} \left\{ \frac{3A\rho^2 \left(2 + \frac{c_\mu}{\tau^2(2\eta-1)} \right) n^{1/(2\eta)} \log n}{(2\eta-1) \sqrt{n/k}} \right\}^q \\
 &\leq O(n^{-1}) + O \left(n^{-\frac{2\eta-1}{2\eta}} \right) + O(1) \cdot \frac{k^{\frac{q}{2}-1} (\log n)^q}{n^{\frac{q}{2}-1-\frac{q}{2\eta}}} \\
 &= O \left(n^{-\frac{2\eta-1}{2\eta}} \right),
 \end{aligned}$$

where the last equality follows from the condition on k .

(iv) Now let $\lambda_n = c_1 n^{1/(2\eta+1)}$. In the upper bounds of Theorem 3.1, we choose $d = \lfloor n^{3/(2\eta-1)} \rfloor$ in every infimum for sufficiently large n , in the same way as in

Part (iii). Therefore, (56) still holds true. Furthermore, since $\lambda_n = c_1 n^{1/(2\eta+1)}$, we have that

$$\begin{aligned}
\gamma(\tau^2 \lambda_n / n) &= \sum_{i=1}^{\infty} \left(1 + \frac{\tau^2 \lambda_n}{n \mu_i} \right)^{-1} \leq \sum_{i=1}^{\infty} \left(1 + \frac{\tau^2 c_1 i^{2\eta}}{c_\mu n^{2\eta/(2\eta+1)}} \right)^{-1} \\
&\leq n^{1/(2\eta+1)} + 1 + \frac{c_\mu n^{\frac{2\eta}{2\eta+1}}}{\tau^2 c_1} \sum_{i=\lfloor n^{\frac{1}{2\eta+1}} \rfloor + 2}^{\infty} \frac{1}{i^{2\eta}} \\
&\leq n^{1/(2\eta+1)} + 1 + \frac{c_\mu n^{\frac{2\eta}{2\eta+1}}}{\tau^2 c_1} \int_{n^{\frac{1}{2\eta+1}}}^{\infty} \frac{1}{z^{2\eta}} dz \\
&= n^{1/(2\eta+1)} + 1 + \frac{c_\mu n^{\frac{2\eta}{2\eta+1}}}{\tau^2 c_1 (2\eta - 1) n^{(2\eta-1)/(2\eta+1)}} \\
(58) \quad &\leq \left(\frac{c_\mu}{\tau^2 (2\eta - 1)} + 1 \right) n^{\frac{1}{2\eta+1}}.
\end{aligned}$$

From (56), (58), and the bounds in Theorem 3.1, we obtain that

$$\begin{aligned}
&\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}, \bar{w}(\mathbf{s}^*) | \mathcal{S}} \{ \bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*) \}^2 \\
&\leq O(\lambda_n / n) + 15 \frac{\tau^2}{n} \gamma \left(\frac{\tau^2 \lambda_n}{n} \right) + \{1 + o(1)\} \frac{2n}{k \lambda_n} \left\{ \frac{Ab(m, d, q) \rho^2 \gamma \left(\frac{\tau^2 \lambda_n}{n} \right)}{\sqrt{m}} \right\}^q \\
&\leq O \left(n^{-\frac{2\eta}{2\eta+1}} \right) + 15 \tau^2 \left(\frac{c_\mu}{\tau^2 (2\eta - 1)} + 1 \right) n^{-\frac{2\eta}{2\eta+1}} \\
&\quad + \{1 + o(1)\} \frac{2n^{\frac{2\eta}{2\eta+1}}}{k} \left\{ \frac{3A \rho^2 \left(\frac{c_\mu}{\tau^2 (2\eta - 1)} + 1 \right) n^{\frac{1}{2\eta+1}} \log n}{(2\eta - 1) \sqrt{n/k}} \right\}^q \\
&\leq O(n^{-1}) + O \left(n^{-\frac{2\eta}{2\eta+1}} \right) + O(1) \cdot \frac{k^{\frac{q}{2}-1} (\log n)^q}{n^{\frac{(2\eta-1)q}{2(2\eta+1)} - \frac{2\eta}{2\eta+1}}} \\
&= O \left(n^{-\frac{2\eta}{2\eta+1}} \right),
\end{aligned}$$

where the last equality follows from the condition on k .

1.3 Extension to Unknown τ^2

In this section, we extend the convergence rates of Bayes L_2 -risk in Theorem 3.2 to the case where the covariance function is parameterized in a different way and is scaled by τ^2 , such that τ^2 is unknown and assigned a prior distribution. We modify the GP prior on $w(\cdot)$ in Equation (11) of the main text to the following

$$\begin{aligned}
y(\mathbf{s}_i) &= w(\mathbf{s}_i) + \epsilon(\mathbf{s}_i), \quad \epsilon(\mathbf{s}_i) \sim N(0, \tau^2), \\
(59) \quad w(\cdot) &\sim \text{GP}\{0, \lambda_n^{-1} \tau^2 C_\alpha(\cdot, \cdot)\};
\end{aligned}$$

that is, C_α is scaled with τ^2 , the same as the error variance. This parameterization has also been used in the application of GP models before. We maintain the same eigen-decomposition of the kernel $C_{\alpha_0}(\cdot, \cdot)$ and the Assumptions A.3 and A.4 as before. We assume that α is still fixed at its truth α_0 , but now impose a prior on τ^2 .

A.5' (Prior) For each of the k subsets, τ^2 is assigned a prior with a bounded support in $(0, \bar{\tau}^2]$ for some finite constants $\bar{\tau}^2 > 0$.

Let $\mathbb{E}_{\tau^2|\mathbf{y}}$ and $\mathbb{E}_{\bar{w}(\mathbf{s}^*)|\tau^2, \mathbf{y}, \mathbf{s}^*}$ be the expectations of $\{\tau_j^2 : j = 1, \dots, k\}$ given $\mathbf{y}$, and $\bar{w}(\mathbf{s}^*)$ given $\mathbf{y}$, $\{\tau_j^2 : j = 1, \dots, k\}$, and $\mathbf{s}^*$, respectively, where τ_j^2 is drawn from the posterior of τ^2 given $\mathbf{y}_j$ from the j th subset posterior. Then the Bayes L_2 -risk of the DISK posterior for $\bar{w}(\cdot)$ can be written as

$$(60) \quad \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2, \mathbf{s}^*} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2.$$

Then, we have the following corollary when a prior distribution is imposed on τ^2 .

Corollary 1.1 *If Assumptions A.1 – A.4 and A.5' hold, then all the convergence rates in the four cases of Theorem 3.2 still hold true for the Bayes L_2 -risk given in (62).*

Proof [Proof of Corollary 1.1] We proceed to prove a similar bound for the Bayes L_2 risk to Theorem 3.1 in the main paper under A.5'. By A.5', we need to account for the randomness in the posterior of $p(\tau_j^2|\mathbf{y}_j)$ across $j = 1, \dots, k$. Based on the model (59), we can see that conditional on the subset posterior draws of τ_j^2 from the subset posterior $p(\tau^2|\mathbf{y}_j)$ for $j = 1, \dots, k$, the DISK posterior draw $\bar{w}(\mathbf{s}^*)$ follows the distribution $N(\bar{m}, \bar{v})$, with

$$(61) \quad \begin{aligned} \bar{m} &= \frac{1}{k} \sum_{j=1}^k \mathbf{c}_{j,*}^T \{ \mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I} \}^{-1} \mathbf{y}_j, \\ \bar{v}^{1/2} &= \frac{1}{k} \sum_{j=1}^k v_j^{1/2}, \quad v_j = \frac{\tau_j^2}{\lambda_n} \left\{ c_{*,*} - \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I})^{-1} \mathbf{c}_{j,*} \right\}, \end{aligned}$$

where $\mathbf{c}_{j,*}$, $\mathbf{C}_{j,j}$, $c_{*,*}$ are defined similarly to those in (3) according to the base kernel C_{α_0} . Notice that $\bar{m}$ does not depend on τ_j^2 due to the rescaled kernel $\tau^2 C_{\alpha_0}$ in (59).

Let $\mathbb{E}_{\mathbf{s}^*}$, $\mathbb{E}_0$, $\mathbb{E}_{\mathcal{S}}$, $\mathbb{E}_{\mathbf{y}|\mathcal{S}}$, and $\mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2}$, $\mathbb{E}_{\tau^2|\mathbf{y}}$ respectively be the expectations with respect to the distributions of $\mathbf{s}^*$, $(\mathcal{S}, \mathbf{y})$, $\mathcal{S}$, $\mathbf{y}$ given $\mathcal{S}$, $\bar{w}(\mathbf{s}^*)$ given $\mathbf{y}$ and $\{\tau_j^2 : j = 1, \dots, k\}$, and $\{\tau_j^2 : j = 1, \dots, k\}$ given $\mathbf{y}$. Then based on A.5', the Bayes L_2 -risk of the DISK posterior for $\bar{w}(\cdot)$ can be written as

$$(62) \quad \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2.$$

To upper bound (62), we apply the law of total variance repeatedly to obtain that

$$(63) \quad \begin{aligned} & \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*)\}^2 \\ &= \left[\mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*)\} - w_0(\mathbf{s}^*) \right]^2 + \text{var}_{\mathbf{y}, \tau^2, \bar{w}(\mathbf{s}^*)|\mathcal{S}} \{\bar{w}(\mathbf{s}^*)\} \\ &= \left[\mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*)\} - w_0(\mathbf{s}^*) \right]^2 \\ & \quad + \text{var}_{\mathbf{y}|\mathcal{S}} \left[\mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*)\} \right] \\ & \quad + \mathbb{E}_{\mathbf{y}|\mathcal{S}} \left(\text{var}_{\tau^2|\mathbf{y}} \left[\mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*)\} \right] \right) \\ & \quad + \mathbb{E}_{\mathbf{y}|\mathcal{S}} \left(\mathbb{E}_{\tau^2|\mathbf{y}} \left[\text{var}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{\bar{w}(\mathbf{s}^*)\} \right] \right). \end{aligned}$$

Using (61), we can derive that

$$\begin{aligned}
& \left[\mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y},\tau^2} \{ \bar{w}(\mathbf{s}^*) \} - w_0(\mathbf{s}^*) \right]^2 \\
&= \left[k^{-1} \sum_{j=1}^k \mathbf{c}_{j,*}^T \{ \mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I} \}^{-1} \mathbf{w}_{0j} - w_0(\mathbf{s}^*) \right]^2, \\
(64) \quad &= \{ \mathbf{c}_*^T (k \mathbf{L} + \lambda_n \mathbf{I})^{-1} \mathbf{w}_0 - w_0(\mathbf{s}^*) \}^2, \\
&\quad \text{var}_{\mathbf{y}|\mathcal{S}} \left[\mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y},\tau^2} \{ \bar{w}(\mathbf{s}^*) \} \right] \\
&= \text{var}_{\mathbf{y}|\mathcal{S}} \left[k^{-1} \sum_{j=1}^k \mathbf{c}_{j,*}^T \{ \mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I} \}^{-1} \mathbf{y}_j \right] \\
(65) \quad &= \tau_0^2 \mathbf{c}_*^T (\mathbf{s}^*) (k \mathbf{L} + \lambda_n \mathbf{I})^{-2} \mathbf{c}_*(\mathbf{s}^*), \\
&\quad \mathbb{E}_{\mathbf{y}|\mathcal{S}} \left(\text{var}_{\tau^2|\mathbf{y}} \left[\mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y},\tau^2} \{ \bar{w}(\mathbf{s}^*) \} \right] \right) \\
(66) \quad &= \mathbb{E}_{\mathbf{y}|\mathcal{S}} \left(\text{var}_{\tau^2|\mathbf{y}} \left[\frac{1}{k} \sum_{j=1}^k \mathbf{c}_{j,*}^T \{ \mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I} \}^{-1} \mathbf{y}_j \right] \right) = 0, \\
(67) \quad &\mathbb{E}_{\mathbf{y}|\mathcal{S}} \left(\mathbb{E}_{\tau^2|\mathbf{y}} \left[\text{var}_{\bar{w}(\mathbf{s}^*)|\mathbf{y},\tau^2} \{ \bar{w}(\mathbf{s}^*) \} \right] \right) = \mathbb{E}_{\mathbf{y}|\mathcal{S}} \{ \mathbb{E}_{\tau^2|\mathbf{y}}(\bar{v}) \},
\end{aligned}$$

where (64) and (67) follow from (61) and (5), (65) follows similarly to (5), and (66) is zero because $\bar{m}$ does not depend on τ_j^2 ($j = 1, \dots, k$). Next, we find upper bound for (64), (65), and (67), respectively.

First, we notice that (64) has the same expression as (5) by setting $\tau^2 = 1$ in (5). Therefore, the proof and the conclusion of Lemma 1.1 still works as before, by setting $\tau^2 = 1$, i.e.

$$\begin{aligned}
& \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \lambda_n \mathbf{I})^{-1} \mathbf{w}_0 - w_0(\mathbf{s}^*) \}^2 \leq \frac{8\lambda_n}{n} \|w_0\|_{\mathbb{H}}^2 \\
(68) \quad &+ \|w_0\|_{\mathbb{H}}^2 \inf_{d \in \mathbb{N}} \left[\frac{8n}{\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) + \mu_1 \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\lambda_n}{n})}{\sqrt{m}} \right\}^q \right].
\end{aligned}$$

Second, we notice that (65) differs from (5) only with the τ^2 outside replaced by the true error variance τ_0^2 , and that $\tau^2 = 1$ in $(k \mathbf{L} + \tau^2 \lambda_n \mathbf{I})^{-2}$. We carefully inspect and modify the proof of Lemma 1.2 to obtain that

$$\begin{aligned}
& \tau_0^2 \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \{ \mathbf{c}_*^T (k \mathbf{L} + \lambda_n \mathbf{I})^{-2} \mathbf{c}_* \} \leq \\
& \left(\frac{2\tau_0^2 n}{k \lambda_n} + \frac{4\|w_0\|_{\mathbb{H}}^2}{k} \right) \inf_{d \in \mathbb{N}} \left[\mu_{d+1} + 12 \frac{n}{\lambda_n} \rho^4 \text{tr}(C_{\alpha}) \text{tr}(C_{\alpha}^d) \right. \\
(69) \quad & \left. + \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\lambda_n}{n})}{\sqrt{m}} \right\}^q \right] + \frac{12\lambda_n}{kn} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau_0^2 \lambda_n}{n} \gamma \left(\frac{\lambda_n}{n} \right).
\end{aligned}$$

Third, $\bar{v}$ (and v_j) in (61) differs from $\bar{v}$ (and v_j) in (3) only in that (61) has a τ_j^2 factor outside and it has $\tau^2 = 1$ inside $\left(\mathbf{C}_{j,j} + \frac{\tau^2 \lambda_n}{k} \mathbf{I} \right)^{-1}$ in (3). Using the expression of v_j in (61) and the upper bound $\tau_j^2 \leq \bar{\tau}^2$ in A.5', we carefully inspect and modify the proof of Lemma 1.2 to obtain that

$$\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}}(\bar{v})$$

$$\begin{aligned}
&\leq \frac{1}{k} \sum_{j=1}^k \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}}(v_j) \\
&\leq \frac{\bar{\tau}^2 \lambda_n^{-1}}{k} \sum_{j=1}^k \mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \left\{ \mathbf{c}_{*,*} - \mathbf{c}_{j,*}^T (\mathbf{C}_{j,j} + \frac{\lambda_n}{k} \mathbf{I})^{-1} \mathbf{c}_{j,*} \right\} \\
&\leq \bar{\tau}^2 \left\{ \frac{3}{n} \gamma \left(\frac{\lambda_n}{n} \right) + \inf_{d \in \mathbb{N}} \left[\left\{ \frac{4n}{\lambda_n^2} \text{tr}(\mathbf{C}_{\alpha}) + \frac{1}{\lambda_n} \right\} \text{tr}(\mathbf{C}_{\alpha}^d) \right. \right. \\
(70) \quad &\quad \left. \left. + \lambda_n^{-1} \text{tr}(\mathbf{C}_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{1}{n})}{\sqrt{m}} \right\}^q \right] \right\}.
\end{aligned}$$

Now we can combine (62), (63), (64), (65), (66), (67), (68), (69), and (70) to obtain that

$$\begin{aligned}
&\mathbb{E}_{\mathbf{s}^*} \mathbb{E}_{\mathcal{S}} \mathbb{E}_{\mathbf{y}|\mathcal{S}} \mathbb{E}_{\tau^2|\mathbf{y}} \mathbb{E}_{\bar{w}(\mathbf{s}^*)|\mathbf{y}, \tau^2} \{ \bar{w}(\mathbf{s}^*) - w_0(\mathbf{s}^*) \}^2 \\
&\leq \frac{8\lambda_n}{n} \|w_0\|_{\mathbb{H}}^2 + \|w_0\|_{\mathbb{H}}^2 \inf_{d \in \mathbb{N}} \left[\frac{8n}{\lambda_n} \rho^4 \text{tr}(\mathbf{C}_{\alpha}) \text{tr}(\mathbf{C}_{\alpha}^d) + \mu_1 \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\lambda_n}{n})}{\sqrt{m}} \right\}^q \right] \\
&\quad + \left(\frac{2\tau_0^2 n}{k\lambda_n} + \frac{4\|w_0\|_{\mathbb{H}}^2}{k} \right) \inf_{d \in \mathbb{N}} \left[\mu_{d+1} + 12 \frac{n}{\tau^2 \lambda_n} \rho^4 \text{tr}(\mathbf{C}_{\alpha}) \text{tr}(\mathbf{C}_{\alpha}^d) \right. \\
&\quad \left. + \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{\lambda_n}{n})}{\sqrt{m}} \right\}^q \right] + \frac{12\lambda_n}{kn} \|w_0\|_{\mathbb{H}}^2 + 12 \frac{\tau_0^2 \lambda_n}{n} \gamma \left(\frac{\lambda_n}{n} \right) \\
&\quad + \bar{\tau}^2 \left\{ \frac{3}{n} \gamma \left(\frac{\lambda_n}{n} \right) + \inf_{d \in \mathbb{N}} \left[\left\{ \frac{4n}{\lambda_n^2} \text{tr}(\mathbf{C}_{\alpha}) + \frac{1}{\lambda_n} \right\} \text{tr}(\mathbf{C}_{\alpha}^d) \right. \right. \\
(71) \quad &\quad \left. \left. + \lambda_n^{-1} \text{tr}(\mathbf{C}_{\alpha}) \left\{ \frac{Ab(m, d, q) \rho^2 \gamma(\frac{1}{n})}{\sqrt{m}} \right\}^q \right] \right\}.
\end{aligned}$$

We notice that the upper bound in (71) differs from the upper bound in Theorem 3.1 only by some multiplicative constants in each term and inside the $\gamma(\cdot)$ functions. In the previous proof of Theorem 3.2, these constants will only change the multiplicative constants and do not affect the convergence rates of the Bayes L_2 -risk of $\bar{w}(\cdot)$. As a result, the convergence rate results of Theorem 3.2 continue to hold for (71) under the various conditions specified in the different cases of Theorem 3.2. This completes the proof. $\blacksquare$

2. SAMPLING FROM THE SUBSET POSTERIOR DISTRIBUTIONS USING A FULL-RANK GP PRIOR

Recall the univariate spatial regression model for the data observed at the i th location in subset j using a GP prior is

$$(72) \quad y(\mathbf{s}_{ji}) = \mathbf{x}(\mathbf{s}_{ji})^T \boldsymbol{\beta} + w(\mathbf{s}_{ji}) + \epsilon(\mathbf{s}_{ji}), \quad j = 1, \dots, k, \quad i = 1, \dots, m_j.$$

For the simulations and real data analysis, we assume that $C_{\alpha}(\mathbf{s}_{ji}, \mathbf{s}_{ji'}) = \sigma^2 \rho(\mathbf{s}_{ji}, \mathbf{s}_{ji'}; \phi)$ and $D_{\alpha}(\mathbf{s}_{ji}, \mathbf{s}_{ji'}) = \mathbf{1}(i = i') \tau^2$, where σ^2, ϕ, τ^2 are positive scalars, $\rho(\cdot, \cdot)$ is a

known positive definite correlation function, and $\mathbf{1}(i = i') = 1$ if $i = i'$ and 0 otherwise. This implies that $\boldsymbol{\alpha} = (\sigma^2, \tau^2, \phi)$. The model in (72) is completed by putting priors on the unknown parameters. The priors distributions on $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$ have the following forms:

$$(73) \quad \boldsymbol{\beta} \sim N(\boldsymbol{\mu}_\beta, \Sigma_\beta), \quad \sigma^2 \sim \text{IG}(a_\sigma, b_\sigma), \quad \tau^2 \sim \text{IG}(a_\tau, b_\tau), \quad \phi \sim \text{U}(a_\phi, b_\phi),$$

where $\boldsymbol{\mu}_\beta, \Sigma_\beta, a_\sigma, b_\sigma, a_\tau, b_\tau, a_\phi$, and b_ϕ are constants, N represents the multi-variate Gaussian distribution of appropriate dimension, $\text{IG}(a, b)$ represents the Inverse-Gamma distribution with mean $a/(b+1)$ and variance $b/\{(a-1)^2(a-2)\}$ for $a > 2$, and $\text{U}(a, b)$ represents the uniform distribution on the interval $[a, b]$. The spatial process $w(\cdot)$ is assigned a GP prior as

$$(74) \quad w(\cdot) \mid \sigma^2, \phi \sim \text{GP}\{0, C_\alpha(\cdot, \cdot)\}, \quad C_\alpha(\cdot, \cdot) = \sigma^2 \rho(\cdot, \cdot; \phi).$$

The training data $\{\mathbf{x}(\mathbf{s}_{j1}), y(\mathbf{s}_{j1})\}, \dots, \{\mathbf{x}(\mathbf{s}_{jm_j}), y(\mathbf{s}_{jm_j})\}$ are observed at the m_j spatial locations and $\mathcal{S}_j = \{\mathbf{s}_{j1}, \dots, \mathbf{s}_{jm_j}\}$ contains the locations in subset j .

Consider the setup for predictions and inferences on subset j . Let $\mathcal{S}^* = \{\mathbf{s}_1^*, \dots, \mathbf{s}_l^*\}$ be the set of locations such that $\mathcal{S}^* \cap \mathcal{S}_j = \emptyset$. If $\mathbf{w}_j^T = \{w(\mathbf{s}_{j1}), \dots, w(\mathbf{s}_{jm_j})\}$ and $\boldsymbol{\epsilon}_j^T = \{\epsilon(\mathbf{s}_{j1}), \dots, \epsilon(\mathbf{s}_{jm_j})\}$, then (72) implies that $\mathbf{w}_j$ apriori follows $N\{\mathbf{0}, \mathbf{C}_{j,j}(\boldsymbol{\alpha})\}$, where $\mathbf{C}_{j,j}(\boldsymbol{\alpha})$ is the block of $\mathbf{C}(\boldsymbol{\alpha})$ that corresponds to the locations in $\mathcal{S}_j$, and $\boldsymbol{\epsilon}_j$ follows $N(\mathbf{0}, \tau^2 \mathbf{I})$, where $\mathbf{I}$ is the identity matrix of appropriate dimension. Given the training data on subset j , our goal is to predict $\mathbf{y}_j^* = \{y(\mathbf{s}_1^*), \dots, y(\mathbf{s}_l^*)\}$ and to perform posterior inference on $\mathbf{w}_j^* = \{w(\mathbf{s}_1), \dots, w(\mathbf{s}_l)\}$, $\boldsymbol{\beta}_j$, and $\boldsymbol{\alpha}_j$, where the subscript j denotes that the predictions and inferences condition only on subset j . Standard Markov chain Monte Carlo (MCMC) algorithms exist to achieve this goal (Banerjee et al., 2014), but conditioning only on subset j ignores the information contained in the other $(k-1)$ subsets, resulting in greater posterior uncertainty compared to the full data posterior distribution.

Stochastic approximation is an approach for proper uncertainty quantification that modifies the likelihood used for sampling from the subset posterior distributions for predictions and inferences. The likelihoods for $\boldsymbol{\beta}$, $\boldsymbol{\alpha}$, and $\mathbf{w}_j$ are raised to the power of k to compensate for the data in the other $(k-1)$ subsets, where we assume that $m_1 = \dots = m_k = m$ and $k = n/m$. First, consider stochastic approximation for the likelihood of $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$. Integrating out $\mathbf{w}_j$ in (72) gives

$$(75) \quad \mathbf{y}_j = \mathbf{X}_j \boldsymbol{\beta} + \boldsymbol{\eta}_j, \quad \boldsymbol{\eta}_j \sim N\{\mathbf{0}, \mathbf{C}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\},$$

where $\mathbf{X}_j = [\mathbf{x}(\mathbf{s}_{j1}) : \dots : \mathbf{x}(\mathbf{s}_{jm_j})]^T \in \mathbb{R}^{m \times p}$ is the design matrix for subset j . The likelihood of $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$ given $\mathbf{y}_j$, $\mathbf{X}_j$ after stochastic approximation is

$$(76) \quad \{l_j(\boldsymbol{\beta}, \boldsymbol{\alpha})\}^k = (2\pi)^{-mk/2} |\mathbf{C}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}|^{-k/2} e^{-\frac{k}{2}(\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})^T \{\mathbf{C}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}^{-1} (\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})}.$$

The prior distribution for $\boldsymbol{\beta}$ in (73), the pseudo likelihood in (76), and Bayes rule implies that the density of the j th subset posterior distribution for $\boldsymbol{\beta}$ given the rest is

$$\boldsymbol{\beta} \mid \text{rest} \propto e^{-\frac{1}{2}(\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})^T [k^{-1} \{\mathbf{C}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}]^{-1} (\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})} e^{-\frac{1}{2}(\boldsymbol{\beta} - \boldsymbol{\mu}_\beta)^T \Sigma_\beta^{-1} (\boldsymbol{\beta} - \boldsymbol{\mu}_\beta)}.$$

This implies that the complete conditional distribution of β_j has density $N(\mathbf{m}_{j\beta}, \mathbf{V}_{j\beta})$, where

$$(77) \quad \begin{aligned} \mathbf{V}_{j\beta} &= \left[k \mathbf{X}_j^T \{ \mathbf{C}_{j,j}(\alpha) + \tau^2 \mathbf{I} \}^{-1} \mathbf{X}_j + \Sigma_\beta^{-1} \right]^{-1}, \\ \mathbf{m}_{j\beta} &= \mathbf{V}_{j\beta} \left[k \mathbf{X}_j^T \{ \mathbf{C}_{j,j}(\alpha) + \tau^2 \mathbf{I} \}^{-1} \mathbf{y}_j + \Sigma_\beta^{-1} \boldsymbol{\mu}_\beta \right]. \end{aligned}$$

If the density of the prior distribution for α is assumed to be $\pi(\sigma^2)\pi(\tau^2)\pi(\phi)$, where the prior densities $\pi(\sigma^2)$, $\pi(\tau^2)$, and $\pi(\phi)$ are defined in (73), then the pseudo likelihood in (76), and Bayes rule implies that the density of the j th subset posterior distribution for α given the rest is

$$(78) \quad \begin{aligned} \alpha \mid \text{rest} &\propto | \mathbf{C}_{j,j}(\alpha) + \tau^2 \mathbf{I} |^{-k/2} e^{-\frac{1}{2}(\mathbf{y}_j - \mathbf{X}_j \beta)^T [k^{-1} \{ \mathbf{C}_{j,j}(\alpha) + \tau^2 \mathbf{I} \}]^{-1} (\mathbf{y}_j - \mathbf{X}_j \beta)} \\ &(\sigma^2)^{-a_\sigma - 1} e^{-b_\sigma / \sigma^2} (\tau^2)^{-a_\tau - 1} e^{-b_\tau / \tau^2} (b_\phi - a_\phi)^{-1}. \end{aligned}$$

This density does not have a standard form, so we use a Metropolis-Hastings step with a normal random walk proposal and sample α_j using the `metrop` function in the R package `mcmc` (R Development Core Team, 2017).

Second, we derive the posterior predictive distribution of $\mathbf{w}_j^*$ given the rest. The GP prior on $(\mathbf{w}_j, \mathbf{w}_j^*)$ implies that the density of $\mathbf{w}_j^*$ given $\mathbf{w}_j$ is

$$(79) \quad \mathbf{w}_j^* \mid \mathbf{w}_j \sim N \left\{ \mathbf{C}_{*,j}(\alpha) \mathbf{C}_{j,j}^{-1}(\alpha) \mathbf{w}_j, \mathbf{C}_{*,*}(\alpha) - \mathbf{C}_{*,j}(\alpha) \mathbf{C}_{j,j}^{-1}(\alpha) \mathbf{C}_{j,*}(\alpha) \right\},$$

where $\text{cov}(\mathbf{w}_j^*, \mathbf{w}_j^*) = \mathbf{C}_{*,*}(\alpha)$, $\text{cov}(\mathbf{w}_j^*, \mathbf{w}_j) = \mathbf{C}_{*,j}(\alpha)$, and $\text{cov}(\mathbf{w}_j, \mathbf{w}_j^*) = \mathbf{C}_{j,*}(\alpha)$. Given α , β , $\mathbf{y}_j$, and $\mathbf{X}_j$, (72) implies that the likelihood of $\mathbf{w}_j$ after stochastic approximation is

$$(80) \quad \{l_j(\mathbf{w}_j)\}^k = (2\pi)^{-mk/2} |\tau^2 \mathbf{I}|^{-k/2} e^{-\frac{k}{2\tau^2}(\mathbf{y}_j - \mathbf{X}_j \beta - \mathbf{w}_j)^T (\mathbf{y}_j - \mathbf{X}_j \beta - \mathbf{w}_j)}.$$

The GP prior on $\mathbf{w}_j$, the pseudo likelihood in (80), and Bayes rule implies that the density of the subset posterior distribution for $\mathbf{w}_j$ given the rest is

$$\mathbf{w}_j \mid \text{rest} \propto e^{-\frac{1}{2\tau^2/k}(\mathbf{y}_j - \mathbf{X}_j \beta - \mathbf{w}_j)^T (\mathbf{y}_j - \mathbf{X}_j \beta - \mathbf{w}_j)} e^{-\frac{1}{2} \mathbf{w}_j^T \mathbf{C}_{j,j}^{-1}(\alpha) \mathbf{w}_j}.$$

This implies that the complete conditional distribution of $\mathbf{w}_j$ has density $N(\mathbf{m}_{\mathbf{w}_j}, \mathbf{V}_{\mathbf{w}_j})$, where

$$(81) \quad \mathbf{V}_{\mathbf{w}_j} = \left\{ \mathbf{C}_{j,j}^{-1}(\alpha) + \frac{k}{\tau^2} \mathbf{I} \right\}^{-1}, \quad \mathbf{m}_{\mathbf{w}_j} = \frac{k}{\tau^2} \mathbf{V}_{\mathbf{w}_j} (\mathbf{y}_j - \mathbf{X}_j \beta);$$

therefore, (79) and (81) imply that the complete conditional distribution of $\mathbf{w}_j^*$ has density $N(\mathbf{m}_{\mathbf{w}_j^*}, \mathbf{V}_{\mathbf{w}_j^*})$, where

$$(82) \quad \begin{aligned} \mathbf{m}_{\mathbf{w}_j^*} &= \mathbb{E}(\mathbf{w}_j^* \mid \text{rest}) = \mathbf{C}_{*,j}(\alpha) \mathbf{C}_{j,j}^{-1}(\alpha) \mathbb{E}(\mathbf{w}_j \mid \text{rest}) \\ &= \mathbf{C}_{*,j}(\alpha) \left\{ \mathbf{C}_{j,j}(\alpha) + \frac{\tau^2}{k} \mathbf{I} \right\}^{-1} (\mathbf{y}_j - \mathbf{X}_j \beta) \end{aligned}$$

and

$$\mathbf{V}_{\mathbf{w}_j^*} = \text{var}(\mathbf{w}_j^* \mid \text{rest}) = \mathbb{E} \{ \text{var}(\mathbf{w}_j^* \mid \mathbf{w}_j) \mid \text{rest} \} + \text{var} \{ \mathbb{E}(\mathbf{w}_j^* \mid \mathbf{w}_j) \mid \text{rest} \}$$

$$(83) \quad = \mathbf{C}_{*,*}(\boldsymbol{\alpha}) - \mathbf{C}_{*,j}(\boldsymbol{\alpha}) \mathbf{C}_{j,j}^{-1}(\boldsymbol{\alpha}) \mathbf{C}_{j,*}(\boldsymbol{\alpha}) + \mathbf{C}_{*,j}(\boldsymbol{\alpha}) \mathbf{C}_{j,j}^{-1}(\boldsymbol{\alpha}) \mathbf{V}_{\mathbf{w}_j} \mathbf{C}_{j,j}^{-1}(\boldsymbol{\alpha}) \mathbf{C}_{j,*}(\boldsymbol{\alpha}).$$

Finally, we derive the posterior predictive distribution of $\mathbf{y}_j^*$ given the rest. If $\boldsymbol{\beta}_j$, τ_j^2 , $\mathbf{w}_j^*$ are the samples from the j th subset posterior distribution of $\boldsymbol{\beta}$, τ^2 , and $\mathbf{w}^*$, then (72) implies that $\mathbf{y}_j^*$ given the rest is sampled as

$$\mathbf{y}_j^* = \mathbf{X}_j \boldsymbol{\beta}_j + \mathbf{w}_j^* + \boldsymbol{\epsilon}_j^*, \quad \boldsymbol{\epsilon}_j^* \sim N(\mathbf{0}, \tau_j^2 \mathbf{I});$$

therefore, the complete conditional distribution of $\mathbf{y}_j^*$ is $N(\boldsymbol{\mu}_{\mathbf{y}_j^*}, \mathbf{V}_{\mathbf{y}_j^*})$, where

$$(84) \quad \boldsymbol{\mu}_{\mathbf{y}_j^*} = \mathbf{X}_j \boldsymbol{\beta}_j + \mathbf{w}_j^*, \quad \mathbf{V}_{\mathbf{y}_j^*} = \tau_j^2 \mathbf{I}.$$

All full conditionals except that of $\boldsymbol{\alpha}$ are analytically tractable in terms of standard distributions in subset j ($j = 1, \dots, k$). The Gibbs sampler with a Metropolis-Hastings step iterates between the following four steps until sufficient number of samples of $\boldsymbol{\beta}_j$, $\boldsymbol{\alpha}_j$, $\mathbf{w}_j^*$, and $\mathbf{y}_j^*$ are drawn post convergence to the stationary distribution:

1. Sample $\boldsymbol{\beta}_j$ from $N(\boldsymbol{\mu}_{j\boldsymbol{\beta}}, \mathbf{V}_{j\boldsymbol{\beta}})$, where $\boldsymbol{\mu}_{j\boldsymbol{\beta}}$ and $\mathbf{V}_{j\boldsymbol{\beta}}$ are defined in (77).
2. Sample $\boldsymbol{\alpha}_j$ using the Metropolis-Hastings algorithm from the j th subset posterior density (up to constants) of $\boldsymbol{\alpha}_j$ in (78) with a normal random walk proposal.
3. Sample $\mathbf{w}_j^*$ from $N(\boldsymbol{\mu}_{\mathbf{w}_j^*}, \mathbf{V}_{\mathbf{w}_j^*})$, where $\boldsymbol{\mu}_{\mathbf{w}_j^*}$ and $\mathbf{V}_{\mathbf{w}_j^*}$ are defined in (82) and (83).
4. Sample $\mathbf{y}_j^*$ from $N(\boldsymbol{\mu}_{\mathbf{y}_j^*}, \mathbf{V}_{\mathbf{y}_j^*})$, where $\boldsymbol{\mu}_{\mathbf{y}_j^*}$ and $\mathbf{V}_{\mathbf{y}_j^*}$ are defined in (84).

3. SAMPLING FROM THE SUBSET POSTERIOR DISTRIBUTIONS USING A LOW-RANK GP PRIOR

For clarity, we focus on the modified predictive process (MPP) prior as a representative example of low-rank GP prior. The Gibbs sampling algorithm derived in this section is easily extended to other low-rank GP priors. Following the setup in Section 2, we assume that $C_{\boldsymbol{\alpha}}(\mathbf{s}_{ji}, \mathbf{s}_{ji'}) = \sigma^2 \rho(\mathbf{s}_{ji}, \mathbf{s}_{ji'}; \phi)$ and $D_{\boldsymbol{\alpha}}(\mathbf{s}_{ji}, \mathbf{s}_{ji'}) = \mathbf{1}(i = i')\tau^2$, $\boldsymbol{\alpha} = (\sigma^2, \tau^2, \phi)$, the prior distributions on $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$ have the same forms as in (73), and $\mathcal{S}_j$ contains the locations in subset j . Following the previous section, we assume that $m_1 = \dots = m_k = m$ and $k = n/m$. The only change in this section is that the spatial process $w(\cdot)$ in (72) is assigned a MPP prior derived from parent GP prior in (74). MPP projects the parent GP $w(\cdot)$ onto a subspace spanned by its realization over a set of r locations, $\mathcal{S}^{(0)} = \{\mathbf{s}_1^{(0)}, \dots, \mathbf{s}_r^{(0)}\}$, known as the “knots”, where no conditions are imposed on $\mathcal{S} \cap \mathcal{S}^{(0)}$. Let $\mathbf{c}(\cdot, \mathcal{S}^{(0)}) = \{C_{\boldsymbol{\alpha}}(\cdot, \mathbf{s}_1^{(0)}), \dots, C_{\boldsymbol{\alpha}}(\cdot, \mathbf{s}_r^{(0)})\}^T$ and $\mathbf{w}^{(0)} = \{w(\mathbf{s}_1^{(0)}), \dots, w(\mathbf{s}_r^{(0)})\}^T$ be $r \times 1$ vectors and $\mathbf{C}(\mathcal{S}^{(0)})$ be an $r \times r$ matrix whose (i, j) th entry is $C_{\boldsymbol{\alpha}}(\mathbf{s}_i^{(0)}, \mathbf{s}_j^{(0)})$. The MPP prior defines

$$(85) \quad \tilde{w}(\cdot) = \mathbf{c}^T(\cdot, \mathcal{S}^{(0)}) \mathbf{C}(\mathcal{S}^{(0)})^{-1} \mathbf{w}^{(0)} + \tilde{\epsilon}(\cdot),$$

where the processes $\tilde{\epsilon}(\cdot)$ and $w(\cdot)$ are mutually independent and $\tilde{\epsilon}(\cdot)$ is a GP with mean 0, $\text{cov}\{\tilde{\epsilon}(\mathbf{a}), \tilde{\epsilon}(\mathbf{b})\} = \delta(\mathbf{a}) \mathbf{1}(\mathbf{a} = \mathbf{b})$ for any $\mathbf{a}, \mathbf{b} \in \mathcal{D}$, and

$$\delta(\mathbf{s}_{ji}) = C_{\boldsymbol{\alpha}}(\mathbf{s}_{ji}, \mathbf{s}_{ji}) - \mathbf{c}^T(\mathbf{s}_{ji}, \mathcal{S}^{(0)}) \mathbf{C}(\mathcal{S}^{(0)})^{-1} \mathbf{c}(\mathbf{s}_{ji}, \mathcal{S}^{(0)}).$$

The process $\tilde{w}(\cdot)$ is a low-rank GP with mean 0 and

$$\text{cov}\{\tilde{w}(\mathbf{a}), \tilde{w}(\mathbf{b})\} = \mathbf{c}^T(\mathbf{a}, \mathcal{S}^{(0)}) \mathbf{C}(\mathcal{S}^{(0)})^{-1} \mathbf{c}(\mathbf{b}, \mathcal{S}^{(0)}) + \delta(\mathbf{a}) \mathbf{1}_{\mathbf{a}=\mathbf{b}}$$

for any $\mathbf{a}, \mathbf{b} \in \mathcal{D}$. If we replace $w(\cdot)$ by $\tilde{w}(\cdot)$ in (72), then

$$(86) \quad y(\mathbf{s}_{ji}) = \mathbf{x}(\mathbf{s}_{ji})^T \boldsymbol{\beta} + \tilde{w}(\mathbf{s}_{ji}) + \epsilon(\mathbf{s}_{ji}), \quad j = 1, \dots, k, \quad i = 1, \dots, m_j.$$

and our definition in (85) implies that $\tilde{w}(\cdot)$ is assigned a MPP prior (Finley et al., 2009).

We start by defining mean and covariance functions specific to univariate spatial regression using MPP. Let $\tilde{\mathbf{w}}_j = \{\tilde{w}(\mathbf{s}_{j1}), \dots, \tilde{w}(\mathbf{s}_{jm})\}$ and $\tilde{\mathbf{w}}_j^* = \{\tilde{w}(\mathbf{s}_1), \dots, \tilde{w}(\mathbf{s}_l)\}$. The MPP prior is identical to the FITC approximation in sparse approximate GP regression, so we use the FITC notations to simplify the description of posterior computations (Quiñonero-Candela and Rasmussen, 2005). Define $\mathbf{Q}_{j,j} = \mathbf{C}_{j,0}(\boldsymbol{\alpha}) \mathbf{C}^{-1}(\mathcal{S}^{(0)}) \mathbf{C}_{0,j}(\boldsymbol{\alpha})$, where $\text{cov}\{w(\mathbf{s}_{ja}), w(\mathbf{s}_b^{(0)})\} = \{\mathbf{C}_{j,0}(\boldsymbol{\alpha})\}_{a,b}$ ($a = 1, \dots, m$; $b = 1, \dots, r$) and $\mathbf{C}_{0,j}(\boldsymbol{\alpha}) = \mathbf{C}_{j,0}^T(\boldsymbol{\alpha})$. The density of $(\tilde{\mathbf{w}}_j, \tilde{\mathbf{w}}_j^*)$ under the GP prior implied by MPP is $N\{\mathbf{0}, \tilde{\mathbf{C}}(\boldsymbol{\alpha})\}$, where 2×2 block form of $\tilde{\mathbf{C}}(\boldsymbol{\alpha})$ is defined using

$$(87) \quad \begin{aligned} \tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) &= \mathbf{Q}_{j,j} + \text{diag}\{\mathbf{C}_{j,j}(\boldsymbol{\alpha}) - \mathbf{Q}_{j,j}\} = \text{cov}(\tilde{\mathbf{w}}_j, \tilde{\mathbf{w}}_j), \\ \tilde{\mathbf{C}}_{j,*}(\boldsymbol{\alpha}) &= \mathbf{Q}_{j,*} = \text{cov}(\tilde{\mathbf{w}}_j, \tilde{\mathbf{w}}_j^*), \\ \tilde{\mathbf{C}}_{*,*}(\boldsymbol{\alpha}) &= \mathbf{Q}_{*,*} + \text{diag}\{\mathbf{C}_{*,*}(\boldsymbol{\alpha}) - \mathbf{Q}_{*,*}\} = \text{cov}(\tilde{\mathbf{w}}_j^*, \tilde{\mathbf{w}}_j^*), \\ \tilde{\mathbf{C}}_{*,j}(\boldsymbol{\alpha}) &= \mathbf{Q}_{*,j} = \text{cov}(\tilde{\mathbf{w}}_j^*, \tilde{\mathbf{w}}_j). \end{aligned}$$

Stochastic approximation is implemented following Section 2. First, consider stochastic approximation for the likelihood of $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$. Integrating out $\tilde{\mathbf{w}}_j$ in (86) gives

$$(88) \quad \mathbf{y}_j = \mathbf{X}_j \boldsymbol{\beta} + \tilde{\boldsymbol{\eta}}_j, \quad \tilde{\boldsymbol{\eta}}_j \sim N\{\mathbf{0}, \tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}.$$

The likelihood of $\boldsymbol{\beta}$ and $\boldsymbol{\alpha}$ given $\mathbf{y}_j, \mathbf{X}_j$ after stochastic approximation is

$$(89) \quad \{l_j(\boldsymbol{\beta}, \boldsymbol{\alpha})\}^k = (2\pi)^{-mk/2} |\tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}|^{-k/2} e^{-\frac{k}{2}(\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})^T \{\tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}^{-1} (\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})}.$$

The prior distribution for $\boldsymbol{\beta}$ in (73), the pseudo likelihood in (89), and Bayes rule implies that the density of the j th subset posterior distribution for $\boldsymbol{\beta}$ given the rest is

$$\boldsymbol{\beta} \mid \text{rest} \propto e^{-\frac{1}{2}(\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})^T [k^{-1}\{\tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}]^{-1} (\mathbf{y}_j - \mathbf{X}_j \boldsymbol{\beta})} e^{-\frac{1}{2}(\boldsymbol{\beta} - \boldsymbol{\mu}_\beta)^T \Sigma_\beta^{-1} (\boldsymbol{\beta} - \boldsymbol{\mu}_\beta)}.$$

This implies that the complete conditional distribution of $\boldsymbol{\beta}_j$ has density $N(\tilde{\mathbf{m}}_{j\beta}, \tilde{\mathbf{V}}_{j\beta})$, where

$$(90) \quad \begin{aligned} \tilde{\mathbf{V}}_{j\beta} &= \left[k \mathbf{X}_j^T \{\tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}^{-1} \mathbf{X}_j + \Sigma_\beta^{-1} \right]^{-1}, \\ \tilde{\mathbf{m}}_{j\beta} &= \tilde{\mathbf{V}}_{j\beta} \left[k \mathbf{X}_j^T \{\tilde{\mathbf{C}}_{j,j}(\boldsymbol{\alpha}) + \tau^2 \mathbf{I}\}^{-1} \mathbf{y}_j + \Sigma_\beta^{-1} \boldsymbol{\mu}_\beta \right]. \end{aligned}$$

Following Section 2, the density of the j th subset posterior distribution for α given the rest is

$$(91) \quad \alpha \mid \text{rest} \propto |\tilde{\mathbf{C}}_{j,j}(\alpha) + \tau^2 \mathbf{I}|^{-k/2} e^{-\frac{1}{2}(\mathbf{y}_j - \mathbf{X}_j \beta)^T [k^{-1} \{\tilde{\mathbf{C}}_{j,j}(\alpha) + \tau^2 \mathbf{I}\}]^{-1} (\mathbf{y}_j - \mathbf{X}_j \beta)} \\ (\sigma^2)^{-a_\sigma - 1} e^{-b_\sigma / \sigma^2} (\tau^2)^{-a_\tau - 1} e^{-b_\tau / \tau^2} (b_\phi - a_\phi)^{-1}.$$

This density does not have a standard form, so we use a Metropolis-Hastings step with a normal random walk proposal and sample α_j using the `metrop` function in the R package `mcmc`.

Second, we derive the posterior predictive distribution of $\tilde{\mathbf{w}}_j^*$ given the rest. The MPP prior on $(\tilde{\mathbf{w}}_j, \tilde{\mathbf{w}}_j^*)$ implies that the density of $\tilde{\mathbf{w}}_j^*$ given $\tilde{\mathbf{w}}_j$ is

$$(92) \quad \tilde{\mathbf{w}}_j^* \mid \tilde{\mathbf{w}}_j \sim N \left\{ \tilde{\mathbf{C}}_{*,j}(\alpha) \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \tilde{\mathbf{w}}_j, \tilde{\mathbf{C}}_{*,*}(\alpha) - \tilde{\mathbf{C}}_{*,j}(\alpha) \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \tilde{\mathbf{C}}_{j,*}(\alpha) \right\}.$$

Given α , β , $\mathbf{y}_j$, and $\mathbf{X}_j$, (86) implies that the likelihood of $\tilde{\mathbf{w}}_j$ after stochastic approximation is

$$(93) \quad \{l_j(\tilde{\mathbf{w}}_j)\}^k = (2\pi)^{-mk/2} |\tau^2 \mathbf{I}|^{-k/2} e^{-\frac{k}{2\tau^2} (\mathbf{y}_j - \mathbf{X}_j \beta - \tilde{\mathbf{w}}_j)^T (\mathbf{y}_j - \mathbf{X}_j \beta - \tilde{\mathbf{w}}_j)}.$$

The MPP prior on $\tilde{\mathbf{w}}_j$, the pseudo likelihood in (93), and Bayes rule implies that the density of the subset posterior distribution for $\tilde{\mathbf{w}}_j$ given the rest is

$$\tilde{\mathbf{w}}_j \mid \text{rest} \propto e^{-\frac{1}{2\tau^2/k} (\mathbf{y}_j - \mathbf{X}_j \beta - \tilde{\mathbf{w}}_j)^T (\mathbf{y}_j - \mathbf{X}_j \beta - \tilde{\mathbf{w}}_j)} e^{-\frac{1}{2} \tilde{\mathbf{w}}_j^T \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \tilde{\mathbf{w}}_j}.$$

This implies that the complete conditional distribution of $\tilde{\mathbf{w}}_j$ has density $N(\mathbf{m}_{\tilde{\mathbf{w}}_j}, \mathbf{V}_{\tilde{\mathbf{w}}_j})$, where

$$(94) \quad \mathbf{V}_{\tilde{\mathbf{w}}_j} = \left\{ \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) + \frac{k}{\tau^2} \mathbf{I} \right\}^{-1}, \quad \mathbf{m}_{\tilde{\mathbf{w}}_j} = \frac{k}{\tau^2} \mathbf{V}_{\tilde{\mathbf{w}}_j} (\mathbf{y}_j - \mathbf{X}_j \beta);$$

therefore, (92) and (94) imply that the complete conditional distribution of $\tilde{\mathbf{w}}_j^*$ has density $N(\mathbf{m}_{\tilde{\mathbf{w}}_j^*}, \mathbf{V}_{\tilde{\mathbf{w}}_j^*})$, where

$$(95) \quad \mathbf{m}_{\tilde{\mathbf{w}}_j^*} = \mathbb{E}(\tilde{\mathbf{w}}_j^* \mid \text{rest}) = \tilde{\mathbf{C}}_{*,j}(\alpha) \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \mathbb{E}(\tilde{\mathbf{w}}_j \mid \text{rest}) \\ = \tilde{\mathbf{C}}_{*,j}(\alpha) \left\{ \tilde{\mathbf{C}}_{j,j}(\alpha) + \frac{\tau^2}{k} \mathbf{I} \right\}^{-1} (\mathbf{y}_j - \mathbf{X}_j \beta)$$

and

$$(96) \quad \mathbf{V}_{\tilde{\mathbf{w}}_j^*} = \text{var}(\tilde{\mathbf{w}}_j^* \mid \text{rest}) = \mathbb{E} \left\{ \text{var}(\tilde{\mathbf{w}}_j^* \mid \tilde{\mathbf{w}}_j) \mid \text{rest} \right\} + \text{var} \left\{ \mathbb{E}(\tilde{\mathbf{w}}_j^* \mid \tilde{\mathbf{w}}_j) \mid \text{rest} \right\} \\ = \tilde{\mathbf{C}}_{*,*}(\alpha) - \tilde{\mathbf{C}}_{*,j}(\alpha) \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \tilde{\mathbf{C}}_{j,*}(\alpha) + \tilde{\mathbf{C}}_{*,j}(\alpha) \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \mathbf{V}_{\tilde{\mathbf{w}}_j} \tilde{\mathbf{C}}_{j,j}^{-1}(\alpha) \tilde{\mathbf{C}}_{j,*}(\alpha).$$

Finally, we derive the posterior predictive distribution of $\mathbf{y}_j^*$ given the rest. If β_j , τ_j^2 , $\tilde{\mathbf{w}}_j^*$ are the samples from the j th subset posterior distribution of β , τ^2 , and $\tilde{\mathbf{w}}^*$, then (86) implies that $\mathbf{y}_j^*$ given the rest is sampled as

$$\mathbf{y}_j^* = \mathbf{X}_j \beta_j + \tilde{\mathbf{w}}_j^* + \epsilon_j^*, \quad \epsilon_j^* \sim N(\mathbf{0}, \tau_j^2 \mathbf{I});$$

therefore, the complete conditional distribution of $\mathbf{y}_j^*$ has density $N(\tilde{\boldsymbol{\mu}}_{\mathbf{y}_j^*}, \tilde{\mathbf{V}}_{\mathbf{y}_j^*})$, where

$$(97) \quad \tilde{\boldsymbol{\mu}}_{\mathbf{y}_j^*} = \mathbf{X}_j \boldsymbol{\beta}_j + \tilde{\mathbf{w}}_j^*, \quad \tilde{\mathbf{V}}_{\mathbf{y}_j^*} = \tau_j^2 \mathbf{I}.$$

All full conditionals except that of $\boldsymbol{\alpha}$ are analytically tractable in terms of standard distributions in subset j ($j = 1, \dots, k$). The Gibbs sampler with a Metropolis-Hastings step iterates between the following four steps until sufficient number of samples of $\boldsymbol{\beta}_j, \boldsymbol{\alpha}_j, \tilde{\mathbf{w}}_j^*$, and $\mathbf{y}_j^*$ are drawn post convergence to the stationary distribution:

1. Sample $\boldsymbol{\beta}_j$ from $N(\tilde{\boldsymbol{\mu}}_{j\boldsymbol{\beta}}, \tilde{\mathbf{V}}_{j\boldsymbol{\beta}})$, where $\tilde{\boldsymbol{\mu}}_{j\boldsymbol{\beta}}$ and $\tilde{\mathbf{V}}_{j\boldsymbol{\beta}}$ are defined in (90).
2. Sample $\boldsymbol{\alpha}_j$ using the Metropolis-Hastings algorithm from the j th subset posterior density (up to constants) of $\boldsymbol{\alpha}_j$ in (91) with a normal random walk proposal.
3. Sample $\tilde{\mathbf{w}}_j^*$ from $N(\boldsymbol{\mu}_{\tilde{\mathbf{w}}_j^*}, \mathbf{V}_{\tilde{\mathbf{w}}_j^*})$, where $\boldsymbol{\mu}_{\tilde{\mathbf{w}}_j^*}$ and $\mathbf{V}_{\tilde{\mathbf{w}}_j^*}$ are defined in (95) and (96).
4. Sample $\mathbf{y}_j^*$ from $N(\tilde{\boldsymbol{\mu}}_{\mathbf{y}_j^*}, \tilde{\mathbf{V}}_{\mathbf{y}_j^*})$, where $\tilde{\boldsymbol{\mu}}_{\mathbf{y}_j^*}$ and $\tilde{\mathbf{V}}_{\mathbf{y}_j^*}$ are defined in (97).

4. COMPARISONS BETWEEN DIVIDE-AND-CONQUER COMPETITORS

4.1 Setup

We compare the four competitors based on the divide-and-conquer technique. Extending Section 4 of the main manuscript, we compare the performance based on learning the process parameters, interpolating the unobserved spatial surface, and predicting the response at new locations. This section presents two simulation studies and one real data analysis. Recall that the first and second simulations (*Simulation 1*) generate data from a spatial linear model where the spatial processes are simulated from a GP and an analytic function with local features, respectively. The number of locations in the two simulations is moderately large with $n = 10,000$. Continuing from the main manuscript, our real data analysis is based on a large data subset of sea surface temperature data with $n = 1,00,000$ locations. For all the three simulations, the response at $(n+l)$ locations is modeled as

$$(98) \quad y(\mathbf{s}_i) = \beta_0 + x(\mathbf{s}_i)\beta_1 + w(\mathbf{s}_i) + \epsilon_i, \quad \epsilon_i \sim N(0, \tau^2), \quad \mathbf{s}_i \in \mathcal{D} \subset \mathbb{R}^2,$$

for $i = 1, \dots, n+l$, where $\mathcal{D}$ is the spatial domain, $y(\mathbf{s}_i)$, $x(\mathbf{s}_i)$, $w(\mathbf{s}_i)$, and ϵ_i are the response, covariate, spatial process, and idiosyncratic error values at the location $\mathbf{s}_i$, β_0 is the intercept, β_1 models the covariate effect, and l is the number of new locations.

The three-step DISK, WASP, DPMC and CMC frameworks are applied using the low-rank MPP priors using the algorithm outlined in Section 3.3 of the main paper with two partitioning schemes. The first partitioning scheme randomly partitions the spatial locations in k groups. In the second partitioning scheme, we divide the spatial domain into sixteen square grid cells and randomly allocate locations in every grid cell into k groups.

We compare the quality of prediction and estimation of spatial surface at predictive locations $\mathcal{S}^* = \{\mathbf{s}_1^*, \dots, \mathbf{s}_l^*\}$. If $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$ are the value of the

spatial surface and response at $\mathbf{s}_{i'}^* \in \mathcal{S}^*$, then the estimation and prediction errors are defined as

$$(99) \quad \text{Est Err}^2 = \frac{1}{l} \sum_{i'=1}^l \{\hat{w}(\mathbf{s}_{i'}^*) - w(\mathbf{s}_{i'}^*)\}^2, \quad \text{Pred Err}^2 = \frac{1}{l} \sum_{i'=1}^l \{\hat{y}(\mathbf{s}_{i'}^*) - y(\mathbf{s}_{i'}^*)\}^2,$$

where $\hat{w}(\mathbf{s}_{i'}^*)$ and $\hat{y}(\mathbf{s}_{i'}^*)$ denote the estimates of $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$ obtained using any distributed or non-distributed methods. For sampling-based methods, we set $\hat{w}(\mathbf{s}_{i'}^*)$ and $\hat{y}(\mathbf{s}_{i'}^*)$ to be the medians of posterior MCMC samples for $w(\mathbf{s}_{i'}^*)$ and $y(\mathbf{s}_{i'}^*)$, respectively, for $i' = 1, \dots, l$. We also estimate the point-wise 95% credible or confidence intervals (CIs) of $w(\mathbf{s}_{i'}^*)$ and predictive intervals (PIs) of $y(\mathbf{s}_{i'}^*)$ for every $\mathbf{s}_{i'} \in \mathcal{S}^*$ and compare the CI and PI coverages and lengths for every method. Finally, we compare the performance of all the methods for parameter estimation using the posterior medians or point estimates and the 95% CIs.

4.2 Simulation 1: Spatial Linear Model Based On GP

This section compares DISK with its divide-and-conquer competitors under the two partitioning schemes and is a continuation of Section 4.2 of the main manuscript. The four divide-and-conquer methods, CMC, DISK, WASP, and DPMC, have similar performance in parameter estimation (Tables 1, 2, and 3). The parameter estimates obtained using all these methods are close to the truth and estimation errors are also very similar. The 95% credible intervals of β_0, β_1, τ^2 in cover the true values and their lower and upper quantiles are very similar. All the four methods underestimate σ^2 and overestimate ϕ slightly. Both results are the impacts of parent MPP prior, which also shows a similar pattern for the two choices of r . We notice that the coverage of CMC is smaller than that of DISK, WASP, and DPMC. More importantly, the choice of r, k , or partitioning scheme has a minimal impact on parameter estimation in DISK, WASP, and DPMC.

The inferential and predictive performance of DISK, WASP, and DPMC are similar, but CMC shows significant differences (Table 4). There are minimal differences in the prediction and estimation errors of CMC, DISK, WASP, and DPMC. This indicates that the estimate of posterior medians are very similar in all the four methods; however, the pointwise coverage of CMC in prediction of the response and inference on the spatial surface is significantly smaller than the nominal value for every choice of r and k . On the other hand, DISK, WASP, and DPMC have nominal coverage in prediction and inference on the spatial surface. Furthermore, their CI and PI coverage values are robust to the choices of r, k , and partitioning scheme.

In summary, the DISK, WASP, and DPMC have similar inferential and predictive performance. While CMC's point estimates are close to those of DISK, WASP, and DPMC, its inferential and predictive performance is worse than its three competitors. The partitioning scheme, random or grid-based, has no impact on the performance of all the four divide-and-conquer methods.

4.3 Simulation 2: Spatial Linear Model Based On Analytic Spatial Surface

This section compares DISK with its divide-and-conquer competitors and is a continuation of Section 4.3 of the main manuscript. Our conclusions remain similar as those observed in the previous section. Specifically, CMC, DISK, WASP, and DPMC have similar performance in parameter estimation (Tables 5, 6, and

TABLE 1

The errors in estimating the parameters $\beta = (\beta_0, \beta_1), \sigma^2, \phi, \tau^2$ in Simulation 1 for the divide-and-conquer methods under random and grid-based partitioning. The parameter estimates for the Bayesian methods $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1), \hat{\sigma}^2, \hat{\phi}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples and their true values are $\beta_0 = (1, 2), \sigma_0^2 = 1, \phi_0 = 4$ and $\tau_0^2 = 0.1$. The entries in the table are averaged across 10 simulation replications.

	$\ \hat{\beta} - \beta_0\ $	$ \hat{\sigma}^2 - \sigma_0^2 $	$ \hat{\phi} - \phi_0 $	$ \hat{\tau}^2 - \tau_0^2 $
Random Partitioning				
CMC ($r = 200, k = 10$)	0.09	0.12	0.68	0.01
CMC ($r = 400, k = 10$)	0.09	0.12	0.75	0.01
CMC ($r = 200, k = 20$)	0.10	0.13	0.95	0.02
CMC ($r = 400, k = 20$)	0.10	0.13	0.82	0.02
DISK ($r = 200, k = 10$)	0.09	0.11	0.64	0.01
DISK ($r = 400, k = 10$)	0.09	0.11	0.64	0.01
DISK ($r = 200, k = 20$)	0.10	0.12	0.66	0.02
DISK ($r = 400, k = 20$)	0.10	0.12	0.66	0.02
WASP ($r = 200, k = 10$)	0.09	0.11	0.64	0.01
WASP ($r = 400, k = 10$)	0.09	0.11	0.63	0.01
WASP ($r = 200, k = 20$)	0.10	0.12	0.66	0.02
WASP ($r = 400, k = 20$)	0.10	0.12	0.66	0.02
DPMC ($r = 200, k = 10$)	0.09	0.11	0.64	0.01
DPMC ($r = 400, k = 10$)	0.09	0.11	0.63	0.01
DPMC ($r = 200, k = 20$)	0.10	0.12	0.66	0.02
DPMC ($r = 400, k = 20$)	0.10	0.12	0.66	0.02
Grid-Based Partitioning				
CMC ($r = 200, k = 10$)	0.09	0.12	0.63	0.01
CMC ($r = 400, k = 10$)	0.09	0.12	0.65	0.01
CMC ($r = 200, k = 20$)	0.10	0.13	0.77	0.01
CMC ($r = 400, k = 20$)	0.10	0.13	0.83	0.01
DISK ($r = 200, k = 10$)	0.09	0.12	0.62	0.01
DISK ($r = 400, k = 10$)	0.09	0.12	0.62	0.01
DISK ($r = 200, k = 20$)	0.10	0.12	0.63	0.01
DISK ($r = 400, k = 20$)	0.10	0.12	0.64	0.01
WASP ($r = 200, k = 10$)	0.09	0.12	0.62	0.01
WASP ($r = 400, k = 10$)	0.09	0.12	0.62	0.01
WASP ($r = 200, k = 20$)	0.10	0.12	0.63	0.01
WASP ($r = 400, k = 20$)	0.10	0.12	0.64	0.01
DPMC ($r = 200, k = 10$)	0.09	0.12	0.62	0.01
DPMC ($r = 400, k = 10$)	0.09	0.12	0.62	0.01
DPMC ($r = 200, k = 20$)	0.10	0.12	0.63	0.01
DPMC ($r = 400, k = 20$)	0.10	0.12	0.64	0.01

7); however, the inferential and predictive performance of DISK, WASP, and DPMC are significantly better than those of CMC (Table 8). The partitioning scheme, random or grid-based, has no impact on the inferential and predictive performance of CMC, DISK, WASP, and DPMC. The results are also robust to the choices of k and r .

4.4 Real data analysis: Sea Surface Temperature data

This section is a continuation of Section 4.3 of the main manuscript and compares DISK with its divide-and-conquer competitors in analyzing the Sea Surface Temperature (SST) data. We have chosen random partitioning based on our conclusions in the previous two simulations. Our results for SST data analysis are also very similar to those in the previous two simulations. Specifically, CMC DISK, WASP, and DPMC have similar performance in parameter estimation, but signif-

TABLE 2

The estimates of parameters $\beta = (\beta_0, \beta_1)$, σ^2 , ϕ , τ^2 and their 95% marginal credible intervals (CIs) in Simulation 1 for the divide-and-conquer methods under random partitioning. The parameter estimates for the Bayesian methods $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1)$, $\hat{\sigma}^2$, $\hat{\phi}$, $\hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications.

	β_0	β_1	σ^2	ϕ	τ^2
Truth	1.00	2.00	1.00	4.00	0.10
Parameter Estimates					
CMC ($r = 200, k = 10$)	1.00	2.00	0.91	4.38	0.10
CMC ($r = 400, k = 10$)	1.00	2.00	0.91	4.41	0.10
CMC ($r = 200, k = 20$)	1.00	2.00	0.90	4.55	0.10
CMC ($r = 400, k = 20$)	1.00	2.00	0.91	4.46	0.10
DISK ($r = 200, k = 10$)	1.00	2.00	0.92	4.35	0.11
DISK ($r = 400, k = 10$)	1.00	2.00	0.92	4.35	0.11
DISK ($r = 200, k = 20$)	1.00	2.00	0.91	4.38	0.11
DISK ($r = 400, k = 20$)	1.00	2.00	0.91	4.38	0.11
WASP ($r = 200, k = 10$)	1.00	2.00	0.92	4.35	0.11
WASP ($r = 400, k = 10$)	1.00	2.00	0.92	4.34	0.11
WASP ($r = 200, k = 20$)	1.00	2.00	0.91	4.38	0.11
WASP ($r = 400, k = 20$)	1.00	2.00	0.91	4.38	0.11
DPMC ($r = 200, k = 10$)	1.00	2.00	0.92	4.35	0.11
DPMC ($r = 400, k = 10$)	1.00	2.00	0.92	4.34	0.11
DPMC ($r = 200, k = 20$)	1.00	2.00	0.91	4.38	0.11
DPMC ($r = 400, k = 20$)	1.00	2.00	0.91	4.38	0.11
95% Credible Intervals					
CMC ($r = 200, k = 10$)	(0.98, 1.02)	(2.00, 2.00)	(0.90, 0.93)	(4.28, 4.49)	(0.10, 0.11)
CMC ($r = 400, k = 10$)	(0.98, 1.02)	(2.00, 2.00)	(0.90, 0.93)	(4.31, 4.52)	(0.10, 0.11)
CMC ($r = 200, k = 20$)	(0.99, 1.01)	(2.00, 2.00)	(0.89, 0.92)	(4.49, 4.61)	(0.10, 0.10)
CMC ($r = 400, k = 20$)	(0.99, 1.01)	(2.00, 2.00)	(0.90, 0.92)	(4.40, 4.53)	(0.10, 0.10)
DISK ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
DISK ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
DISK ($r = 200, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.07, 4.67)	(0.09, 0.13)
DISK ($r = 400, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.07, 4.68)	(0.09, 0.13)
WASP ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
WASP ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.69)	(0.09, 0.12)
WASP ($r = 200, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.07, 4.67)	(0.09, 0.13)
WASP ($r = 400, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.07, 4.68)	(0.09, 0.13)
DPMC ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(4.00, 4.70)	(0.09, 0.12)
DPMC ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.99, 4.70)	(0.09, 0.12)
DPMC ($r = 200, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.06, 4.68)	(0.09, 0.13)
DPMC ($r = 400, k = 20$)	(0.94, 1.06)	(1.98, 2.01)	(0.86, 0.96)	(4.06, 4.69)	(0.09, 0.13)

icant differences exist in their predictive performance (Table 9). CMC's predictive coverage is much smaller than the nominal value, which matches our conclusions in the previous two simulations. DISK outperforms WASP and DPMC in predictions in that its MSPE is the smallest among them. DISK also has better nominal predictive coverage than WASP and DPMC while having comparable 95% PI lengths. The results are also robust to the choices of r . We conclude that DISK performs better than its divide-and-conquer competitors in SST data analysis.

4.5 Computation time comparisons

We report the run-times of all the methods used in the simulated and real data analysis in Section 4 of the main paper. Since distributed methods partition the data into the same subset size and fit the same MPP model for subset posterior inference, the run times are identical for any method in Simulation 1 and 2. Thus, we only present run times for simulation and for the sea surface temperature data; see Tables 10 and 11 for the run-times in $\log_{10}$ seconds for Simulation

TABLE 3

The estimates of parameters $\beta = (\beta_0, \beta_1)$, σ^2 , ϕ , τ^2 and their 95% marginal credible intervals (CIs) in Simulation 1 for the divide-and-conquer methods under grid-based partitioning. The parameter estimates for the Bayesian methods $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1)$, $\hat{\sigma}^2$, $\hat{\phi}$, $\hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications.

	β_0	β_1	σ^2	ϕ	τ^2
Truth	1.00	2.00	1.00	4.00	0.10
Parameter Estimates					
CMC ($r = 200, k = 10$)	1.00	2.00	0.91	4.37	0.10
CMC ($r = 400, k = 10$)	1.00	2.00	0.91	4.37	0.10
CMC ($r = 200, k = 20$)	1.00	2.00	0.91	4.44	0.10
CMC ($r = 400, k = 20$)	1.00	2.00	0.90	4.48	0.10
DISK ($r = 200, k = 10$)	1.00	2.00	0.91	4.35	0.11
DISK ($r = 400, k = 10$)	1.00	2.00	0.91	4.35	0.11
DISK ($r = 200, k = 20$)	1.00	2.00	0.91	4.37	0.11
DISK ($r = 400, k = 20$)	1.00	2.00	0.91	4.37	0.11
WASP ($r = 200, k = 10$)	1.00	2.00	0.91	4.35	0.11
WASP ($r = 400, k = 10$)	1.00	2.00	0.91	4.34	0.11
WASP ($r = 200, k = 20$)	1.00	2.00	0.91	4.37	0.11
WASP ($r = 400, k = 20$)	1.00	2.00	0.91	4.37	0.11
DPMC ($r = 200, k = 10$)	1.00	2.00	0.91	4.35	0.11
DPMC ($r = 400, k = 10$)	1.00	2.00	0.91	4.34	0.11
DPMC ($r = 200, k = 20$)	1.00	2.00	0.91	4.37	0.11
DPMC ($r = 400, k = 20$)	1.00	2.00	0.91	4.37	0.11
95% Credible Intervals					
CMC ($r = 200, k = 10$)	(0.98, 1.02)	(2.00, 2.00)	(0.89, 0.93)	(4.27, 4.47)	(0.10, 0.11)
CMC ($r = 400, k = 10$)	(0.98, 1.02)	(2.00, 2.00)	(0.89, 0.93)	(4.27, 4.48)	(0.10, 0.11)
CMC ($r = 200, k = 20$)	(0.99, 1.01)	(2.00, 2.00)	(0.90, 0.92)	(4.38, 4.50)	(0.10, 0.11)
CMC ($r = 400, k = 20$)	(0.99, 1.01)	(2.00, 2.00)	(0.89, 0.92)	(4.42, 4.54)	(0.10, 0.11)
DISK ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.99, 4.69)	(0.09, 0.12)
DISK ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.99, 4.70)	(0.09, 0.12)
DISK ($r = 200, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.06, 4.67)	(0.09, 0.13)
DISK ($r = 400, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.07, 4.67)	(0.09, 0.13)
WASP ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.99, 4.69)	(0.10, 0.12)
WASP ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.98, 4.70)	(0.09, 0.12)
WASP ($r = 200, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.06, 4.67)	(0.09, 0.13)
WASP ($r = 400, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.07, 4.67)	(0.09, 0.13)
DPMC ($r = 200, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.98, 4.70)	(0.09, 0.12)
DPMC ($r = 400, k = 10$)	(0.92, 1.08)	(1.99, 2.01)	(0.86, 0.98)	(3.98, 4.71)	(0.09, 0.12)
DPMC ($r = 200, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.05, 4.69)	(0.09, 0.13)
DPMC ($r = 400, k = 20$)	(0.94, 1.06)	(1.99, 2.01)	(0.86, 0.96)	(4.07, 4.68)	(0.09, 0.13)

1 and in $\log_{10}$ hours for sea surface temperature data analysis, respectively. Similar to our observations in the performance comparisons, the run-times for the distributed methods are independent of the partitioning schemes. The run-times cannot be compared directly from the tables due to the differences in implementation. Specifically, distributed methods are implemented in R for all values of r and k , whereas most non-distributed methods are implemented in R and a higher-level language, such as C/C++ and Fortran.

The combination step in any distributed method requires a very small time compared to the time required for sampling on the subsets. For example, the time required for combination using the WASP is the largest among all the four distributed methods, but the maximum of WASP's combination time is only 8% of the maximum time required for sampling on the subsets. On an average, the combination steps of the other three methods require less than 1% of the time required for sampling on the subsets. This implies that run-times for all the four distributed methods in the two simulations are fairly similar (Table 10). In the real data analysis, the combination steps of the all the four distributed methods

TABLE 4

Inference on the values of spatial surface and response at the locations in $\mathcal{S}_$ in Simulation 1 for the divide-and-conquer methods under random and grid-based partitioning. The estimation and prediction errors are defined in (99) and coverage and credible intervals are calculated pointwise for the locations in $\mathcal{S}_*$. The entries in the table are averaged over 10 simulation replications.*

	Est Err	Pred Err	95% CI Coverage		95% CI Length	
	GP	Y	GP	Y	GP	Y
Random Partitioning						
CMC ($r = 200, k = 10$)	0.56	0.64	0.38	0.39	0.81	0.81
CMC ($r = 400, k = 10$)	0.43	0.52	0.40	0.41	0.74	0.74
CMC ($r = 200, k = 20$)	0.58	0.67	0.27	0.28	0.57	0.57
CMC ($r = 400, k = 20$)	0.46	0.55	0.28	0.29	0.52	0.52
DISK ($r = 200, k = 10$)	0.55	0.64	0.97	0.97	3.20	3.45
DISK ($r = 400, k = 10$)	0.42	0.51	0.97	0.97	2.88	3.15
DISK ($r = 200, k = 20$)	0.58	0.67	0.97	0.97	3.25	3.51
DISK ($r = 400, k = 20$)	0.46	0.55	0.97	0.97	2.98	3.25
WASP ($r = 200, k = 10$)	0.55	0.64	0.96	0.96	3.25	3.25
WASP ($r = 400, k = 10$)	0.42	0.51	0.96	0.96	2.97	2.97
WASP ($r = 200, k = 20$)	0.58	0.67	0.96	0.96	3.30	3.30
WASP ($r = 400, k = 20$)	0.46	0.55	0.96	0.96	3.06	3.06
DPMC ($r = 200, k = 10$)	0.55	0.64	0.97	0.97	3.46	3.46
DPMC ($r = 400, k = 10$)	0.42	0.51	0.97	0.97	3.17	3.17
DPMC ($r = 200, k = 20$)	0.58	0.67	0.97	0.97	3.53	3.53
DPMC ($r = 400, k = 20$)	0.46	0.55	0.97	0.97	3.28	3.28
Grid-Based Partitioning						
CMC ($r = 200, k = 10$)	0.75	0.80	0.38	0.39	0.81	0.81
CMC ($r = 400, k = 10$)	0.65	0.72	0.40	0.40	0.74	0.74
CMC ($r = 200, k = 20$)	0.76	0.82	0.27	0.28	0.57	0.57
CMC ($r = 400, k = 20$)	0.68	0.74	0.28	0.28	0.52	0.52
DISK ($r = 200, k = 10$)	0.75	0.80	0.97	0.97	3.45	3.45
DISK ($r = 400, k = 10$)	0.65	0.72	0.97	0.97	3.15	3.15
DISK ($r = 200, k = 20$)	0.76	0.82	0.97	0.97	3.51	3.51
DISK ($r = 400, k = 20$)	0.68	0.74	0.97	0.97	3.26	3.26
WASP ($r = 200, k = 10$)	0.74	0.80	0.96	0.96	3.25	3.25
WASP ($r = 400, k = 10$)	0.65	0.72	0.96	0.96	2.97	2.97
WASP ($r = 200, k = 20$)	0.76	0.82	0.96	0.95	3.30	3.30
WASP ($r = 400, k = 20$)	0.68	0.74	0.96	0.96	3.06	3.06
DPMC ($r = 200, k = 10$)	0.74	0.80	0.97	0.97	3.46	3.46
DPMC ($r = 400, k = 10$)	0.65	0.72	0.97	0.97	3.16	3.16
DPMC ($r = 200, k = 20$)	0.76	0.82	0.97	0.97	3.53	3.53
DPMC ($r = 400, k = 20$)	0.68	0.74	0.97	0.97	3.28	3.28

require less than 1% of the time required for sampling on the subsets, so all of them have nearly identical run-times (Table 11).

5. MARKOV CHAINS ON THE SUBSETS IN DISK

Any divide-and-conquer method runs modified Markov chain Monte Carlo algorithms in parallel on the subsets to obtain draws from the subset posterior distributions. In our context, we draw parameter and response values from the respective posterior distributions on every subset. There are no theoretical results that guarantee convergence of the Markov chain produced by the sampling algorithms to its stationary distribution in a spatial linear model with MPP prior. This further complicates the theoretical analysis of the Markov chain produced on the subsets in DISK, where the likelihood is modified. We are not aware any rigorous approach for comparing the Markov chains obtained from the subset and true posterior distributions. We use heuristics based on trace plots and autocorrelation functions of the Markov chains for parameters, spatial surface, and predictive surface to judge “convergence” to the respective subset posterior dis-

TABLE 5

The errors in estimating the parameters β, τ^2 in Simulation 2 for the divide-and-conquer methods under random and grid-based partitioning. The parameter estimates for the Bayesian methods $\hat{\beta}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples and $\beta_0 = 1$ and $\tau_0^2 = 0.01$. The entries in the table are averaged across 10 simulation replications.

	$\ \hat{\beta} - \beta_0\ $	$ \hat{\tau}^2 - \tau_0^2 $
Random Partitioning		
CMC ($r = 200, k = 10$)	0.03	0.00
CMC ($r = 400, k = 10$)	0.03	0.09
CMC ($r = 200, k = 20$)	1.41	0.09
CMC ($r = 400, k = 20$)	1.41	0.09
DISK ($r = 200, k = 10$)	0.18	0.04
DISK ($r = 400, k = 10$)	0.13	0.04
DISK ($r = 200, k = 20$)	0.18	0.04
DISK ($r = 400, k = 20$)	0.13	0.04
WASP ($r = 200, k = 10$)	0.68	0.09
WASP ($r = 400, k = 10$)	0.68	0.09
WASP ($r = 200, k = 20$)	0.72	0.09
WASP ($r = 400, k = 20$)	0.72	0.09
DPMC ($r = 200, k = 10$)	0.68	0.09
DPMC ($r = 400, k = 10$)	0.68	0.09
DPMC ($r = 200, k = 20$)	0.72	0.09
DPMC ($r = 400, k = 20$)	0.72	0.09
Grid-Based Partitioning		
CMC ($r = 200, k = 10$)	0.03	0.09
CMC ($r = 400, k = 10$)	0.03	0.09
CMC ($r = 200, k = 20$)	0.02	0.09
CMC ($r = 400, k = 20$)	0.02	0.09
DISK ($r = 200, k = 10$)	0.03	0.09
DISK ($r = 400, k = 10$)	0.03	0.09
DISK ($r = 200, k = 20$)	0.02	0.09
DISK ($r = 400, k = 20$)	0.02	0.09
WASP ($r = 200, k = 10$)	0.03	0.09
WASP ($r = 400, k = 10$)	0.03	0.09
WASP ($r = 200, k = 20$)	0.02	0.09
WASP ($r = 400, k = 20$)	0.02	0.09
DPMC ($r = 200, k = 10$)	0.03	0.09
DPMC ($r = 400, k = 10$)	0.03	0.09
DPMC ($r = 200, k = 20$)	0.02	0.09
DPMC ($r = 400, k = 20$)	0.02	0.09

tributions.

Unfortunately, it is impractical to compare trace plots and auto correlation functions obtained using subset and true posterior distributions; therefore, we compare the effective sample sizes of Markov chains for the parameters, spatial surface, and response obtained on the subsets using an MPP prior relative to those obtained using the full data and the same MPP prior. The number of posterior samples in both cases is 1000, which are obtained from a Markov chain of 10000 draws after using a burn-in of 5000 and collecting every fifth sample. The `effectiveSize` command coda R package is used for estimating the effective sample sizes for every choice of k and r (Plummer et al., 2006). We compute the ratio of the effective sample sizes of the Markov chains produced on the subsets in DISK to those obtained using the MPP prior and the full data. For two- or higher-dimensional parameters, spatial surface, and predictive surface, we average the ratio of the effective sample sizes across all the dimensions. While there are no theoretical justifying the convergence of the Markov chain to the stationary distribution, we still assume so because MPP has been used extensively for analyzing spatial data. This heuristic shows that the Markov chains obtained using the data subsets and full data are “similar” in that their effective sample

TABLE 6

The estimates of parameters $\beta, \sigma^2, \phi, \tau^2$ and their 95% marginal credible intervals (CIs) in Simulation 2 for the divide-and-conquer methods under random partitioning. The parameter estimates for the Bayesian methods $\hat{\beta}, \hat{\sigma}^2, \hat{\phi}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications

	β	σ^2	ϕ	τ^2
Truth	1.00	-	-	0.01
Parameter Estimates				
CMC ($r = 200, k = 10$)	1.03	0.22	0.11	0.01
CMC ($r = 400, k = 10$)	1.03	0.22	0.11	0.01
CMC ($r = 200, k = 20$)	0.98	0.23	0.13	0.01
CMC ($r = 400, k = 20$)	0.98	0.23	0.13	0.01
DISK ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
DISK ($r = 400, k = 10$)	0.98	0.22	0.14	0.01
DISK ($r = 200, k = 20$)	1.03	0.21	0.12	0.01
DISK ($r = 400, k = 20$)	0.98	0.22	0.14	0.01
WASP ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
WASP ($r = 400, k = 10$)	1.03	0.21	0.12	0.01
WASP ($r = 200, k = 20$)	0.98	0.22	0.14	0.01
WASP ($r = 400, k = 20$)	0.98	0.22	0.14	0.01
DPMC ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
DPMC ($r = 400, k = 10$)	1.03	0.21	0.12	0.01
DPMC ($r = 200, k = 20$)	0.98	0.22	0.14	0.01
DPMC ($r = 400, k = 20$)	0.98	0.22	0.14	0.01
95% Credible Intervals				
CMC ($r = 200, k = 10$)	(0.96, 1.11)	(0.22, 0.23)	(0.11, 0.12)	(0.01, 0.01)
CMC ($r = 400, k = 10$)	(0.96, 1.11)	(0.22, 0.23)	(0.11, 0.12)	(0.01, 0.01)
CMC ($r = 200, k = 20$)	(0.94, 1.02)	(0.22, 0.24)	(0.13, 0.13)	(0.01, 0.01)
CMC ($r = 400, k = 20$)	(0.94, 1.02)	(0.22, 0.24)	(0.12, 0.13)	(0.01, 0.01)
DISK ($r = 200, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
DISK ($r = 400, k = 10$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
DISK ($r = 200, k = 20$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
DISK ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
WASP ($r = 200, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
WASP ($r = 400, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
WASP ($r = 200, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
WASP ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
DPMC ($r = 200, k = 10$)	(0.80, 1.27)	(0.17, 0.25)	(0.10, 0.15)	(0.01, 0.01)
DPMC ($r = 400, k = 10$)	(0.80, 1.27)	(0.17, 0.25)	(0.10, 0.15)	(0.01, 0.01)
DPMC ($r = 200, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.11, 0.18)	(0.01, 0.01)
DPMC ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.11, 0.19)	(0.01, 0.01)

sizes are very close.

The effective sample sizes of the Markov chains for the parameters and the spatial surface and response at the locations in $\mathcal{S}^*$ are very similar to those obtained using the full data and the same MPP prior in Simulation 1 (Table 12). The effective sample sizes decrease with k in Simulation 2 slightly for the spatial surface and response at the locations in $\mathcal{S}^*$ (Table 13); however, this spatial surface is not simulated from a GP in this simulation, so the comparisons are less reliable. The partitioning scheme, random or grid-based, has a minimal impact on the effective sample sizes. The ratio of the effective sample sizes are equal for the β , spatial surface, and predictions in Simulation 1; however, there are differences in the effective sample sizes of the Markov chains for σ^2, ϕ, τ^2 in both simulations. These differences mainly arise due to the non-identifiability of the covariance function parameters. In most spatial applications, the main interest lies in inference and prediction, where the effective sample sizes on the subsets are very similar to their full data benchmarks; therefore, we conclude that the Markov chains produced on the subsets in DISK have similar properties as their full data

TABLE 7

The estimates of parameters $\beta, \sigma^2, \phi, \tau^2$ and their 95% marginal credible intervals (CIs) in Simulation 2 for the divide-and-conquer methods under grid-based partitioning. The parameter estimates for the Bayesian methods $\hat{\beta}, \hat{\sigma}^2, \hat{\phi}, \hat{\tau}^2$ are defined as the posterior medians of their respective MCMC samples. The parameter estimates and upper and lower quantiles of 95% CIs are averaged over 10 simulation replications

	β	σ^2	ϕ	τ^2
Truth	1.00	-	-	0.01
Parameter Estimates				
CMC ($r = 200, k = 10$)	1.03	0.22	0.11	0.01
CMC ($r = 400, k = 10$)	1.03	0.22	0.11	0.01
CMC ($r = 200, k = 20$)	0.98	0.23	0.13	0.01
CMC ($r = 400, k = 20$)	0.98	0.23	0.13	0.01
DISK ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
DISK ($r = 400, k = 10$)	1.03	0.21	0.12	0.01
DISK ($r = 200, k = 20$)	0.98	0.22	0.14	0.01
DISK ($r = 400, k = 20$)	0.98	0.22	0.14	0.01
WASP ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
WASP ($r = 400, k = 10$)	1.03	0.21	0.12	0.01
WASP ($r = 200, k = 20$)	0.98	0.22	0.14	0.01
WASP ($r = 400, k = 20$)	0.99	0.22	0.14	0.01
DPMC ($r = 200, k = 10$)	1.03	0.21	0.12	0.01
DPMC ($r = 400, k = 10$)	1.03	0.21	0.12	0.01
DPMC ($r = 200, k = 20$)	0.98	0.22	0.14	0.01
DPMC ($r = 400, k = 20$)	0.99	0.22	0.14	0.01
95% Credible Intervals				
CMC ($r = 200, k = 10$)	(0.96, 1.10)	(0.21, 0.23)	(0.11, 0.12)	(0.01, 0.01)
CMC ($r = 400, k = 10$)	(0.95, 1.10)	(0.21, 0.23)	(0.11, 0.12)	(0.01, 0.01)
CMC ($r = 200, k = 20$)	(0.94, 1.02)	(0.22, 0.23)	(0.13, 0.14)	(0.01, 0.01)
CMC ($r = 400, k = 20$)	(0.94, 1.02)	(0.22, 0.24)	(0.13, 0.13)	(0.01, 0.01)
DISK ($r = 200, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
DISK ($r = 400, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
DISK ($r = 200, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
DISK ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
WASP ($r = 200, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
WASP ($r = 400, k = 10$)	(0.80, 1.27)	(0.18, 0.24)	(0.11, 0.14)	(0.01, 0.01)
WASP ($r = 200, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
WASP ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.12, 0.18)	(0.01, 0.01)
DPMC ($r = 200, k = 10$)	(0.80, 1.27)	(0.17, 0.24)	(0.10, 0.15)	(0.01, 0.01)
DPMC ($r = 400, k = 10$)	(0.80, 1.27)	(0.17, 0.25)	(0.10, 0.15)	(0.01, 0.01)
DPMC ($r = 200, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.11, 0.18)	(0.01, 0.01)
DPMC ($r = 400, k = 20$)	(0.82, 1.16)	(0.17, 0.26)	(0.11, 0.18)	(0.01, 0.01)

versions in Simulations 1 and 2 in terms of effective sample size comparisons.

REFERENCES

- Banerjee, S., B. P. Carlin, and A. E. Gelfand (2014). *Hierarchical Modeling and Analysis for Spatial Data*. CRC Press.
- Finley, A. O., H. Sang, S. Banerjee, and A. E. Gelfand (2009). Improving the performance of predictive process modeling for large datasets. *Computational Statistics & Data Analysis* 53(8), 2873–2884.
- Harville, D. A. (1997). *Matrix algebra from a statistician's perspective*, Volume 1. Springer.
- Plummer, M., N. Best, K. Cowles, and K. Vines (2006). Coda: convergence diagnosis and output analysis for mcmc. *R news* 6(1), 7–11.
- Quiñonero-Candela, J. and C. E. Rasmussen (2005). A unifying view of sparse approximate gaussian process regression. *Journal of Machine Learning Research* 6(Dec), 1939–1959.
- R Development Core Team (2017). *R: A Language and Environment for Statistical Computing*. Vienna, Austria: R Foundation for Statistical Computing.
- Zhang, Y., J. C. Duchi, and M. J. Wainwright (2015). Divide and conquer kernel ridge regression: a distributed algorithm with minimax optimal rates. *Journal of Machine Learning Research* 16, 3299–3340.

TABLE 8

Inference on the values of spatial surface and response at the locations in $\mathcal{S}_$ in Simulation 2 for the divide-and-conquer methods under random and grid-based partitioning. The estimation and prediction errors are defined in (99) and coverage and credible intervals are calculated pointwise for the locations in $\mathcal{S}_*$. The entries in the table are averaged over 10 simulation replications.*

	Est Err	Pred Err	95% CI Coverage		95% CI Length	
	GP	Y	GP	Y	GP	Y
Random Partitioning						
CMC ($r = 200, k = 10$)	0.56	0.64	0.38	0.39	0.10	0.10
CMC ($r = 400, k = 10$)	0.43	0.52	0.40	0.41	0.10	0.10
CMC ($r = 200, k = 20$)	0.58	0.67	0.27	0.28	0.07	0.07
CMC ($r = 400, k = 20$)	0.46	0.55	0.28	0.29	0.07	0.07
DISK ($r = 200, k = 10$)	0.00	0.01	1.00	0.97	0.54	0.45
DISK ($r = 400, k = 10$)	0.00	0.01	1.00	0.97	0.45	0.47
DISK ($r = 200, k = 20$)	0.00	0.01	1.00	0.97	0.52	0.43
DISK ($r = 400, k = 20$)	0.00	0.01	1.00	0.97	0.43	0.44
WASP ($r = 200, k = 10$)	0.55	0.64	0.96	0.96	0.42	0.42
WASP ($r = 400, k = 10$)	0.42	0.51	0.96	0.96	0.40	0.40
WASP ($r = 200, k = 20$)	0.58	0.67	0.96	0.96	0.43	0.43
WASP ($r = 400, k = 20$)	0.46	0.55	0.96	0.96	0.41	0.41
DPMC ($r = 200, k = 10$)	0.55	0.64	0.97	0.97	0.45	0.45
DPMC ($r = 400, k = 10$)	0.42	0.51	0.97	0.97	0.43	0.43
DPMC ($r = 200, k = 20$)	0.58	0.67	0.97	0.97	0.46	0.46
DPMC ($r = 400, k = 20$)	0.46	0.55	0.97	0.97	0.44	0.44
Grid-Based Partitioning						
CMC ($r = 200, k = 10$)	0.05	0.10	0.80	0.38	0.10	0.10
CMC ($r = 400, k = 10$)	0.04	0.10	0.85	0.37	0.10	0.10
CMC ($r = 200, k = 20$)	0.03	0.10	0.71	0.28	0.07	0.07
CMC ($r = 400, k = 20$)	0.03	0.10	0.70	0.28	0.07	0.07
DISK ($r = 200, k = 10$)	0.04	0.10	1.00	0.97	0.45	0.45
DISK ($r = 400, k = 10$)	0.04	0.10	1.00	0.96	0.42	0.42
DISK ($r = 200, k = 20$)	0.03	0.10	1.00	0.97	0.46	0.46
DISK ($r = 400, k = 20$)	0.03	0.10	1.00	0.96	0.44	0.44
WASP ($r = 200, k = 10$)	0.04	0.10	1.00	0.95	0.42	0.42
WASP ($r = 400, k = 10$)	0.04	0.10	1.00	0.94	0.40	0.40
WASP ($r = 200, k = 20$)	0.03	0.10	1.00	0.96	0.43	0.43
WASP ($r = 400, k = 20$)	0.03	0.10	1.00	0.95	0.41	0.41
DPMC ($r = 200, k = 10$)	0.04	0.10	1.00	0.97	0.45	0.45
DPMC ($r = 400, k = 10$)	0.04	0.10	1.00	0.96	0.43	0.43
DPMC ($r = 200, k = 20$)	0.03	0.10	1.00	0.97	0.46	0.46
DPMC ($r = 400, k = 20$)	0.03	0.10	1.00	0.97	0.44	0.44

TABLE 9

Parametric inference and prediction in SST data using the divide-and-conquer methods and MPP-based modeling with $r = 400, 600$ knots on $k = 300$ subsets. For parametric inference posterior medians are provided along with the 95% credible intervals (CIs) in the parentheses. Similarly, mean squared prediction errors (MSPEs) along with length and coverage of 95% predictive intervals (PIs) are presented. The upper and lower quantiles of 95% CIs and PIs are averaged over 10 simulation replications.

	β_0	β_1	σ^2	ϕ	τ^2
Parameter Estimate					
CMC ($r = 400, k = 300$)	32.37	-0.32	12.38	0.03	0.18
CMC ($r = 600, k = 300$)	32.36	-0.32	12.31	0.03	0.18
DISK ($r = 400, k = 300$)	32.33	-0.32	11.82	0.04	0.18
DISK ($r = 600, k = 300$)	32.33	-0.32	11.85	0.04	0.18
WASP ($r = 400, k = 300$)	32.33	-0.32	11.82	0.04	0.18
WASP ($r = 600, k = 300$)	32.33	-0.32	11.85	0.04	0.18
DPMC ($r = 400, k = 300$)	32.33	-0.32	11.82	0.04	0.18
DPMC ($r = 600, k = 300$)	32.33	-0.32	11.85	0.04	0.18
95% Credible Interval					
CMC ($r = 400, k = 300$)	(32.33, 32.4)	(-0.32, -0.32)	(12.37, 12.39)	(0.0339, 0.0340)	(0.18, 0.18)
CMC ($r = 600, k = 300$)	(32.33, 32.4)	(-0.32, -0.32)	(12.3, 12.31)	(0.0342, 0.0343)	(0.18, 0.18)
DISK ($r = 400, k = 300$)	(31.72, 32.93)	(-0.33, -0.31)	(11.24, 12.43)	(0.0373, 0.0412)	(0.18, 0.19)
DISK ($r = 600, k = 300$)	(31.72, 32.93)	(-0.33, -0.31)	(11.25, 12.45)	(0.0372, 0.0413)	(0.18, 0.19)
WASP ($r = 400, k = 300$)	(31.72, 32.93)	(-0.33, -0.31)	(11.22, 12.46)	(0.0372, 0.0413)	(0.18, 0.19)
WASP ($r = 600, k = 300$)	(31.72, 32.93)	(-0.33, -0.31)	(11.24, 12.47)	(0.0372, 0.0413)	(0.18, 0.19)
DPMC ($r = 400, k = 300$)	(31.72, 32.94)	(-0.33, -0.31)	(11.09, 12.55)	(0.0369, 0.0416)	(0.18, 0.19)
DPMC ($r = 600, k = 300$)	(31.72, 32.94)	(-0.33, -0.31)	(11.14, 12.56)	(0.0368, 0.0416)	(0.18, 0.19)
Predictions					
	MSPE	95% PI Coverage	95% PI Length		
CMC ($r = 400, k = 300$)	0.74	0.05	0.08		
CMC ($r = 600, k = 300$)	0.67	0.05	0.07		
DISK ($r = 400, k = 300$)	0.43	0.95	2.65		
DISK ($r = 600, k = 300$)	0.36	0.95	2.34		
WASP ($r = 400, k = 300$)	0.66	0.93	2.39		
WASP ($r = 600, k = 300$)	0.60	0.92	2.11		
DPMC ($r = 400, k = 300$)	0.66	0.95	2.67		
DPMC ($r = 600, k = 300$)	0.60	0.94	2.36		

TABLE 10

Run-time (in $\log_{10}$ seconds) of the non-distributed methods and distributed methods under the random and grid-based partitioning schemes in Simulations 1 and 2, where MPP prior is used on the subsets.

	INLA	LaGP	NNGP ($m = 10$)	NNGP ($m = 20$)	NNGP ($m = 30$) ($m = 30$)	LatticeKrig	GpGp	
	1.08	0.08	2.96	3.42	3.74	2.03	0.96	
	Vecchia ($m = 10$)	Vecchia ($m = 20$)	Vecchia ($m = 30$)	MPP ($r = 200$)	MPP ($r = 400$)			
	2.76	3.20	3.50	3.97	4.31			
	$k = 10$							
	CMC		DISK		WASP		DPMC	
	$r = 200$	$r = 400$	$r = 200$	$r = 400$	$r = 200$	$r = 400$	$r = 200$	$r = 400$
Random	3.18	3.18	3.18	3.18	3.20	3.20	3.18	3.18
Grid	3.18	3.18	3.18	3.18	3.20	3.20	3.18	3.18
	$k = 20$							
	CMC		DISK		WASP		DPMC	
	$r = 200$	$r = 400$	$r = 200$	$r = 400$	$r = 200$	$r = 400$	$r = 200$	$r = 400$
Random	3.17	3.17	3.17	3.17	3.20	3.20	3.17	3.17
Grid	3.17	3.17	3.17	3.17	3.20	3.20	3.17	3.17

TABLE 11

Run-time (in $\log_{10}$ hours) of laGP and the distributed methods in the sea surface temperature data analysis, where MPP prior is used on the subsets.

laGP	MPP, $r = 400$, $k = 300$				MPP, $r = 600$, $k = 300$			
	CMC	DISK	WASP	DPMC	CMC	DISK	WASP	DPMC
-1.32	1.67	1.67	1.67	1.67	1.69	1.69	1.69	1.69

TABLE 12

The ratio of effective sample sizes of the Markov chains produced on the subsets using the MPP prior and those obtained using the full data and the same MPP prior in Simulation 1 under random and grid-based partitioning. The effective sample sizes have been averaged over the parameter dimensions and over 10 simulation replications.

	β	σ^2	ϕ	τ^2	GP	Y
Random Partitioning						
$k = 10$ and $r = 200$	0.99	0.35	3.24	0.53	1.00	1.00
$k = 20$ and $r = 200$	1.00	0.61	3.92	0.40	1.00	1.00
$k = 10$ and $r = 400$	1.0	0.93	2.53	0.57	1.00	1.00
$k = 20$ and $r = 400$	1.11	1.45	2.88	0.43	1.00	1.00
Grid-Based Partitioning						
$k = 10$ and $r = 200$	1.00	0.34	3.37	0.55	1.00	1.00
$k = 20$ and $r = 200$	1.00	0.61	3.89	0.39	1.00	1.00
$k = 10$ and $r = 400$	1.11	1.00	2.44	0.55	1.00	1.00
$k = 20$ and $r = 400$	1.11	1.65	2.97	0.41	1.00	1.00

TABLE 13

The ratio of effective sample sizes of the Markov chains produced on the subsets using the MPP prior and those obtained using the full data and the same MPP prior in Simulation 2 under random and grid-based partitioning. The effective sample sizes have been averaged over the parameter dimensions and over 10 simulation replications.

	β	σ^2	ϕ	τ^2	GP	Y
Random Partitioning						
$k = 10$ and $r = 200$	0.98	0.46	1.32	3.45	0.93	1.00
$k = 20$ and $r = 200$	0.69	0.20	1.25	3.30	0.79	1.00
$k = 10$ and $r = 400$	0.89	1.65	1.81	2.84	0.91	1.00
$k = 20$ and $r = 400$	0.57	1.05	1.94	2.60	0.73	1.00
Grid-Based Partitioning						
$k = 10$ and $r = 200$	0.98	0.62	1.27	3.52	0.94	1.00
$k = 20$ and $r = 200$	0.65	0.24	1.33	3.32	0.79	1.00
$k = 10$ and $r = 400$	0.88	1.81	1.97	2.73	0.91	1.00
$k = 20$ and $r = 400$	0.63	0.94	2.05	2.70	0.77	1.00