

Online Supplement for “Sparse Pseudo-input Local Kriging for Large Spatial Datasets with Exogenous Variables”

by Babak Farmanesh and Arash Pourhabib

Appendix A Solving optimization problem (11)

Due to the convex objective function and affine constraints of optimization problem (11), the duality gap between the primal and dual problems of (11) is zero by Lagrange duality principle (Bazaraa et al., 2013). This allows us to transform the optimization problem (11) to an unconstrained optimization problem and maximize the Lagrangian of (11) instead,

$$\begin{aligned} \max_{\mathbf{u}_s(\mathbf{x}_*), \boldsymbol{\lambda}_s(\mathbf{x}_*)} \mathcal{L}(\mathbf{u}_s(\mathbf{x}_*), \boldsymbol{\lambda}_s(\mathbf{x}_*)) &= \mathbf{u}_s(\mathbf{x}_*)^T (\tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s + \text{diag}(\mathbf{K}_{\mathbf{X}_s \mathbf{X}_s} - \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s) + \sigma_s^2 \mathbf{I}_s) \mathbf{u}_s(\mathbf{x}_*) \\ &\quad - 2\mathbf{u}_s(\mathbf{x}_*)^T \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s - \sum_{i=1:|\mathbf{B}_s|} \lambda_{is}(\mathbf{x}_*) (\mathbf{u}_s(\mathbf{b}_i)^T \mathbf{y}_s - \mathcal{R}(\mathbf{b}_i)), \end{aligned} \quad (23)$$

where $|\mathbf{B}_s|$ is the number of all the control points located on the boundaries of subdomain Ω_s , and $\boldsymbol{\lambda}_s(\mathbf{x}_*) = [\lambda_{1s}(\mathbf{x}_*), \dots, \lambda_{|\mathbf{B}_s|s}(\mathbf{x}_*)]^T$ is the vector of the Lagrange multipliers.

Assuming $\mathbf{u}_s(\mathbf{x}_*)$ depends on the covariance between $\mathbf{x}_*$ and $\mathbf{X}_s$, and $\lambda_{is}(\mathbf{x}_*)$ depends on the covariance of $\mathbf{b}_i$ and $\mathbf{x}_*$, we write $\mathbf{u}_s(\mathbf{x}_*) = \mathbf{H}_s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s$ and $\lambda_{is}(\mathbf{x}_*) = \beta_{is} \tilde{k}_{\mathbf{b}_i \mathbf{x}_*}^s$ as suggested in (Park et al., 2011), where $\mathbf{H}_j$ is a squared matrix with size equal to the number of data points in Ω_s , and β_{is} is the Lagrange parameter associated with λ_{is} that does not depend on $\mathbf{x}_*$. Consequently, we rewrite Lagrangian (23) as

$$\begin{aligned} \max_{\mathbf{H}_s, \boldsymbol{\beta}_s} \mathcal{L}(\mathbf{H}_s, \boldsymbol{\beta}_s) &= \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s \mathbf{H}_s^T (\tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s + \text{diag}(\mathbf{K}_{\mathbf{X}_s \mathbf{X}_s} - \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s) + \sigma_s^2 \mathbf{I}_s) \mathbf{H}_s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s \\ &\quad - 2\tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s \mathbf{H}_s^T \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s - \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{B}_s}^s \boldsymbol{\beta}_s (\tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s \mathbf{H}_s^T \mathbf{y}_s - \mathbf{r}_s), \end{aligned} \quad (24)$$

where $\boldsymbol{\beta}_s$ is a diagonal matrix with diagonal elements $\beta_{1s}, \dots, \beta_{|\mathbf{B}_s|s}$, and $\mathbf{r}_s = [\mathcal{R}(\mathbf{b}_1), \dots, \mathcal{R}(\mathbf{b}_{|\mathbf{B}_s|})]^T$ is the vectors of boundary values of Ω_s .

Due to convexity of function (24) we can calculate the optimal values of $\mathbf{H}_s$ and $\boldsymbol{\beta}_s$ analytically

by writing out the first order necessary conditions,

$$\frac{d\mathcal{L}(\mathbf{H}_s, \beta_s)}{d\mathbf{H}_s} = 2(\mathbf{G}_s \mathbf{H}_s - \mathbf{I}_s) \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s - \mathbf{y}_s \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{B}_s}^s \beta_s \tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s = 0, \quad (25)$$

$$\frac{d\mathcal{L}(\mathbf{H}_s, \beta_s)}{d\beta_{is}} = \tilde{\mathbf{k}}_{\mathbf{b}_i \mathbf{X}_s} \mathbf{H}_s^T \mathbf{y}_j - r_{is} = 0 \quad \forall i \in [|\mathbf{B}_s|], \quad (26)$$

where $\mathbf{G}_s = (\tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s + \text{diag}(\mathbf{K}_{\mathbf{X}_s \mathbf{X}_s} - \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{X}_s}^s) + \sigma_s^2 \mathbf{I}_s)$, and r_{is} is the i^{th} element of the vector $\mathbf{r}_s$.

Reordering equation (25),

$$(\tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s + 0.5(\tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^j \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^j)^{-1} \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^j \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{B}_s}^j \beta_s \tilde{\mathbf{k}}_{\mathbf{B}_s \mathbf{x}_*}^j \mathbf{y}_s^T) \mathbf{G}_s^{-1} \mathbf{y}_s = \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s \mathbf{H}_s^T \mathbf{y}_s, \quad (27)$$

and evaluating it at the boundary locations gives the system of equations with $|\mathbf{B}_s|$ equations and Lagrangian parameters,

$$(\tilde{\mathbf{k}}_{\mathbf{b}_i \mathbf{X}_s}^s + 0.5(\tilde{\mathbf{k}}_{\mathbf{b}_i \mathbf{X}_s}^s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{b}_i}^s)^{-1} \tilde{\mathbf{k}}_{\mathbf{b}_i \mathbf{X}_s}^s \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{B}_s}^s \beta_s \tilde{\mathbf{k}}_{\mathbf{B}_s \mathbf{b}_i}^s \mathbf{y}_s^T) \mathbf{G}_s^{-1} \mathbf{y}_s = r_{is} \quad \forall i \in [|\mathbf{B}_s|]. \quad (28)$$

After some simple matrix algebra, we obtain the solution to the system of linear equations (28),

$$\beta_s = \frac{\mathbf{I}_s(\mathbf{r}_s - \tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s \mathbf{G}_s^{-1} \mathbf{y}_s) \{[(\text{diag}(\tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{B}_s}^s))^{-1} (\tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s \tilde{\mathbf{K}}_{\mathbf{X}_s \mathbf{B}_s}^s)] \circ \mathbf{K}_{\mathbf{B}_s \mathbf{B}_s}^s\}^{-1}}{0.5 \mathbf{y}_s^T \mathbf{G}_s^{-1} \mathbf{y}_s}. \quad (29)$$

Using the values of β_s from (29), we can easily obtain the solution to $\mathbf{u}(\mathbf{x}_*)$ from (25),

$$\mathbf{u}_s^*(\mathbf{x}_*) = \mathbf{H}_s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s = \mathbf{G}_s^{-1} (\tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s + \mathbf{w}_s), \quad (30)$$

where $\mathbf{w}_s = 0.5(\tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{X}_s}^s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s)^{-1} \mathbf{y}_s \tilde{\mathbf{k}}_{\mathbf{x}_* \mathbf{B}_s}^s \beta_s \tilde{\mathbf{K}}_{\mathbf{B}_s \mathbf{X}_s}^s \tilde{\mathbf{k}}_{\mathbf{X}_s \mathbf{x}_*}^s$.

Appendix B Derivation of low-rank covariance approximation error

We follow the procedure proposed in (Smola and Schölkopf, 2000) to derive the low-rank covariance approximation error in each subdomain Ω_s . In this derivation, given the covariance function $\phi(\cdot, \cdot) : \Omega_s \times \Omega_s \rightarrow \mathbb{R}$ as a symmetric positive semidefinite kernel, we intend to approximate the kernel

$\phi(\mathbf{x}, \cdot) : \Omega_s \rightarrow \mathbb{R}^{\Omega_s}$ centered at $\mathbf{z} \in \Omega_s$ as a linear combination of kernels centered at each element of $\mathbf{X}_s$, i.e.,

$$\phi(\mathbf{z}, \cdot) \approx \sum_{i \in [m_s]} c_i \phi(\tilde{\mathbf{x}}_i, \cdot). \quad (31)$$

To this end, let $\mathcal{H}$ be a reproducing kernel Hilbert space (RKHS) that is defined as the space of functions constructed by the span of $\phi(\mathbf{x}, \cdot)$ centered at a finite number of elements of Ω_s , i.e.,

$$\left\{ \sum_{i \in [n]} \alpha_i \phi(\mathbf{x}_i, \cdot) : n \in \mathbb{N}, \mathbf{x}_i \in \Omega_s, c_i \in \mathbf{R} \right\}.$$

$\mathcal{H}$ is also equipped with the inner product

$$\left\langle \sum_{i \in [n_1]} \alpha_i \phi(\mathbf{x}_i, \cdot), \sum_{j \in [n_2]} \beta_j \phi(\mathbf{x}_j, \cdot) \right\rangle_{\mathcal{H}} = \sum_{i \in [n_1]} \sum_{j \in [n_2]} \alpha_i \beta_j \phi(\mathbf{x}_i, \mathbf{x}_j), \quad (32)$$

which, for any function $f \in \mathcal{H}$, induces the norm

$$\|f\|_{\mathcal{H}}^2 = \langle f, f \rangle_{\mathcal{H}}. \quad (33)$$

Given such $\mathcal{H}$, a natural criterion to find an approximation for the covariance function is to minimize the norm of function $\phi(\mathbf{z}, \cdot) - \sum_{i \in [m_s]} c_i \phi(\tilde{\mathbf{x}}_i, \cdot)$, which belongs to $\mathcal{H}$, that is

$$\min_{\mathbf{c}} \left\| \phi(\mathbf{z}, \cdot) - \sum_{i \in [m_s]} c_i \phi(\tilde{\mathbf{x}}_i, \cdot) \right\|_{\mathcal{H}}^2, \quad (34)$$

where $\mathbf{c} = [c_1, \dots, c_{m_s}]^T$. Assuming $\phi(\mathbf{z}, \mathbf{z}) = h$, objective function (34) can be expanded after plugging in (32) and (33) as

$$\min_{\mathbf{c}} h - 2\mathbf{c}^T \mathbf{k}_{\tilde{\mathbf{X}}_s \mathbf{z}} + \mathbf{c}^T \mathbf{K}_{\tilde{\mathbf{X}}_s \tilde{\mathbf{X}}_s} \mathbf{c},$$

which has the solution $\mathbf{c}^* = \mathbf{K}_{\tilde{\mathbf{X}}_s \tilde{\mathbf{X}}_s}^{-1} \mathbf{k}_{\tilde{\mathbf{X}}_s \mathbf{z}}$. Therefore, the approximation of $\phi(\mathbf{z}, \cdot)$ becomes $\mathbf{k}_{\mathbf{z} \tilde{\mathbf{X}}_s} \mathbf{K}_{\tilde{\mathbf{X}}_s \tilde{\mathbf{X}}_s}^{-1} \mathbf{k}_{\tilde{\mathbf{X}}_s \mathbf{z}}$, and the error of covariance approximation becomes

$$h - \mathbf{k}_{\mathbf{z} \tilde{\mathbf{X}}_s} \mathbf{K}_{\tilde{\mathbf{X}}_s \tilde{\mathbf{X}}_s}^{-1} \mathbf{k}_{\tilde{\mathbf{X}}_s \mathbf{z}}.$$

We finally note that using $\mathbf{z} = \mathbf{x}_i$ for all $\mathbf{x}_i \in \mathbf{X}_s$ in objective function (34) and minimizing the sum over all terms obtains $\mathbf{K}_{\mathbf{X}_s \tilde{\mathbf{X}}_s} \mathbf{K}_{\tilde{\mathbf{X}}_s \tilde{\mathbf{X}}_s}^{-1} \mathbf{K}_{\tilde{\mathbf{X}}_s \mathbf{X}_s}$, which is the low-rank approximation of $\mathbf{K}_{\mathbf{X}_s \mathbf{X}_s}$ in equation (7).

Appendix C Proof of Theorems

C.1 Proof of Proposition 1

Proof. For any $i \in [m_s]$, let $\mathbf{u}_i$ denote the covariance vector between $\mathbf{z}$ and the first i elements of $\tilde{\mathbf{X}}_s$, and let $\mathbf{v}_i$ denote the covariance vector between the $(i+1)^{\text{th}}$ element of $\tilde{\mathbf{X}}_s$ and the first i elements of $\tilde{\mathbf{X}}_s$. That is, $\mathbf{u}_i = [\phi(\mathbf{z}, \tilde{\mathbf{x}}_1), \dots, \phi(\mathbf{z}, \tilde{\mathbf{x}}_i)]^T$, and $\mathbf{v}_i = [\phi(\tilde{\mathbf{x}}_{i+1}, \tilde{\mathbf{x}}_1), \dots, \phi(\tilde{\mathbf{x}}_{i+1}, \tilde{\mathbf{x}}_i)]^T$. Also let $\mathbf{K}_i$ denote the covariance matrix between the first i elements of $\tilde{\mathbf{X}}_s$ themselves. We now prove by induction on i . For the base case, i.e, $i = 1$, the claim clearly holds,

$$\mathbb{E}_{\Omega_s}(\mathbf{u}_1^T \mathbf{K}_1^{-1} \mathbf{u}_1) = \mathbb{E}_{\Omega_s}(\phi(\mathbf{z}, \tilde{\mathbf{x}}_1) \phi(\tilde{\mathbf{x}}_1, \tilde{\mathbf{x}}_1)^{-1} \phi(\mathbf{z}, \tilde{\mathbf{x}}_1)) = \frac{1}{h} \mathbb{E}_{\Omega_s}(\phi^2(\mathbf{z}, \tilde{\mathbf{x}}_1)) = \frac{1}{h} \mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}')). \quad (35)$$

Suppose the claim holds for $m_s - 1$, we show that it also holds for m_s . Expanding $\mathbf{u}_{m_s}^T \mathbf{K}_{m_s}^{-1} \mathbf{u}_{m_s}$ gives

$$\mathbf{u}_{m_s}^T \mathbf{K}_{m_s}^{-1} \mathbf{u}_{m_s} = \begin{bmatrix} \mathbf{u}_{m_s-1}^T & \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}) \end{bmatrix} \begin{bmatrix} \mathbf{K}_{m_s-1} & \mathbf{v}_{m_s-1} \\ \mathbf{v}_{m_s-1}^T & h \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{u}_{m_s-1} \\ \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}) \end{bmatrix} \quad (36a)$$

$$= \begin{bmatrix} \mathbf{u}_{m_s-1}^T & \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}) \end{bmatrix} \begin{bmatrix} \mathbf{K}_{m_s-1}^{-1} + c \mathbf{K}_{m_s-1}^{-1} \mathbf{v}_{m_s-1} \mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} & -c \mathbf{K}_{m_s-1}^{-1} \mathbf{v}_{m_s-1} \\ -c \mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} & c \end{bmatrix} \begin{bmatrix} \mathbf{u}_{m_s-1} \\ \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}) \end{bmatrix} \quad (36b)$$

$$= \mathbf{u}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1} + \frac{(\mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1})^2 + \phi^2(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}) - 2 \mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1} \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s})}{c} \quad (36c)$$

$$= \mathbf{u}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1} + \frac{(\mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1} - \phi(\mathbf{z}, \tilde{\mathbf{x}}_{m_s}))^2}{c} \quad (36d)$$

$$\geq \mathbf{u}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1}. \quad (36e)$$

where equality (36b) follows from the block matrix inversion lemma (Hager, 1989), and $c = (h - \mathbf{v}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{v}_{m_s-1})^{-1}$, which is always non-negative.

By (36) and the induction step,

$$\mathbb{E}_{\Omega_s}(\mathbf{u}_{m_s}^T \mathbf{K}_{m_s}^{-1} \mathbf{u}_{m_s}) \geq \mathbb{E}_{\Omega_s}(\mathbf{u}_{m_s-1}^T \mathbf{K}_{m_s-1}^{-1} \mathbf{u}_{m_s-1}) \geq \frac{1}{h} \mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}')). \quad (37)$$

□

C.2 Proof of Theorem 1

First, we prove the following lemma

Lemma 3. *For the random variables $z_1 \sim \mathcal{U}(a, a + e)$ and $z_2 \sim \mathcal{U}(b, b + e)$, where $a \leq b$ and $a, b, e \geq 0$, define $v = (z_1 - z_2)^2$. Then $\mathbb{E}_v(\exp(-cv)) \leq \mathbb{E}_v(\exp(-cv) \mid a = b)$ for any $c > 0$.*

Proof. Let $z = z_1 - z_2$, then by convolution of probability distributions, we have:

$$f_z(t) = \int_{-\infty}^{+\infty} f_{z_1}(t + z_2) f_{z_2}(z_2) dz_2 = \frac{1}{e} \int_b^{b+e} f_{z_1}(t + z_2) dz_2, \quad (38)$$

where the last equation follows from the fact that $f_{z_2} = \frac{1}{e}$ if $b \leq z_2 \leq b + e$. Note that the integrand $f_{z_1}(z + z_2)$ is zero unless $a \leq t + z_2 \leq a + e$, which implies $a - t \leq z_2 \leq a + e - t$. Figure 6 shows the region defined by $a - t \leq z_2 \leq a + e - t$ and $b \leq z_2 \leq b + e$, for the case that $a + b < e$ and $a + e > b$. In the both cases, integration (38) can be calculated as follows:

$$f_z(t) = \begin{cases} \frac{1}{e^2} \int_{a-e-t}^t dz_2 & a - b - e \leq t < a - b \\ \frac{1}{e^2} \int_t^{a-e} dz_2 & a - b \leq t \leq a - b + e \end{cases} = \begin{cases} \frac{1}{e^2} (t + b - a + e) & a - b - e \leq t < a - b \\ \frac{-1}{e^2} (t + b - a - e) & a - b \leq t \leq a - b + e \end{cases} \quad (39)$$

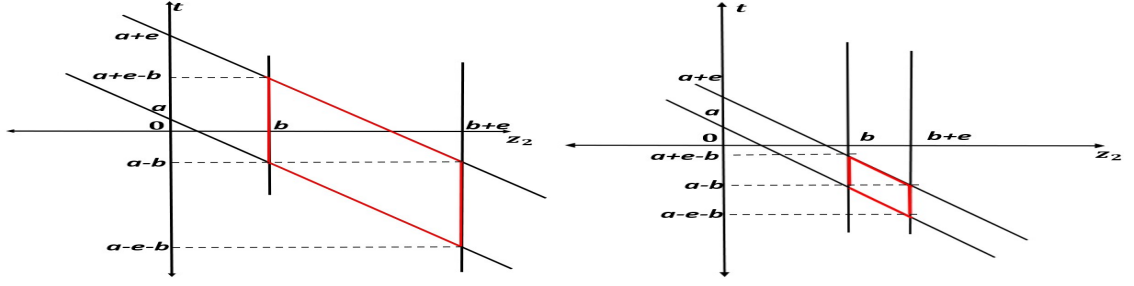


Figure 6: The region defined by $a - t \leq z_2 \leq a + e - t$ and $b \leq z_2 \leq b + e$. Left panel corresponds to the case when $a + b > e$ and right panel corresponds to the case when $a + e < b$.

Hence, $F_v(t) = p(v \leq t) = p(z^2 \leq t) = p(\sqrt{t} \leq z \leq \sqrt{t})$ can be written as

$$F_v(t) = \begin{cases} \frac{2\sqrt{t}}{e^2}(a - b + e) & 0 \leq \sqrt{t} < b - a, \\ \frac{1}{e^2}(2\sqrt{t}e - t - (a - b)^2) & b - a \leq \sqrt{t} < a - b + e, \\ 1 - \frac{1}{2e^2}(\sqrt{t} + a - b - e)^2 & a - b + e \leq \sqrt{t} \leq b - a + e. \end{cases} \quad (40)$$

Moreover, $G_v(t) = p(v \leq t \mid a = b) = p(z^2 \leq t \mid a = b) = p(\sqrt{t} \leq z \leq \sqrt{t} \mid a = b)$ can be derived by setting $a = b$ in CDF (40)

$$G_v(t) = \frac{1}{e^2}(2\sqrt{t}e - t) \quad 0 \leq \sqrt{t} \leq e. \quad (41)$$

Comparing $G_v(t)$ and $F_v(t)$ for all possible values of t gives

- $\sqrt{t} < 0$: $G_v(t) = F_v(t) = 0$.
- $0 \leq \sqrt{t} < b - a$: then $F_v(t) - G_v(t) = \frac{1}{e^2}(2\sqrt{t}(a - b) + t)$. Since $\sqrt{t} < b - a \Rightarrow t < \sqrt{t}(b - a) \Rightarrow t + \sqrt{t}(a - b) < 0 \Rightarrow t + 2\sqrt{t}(a - b) < 0 \Rightarrow F_v(t) - G_v(t) < 0 \Rightarrow F_v(t) < G_v(t)$.
- $b - a \leq \sqrt{t} < a - b + e$: then $F_v(t) - G_v(t) = -\frac{(a - b)^2}{e^2} < 0 \Rightarrow F_v(t) - G_v(t) < 0 \Rightarrow F_v(t) < G_v(t)$.
- $a - b + e \leq \sqrt{t} < e$: then $F_v(t) - G_v(t) = 1 - \frac{1}{2e^2}(\sqrt{t} + a - b - e)^2 + \frac{1}{e^2}(t - 2\sqrt{t}e)$.

Note that $\frac{e(F_v(t) - G_v(t))}{et} = \frac{1}{2e^2}(1 - \frac{a - b + e}{\sqrt{t}}) > 0$, and therefore, $F_v(t) - G_v(t)$ is a monotonically increasing function. Due to the monotonicity of $F_v(t) - G_v(t)$, the maximum occurs at e , so

$\max_t F_v(t) - G_v(t) = F_v(e) - G_v(e) = -\frac{(a-b)^2}{e^2} < 0$. Therefore, $F_v(t) - G_v(t) \leq F_v(e) - G_v(e) < 0 \Rightarrow F_v(t) \leq G_v(t)$.

- $e \leq \sqrt{t} < b - a + e$: in this case $G_v(t)$ is always 1, hence, $F_v(t) \leq G_v(t)$.
- $b - a + e \leq \sqrt{t}$: in this case $G_v(t) = F_v(t) = 1$

Therefore, we can conclude that

$$p(v \leq t) \leq p(v \leq t \mid a = b) \quad \forall t \in \mathbb{R} \Rightarrow p(-cv \geq t') \leq p(-cv \geq t' \mid a = b) \quad \forall t' \in \mathbb{R} \text{ and } c > 0,$$

which implies that random variable $(-cv)$ is stochastically less than random variable $(-cv \mid a = b)$, i.e., $-cv \preceq_{st} -cv \mid a = b$. Consequently, the expectation of any non-decreasing function of these two variables are ordered, i.e., $\mathbb{E}_v(\exp(-cv)) \leq \mathbb{E}_v(\exp(-cv) \mid a = b)$ for any $c > 0$. $\square$

To proceed to the proof of Theorem 1, we use the following characterization for the cutting hyperplanes and subdomains. Assuming that the cutting hyperplanes are equidistant with distant $W = L/S$ from each other, we can characterize the ℓ^{th} $\in [S - 1]$ cutting hyperplane on Ω with respect to k^{th} primary axis of $\mathbb{R}^p$ using the vector of angles $\boldsymbol{\theta} = \{\theta_1, \dots, \theta_p\} \setminus \{\theta_k\}$,

$$H_{\boldsymbol{\theta}, k, W, \ell} = \{\mathbf{x} \in \Omega \mid x_k - \sum_{j \in [p] \setminus \{k\}} \tan(\theta_j) x_j - \ell W = 0\} \quad \forall \ell \in [S - 1]. \quad (42)$$

Note that this cutting hyperplane is orthogonal to the axis k only if $\boldsymbol{\theta} = \mathbf{0}$, that is $\theta_j = 0$ for $j \in [p] \setminus \{k\}$.

Denoting, respectively, the hyperplanes containing the “bottom” and the “top” faces of Ω as

$$H_{\boldsymbol{\theta}, k, W, 0} = \{\mathbf{x} \in \Omega \mid x_k = 0\} \quad \text{and} \quad H_{\boldsymbol{\theta}, k, W, S} = \{\mathbf{x} \in \Omega \mid x_k - L = 0\},$$

we define the s^{th} subdomain as the intersection of area between two consecutive hyperplanes and Ω , specifically,

$$\Omega_{\boldsymbol{\theta}, k, W, s} = \{\mathbf{x} \in \Omega \mid \min_{\mathbf{x}' \in H_{\boldsymbol{\theta}, k, W, s-1}} \|\mathbf{x} - \mathbf{x}'\|_2 \leq W \quad \text{and} \quad \min_{\mathbf{x}' \in H_{\boldsymbol{\theta}, k, W, s}} \|\mathbf{x} - \mathbf{x}'\|_2 \leq W\}, \quad (43)$$

where $\|\cdot\|_2$ denotes the Euclidean norm.

Proof of Theorem 1. Let $\mathbf{x}_{\{k\}} = \{x_1, \dots, x_p\} \setminus \{x_k\}$ for any $\mathbf{x} \in \Omega$. Then, based on how each $\Omega_{\boldsymbol{\theta}, k, W, s}$ in (43) is constructed and considering the distribution of the data points in Ω according to (16), all variables $x_j \in \mathbf{x}_{\{i\}}$ are independent and have the uniform distribution $\mathcal{U}(0, L)$. Moreover, by the definition of the hyperplanes in (42), and given $\mathbf{x}_{\{k\}}$, the corresponding values of the variable x_k on the hyperplanes $H_{\boldsymbol{\theta}, k, W, s-1}$ and $H_{\boldsymbol{\theta}, k, W, s}$ are

$$\sum_{j \in [p] \setminus \{k\}} \tan(\theta_j) x_j + (s-1)w \quad \& \quad \sum_{j \in [p] \setminus \{k\}} \tan(\theta_j) x_j + sw. \quad (44)$$

Therefore, the conditional distribution $x_k | \mathbf{x}_{\{k\}}$ in the parallelogram subdomain $\Omega_{\boldsymbol{\theta}, k, W, s}$ has a uniform distribution whose support is bounded by the values calculated in (44). Consequently, given a parallelogram subdomain $\Omega_{\boldsymbol{\theta}, k, W, s}$, for any $\mathbf{x} \in \Omega_{\boldsymbol{\theta}, k, W, s-1}$,

$$x_j \sim \mathcal{U}(0, L) \quad \forall j \in [p] \setminus \{k\}, \quad (45a)$$

$$x_i | \mathbf{x}^i \sim \mathcal{U} \left(\sum_{j \in [p] \setminus \{k\}} \tan(\theta_j) x_j + (s-1)w, \sum_{j \in [p] \setminus \{k\}} \tan(\theta_j) x_j + sw \right). \quad (45b)$$

Now that we have the distribution (45), we expand $\mathbb{E}_{\Omega_{\boldsymbol{\theta}, k, W, s}}(\phi(\mathbf{x}, \mathbf{x}'))$ by conditioning, that is

$$\mathbb{E}_{\Omega_{\boldsymbol{\theta}, k, W, s}}(\phi(\mathbf{x}, \mathbf{x}')) = \mathbb{E}_{\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}} \left(\mathbb{E}_{x_k, x'_k}(\phi(\mathbf{x}, \mathbf{x}') \mid \mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}) \right) \quad (46a)$$

$$= \mathbb{E}_{\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}} \left(\exp \left(- \sum_{j \in [p] \setminus \{k\}} \gamma_j (x_j - x'_j)^2 \right) \mathbb{E}_{x_k, x'_k} \left(\exp(-\gamma_k (x_k - x'_k)^2) \mid \mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}} \right) \right) \quad (46b)$$

$$= \mathbb{E}_{\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}} (g(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}) h(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}})). \quad (46c)$$

Note that the function $g(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}})$ is always positive and independent of $\boldsymbol{\theta}$, and function $h(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}})$ is positive that attains its maximum for any given $\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}$ at $\boldsymbol{\theta} = \mathbf{0}$ by Lemma (3). Therefore, $\boldsymbol{\theta} = \mathbf{0}$,

$$g(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}) h(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}) \leq g(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}}) h(\mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}} \mid \boldsymbol{\theta} = \mathbf{0}) \quad \forall \mathbf{x}_{\{k\}}, \mathbf{x}'_{\{k\}},$$

which results in

$$\mathbb{E}_{\Omega_{\boldsymbol{\theta}, k, W, s}}(\phi(\mathbf{x}, \mathbf{x}')) \leq \mathbb{E}_{\Omega_{\boldsymbol{\theta}, k, W, s}}(\phi(\mathbf{x}, \mathbf{x}') \mid \boldsymbol{\theta} = \mathbf{0})$$

$$\Rightarrow \arg \max_{\boldsymbol{\theta}} \mathbb{E}_{\Omega_{\boldsymbol{\theta},k,W,s}}(\phi(\mathbf{x}, \mathbf{x}')) = \mathbf{0}.$$

□

C.3 Proof of Theorem 2

First, we prove the following lemma

Lemma 4. $\mathbb{E}_{z_1, z_2} \left(\exp(-c(z_1 - z_2)^2) \right) = \int_0^{b^2} \exp(-ct) \left(\frac{1}{b\sqrt{t}} - \frac{1}{b^2} \right) dt$, where $z_1, z_2 \stackrel{i.i.d}{\sim} \mathcal{U}(a, a+b)$.

Proof. Let $v = (z_1 - z_2)^2$, then Philip (2007) shows that v has the following PDF:

$$f_v(t) = \frac{1}{\sqrt{t}b} - \frac{1}{b^2} \quad \forall 0 \leq t \leq b^2;$$

therefore,

$$\begin{aligned} \mathbb{E}_{z_1, z_2} \left(\exp(-c(z_1 - z_2)^2) \right) &= \mathbb{E}_v \left(\exp(-cv) \right) = \\ &= \int_0^{b^2} \exp(-ct) f_s(t) dt = \int_0^{b^2} \exp(-ct) \left(\frac{1}{b\sqrt{t}} - \frac{1}{b^2} \right) dt. \end{aligned}$$

□

Proof of Theorem 2. By the assumptions of uniform distribution of points in Ω (16), and independence of the dimensions due to geometry of $\Omega_{\mathbf{0},k,W,s}$, for any $\mathbf{x} \in \Omega_{\mathbf{0},i,W,s}$,

$$x_k \sim \mathcal{U}((s-1)W, sW) \quad \& \quad x_j \sim \mathcal{U}(0, L) \quad \forall j \in [p] \setminus \{k\}. \quad (47)$$

Letting $G_k = \mathbb{E}_{\Omega_{\mathbf{0},k,W,s}}(\phi(\mathbf{x}, \mathbf{x}'))$, and using distribution (47),

$$G_k = \mathbb{E}_{x_k} \left(\exp(-\gamma_k(x_k - x'_k)^2) \right) \prod_{j \in [p] \setminus \{k\}} \mathbb{E}_{x_j} \left(\exp(-\gamma_j(x_j - x'_j)^2) \right) \quad (48a)$$

$$= \mathbb{E}_{v_k} \left(\exp(-\gamma_k v_k) \right) \prod_{j \in [p] \setminus \{k\}} \mathbb{E}_{v_j} \left(\exp(-\gamma_j v_j) \right) \quad (48b)$$

$$= \left(\int_0^{W^2} \exp(-\gamma_k t) \left(\frac{1}{W\sqrt{t}} - \frac{1}{W^2} \right) dt \right) \left(\prod_{j \in [p] \setminus \{k\}} \left(\int_0^{L^2} \exp(-\gamma_j t) \left(\frac{1}{L\sqrt{t}} - \frac{1}{L^2} \right) dt \right) \right) \quad (48c)$$

$$= \left(\int_0^{W^2} g_k^W(t) dt \right) \left(\prod_{j \in [p] \setminus \{k\}} \left(\int_0^{L^2} g_j^L(t) dt \right) \right), \quad (48d)$$

where equality (48a) follows from the independence of dimensions in each $\Omega_{\mathbf{0},i,W,s}$, equalities (48b) and (48c) follow from Lemma (4) with $f_{v_k}(t) = \frac{1}{\sqrt{t}W} - \frac{1}{W^2}$ $0 \leq t \leq W^2$ and $f_{v_j}(t) = \frac{1}{\sqrt{t}L} - \frac{1}{L^2}$ $0 \leq t \leq L^2$, and $g_\ell^m(t) = \exp(-\gamma_\ell t)(\frac{1}{m\sqrt{t}} - \frac{1}{m^2})$ in (48d).

To show that $G_p - G_k \geq 0$ for any $k \in [p]$, We first expand $G_p - G_k$,

$$\begin{aligned} G_p - G_k &= \left(\int_0^{W^2} g_p^W(t) dt \right) \left(\prod_{j \in [p] \setminus \{p\}} \left(\int_0^{L^2} g_j^L(t) dt \right) \right) - \left(\int_0^{W^2} g_k^W(t) dt \right) \left(\prod_{j \in [p] \setminus \{k\}} \left(\int_0^{L^2} g_j^L(t) dt \right) \right) \\ &= \left(\prod_{j \in [p] \setminus \{k,p\}} \left(\int_0^{L^2} g_j^L(t) dt \right) \right) \left(\int_0^{W^2} g_p^W(t) dt \int_0^{L^2} g_k^L(t) dt - \int_0^{W^2} g_k^W(t) dt \int_0^{L^2} g_p^L(t) dt \right) = A * B. \end{aligned}$$

Note that A is always positive, since each $\int_0^{L^2} g_j^L(t) dt$ is the expectation of the random variable $\exp(-\gamma_j v_j)$ which is positive. Hence, it is enough to show that B is positive. Expanding B further,

$$B = \left(\int_0^{W^2} g_p^W(t) dt \right) \left(\int_0^{W^2} g_k^L(t) dt + \int_{W^2}^{L^2} g_k^L(t) dt \right) - \left(\int_0^{W^2} g_k^W(t) dt \right) \left(\int_0^{W^2} g_p^L(t) dt + \int_{W^2}^{L^2} g_p^L(t) dt \right) \quad (49a)$$

$$\begin{aligned} &= \int_{t_k:0}^{W^2} \int_{t_p:0}^{W^2} g_p^W(t_k) g_k^L(t_p) dt_k dt_p + \int_{t_k:0}^{W^2} \int_{t_p:W^2}^{L^2} g_p^W(t_k) g_k^L(t_p) dt_k dt_p \\ &\quad - \int_{t_k:0}^{W^2} \int_{t_p:0}^{W^2} g_k^W(t_k) g_p^L(t_p) dt_k dt_p - \int_{t_k:0}^{W^2} \int_{t_p:W^2}^{L^2} g_k^W(t_k) g_p^L(t_p) dt_k dt_p \end{aligned} \quad (49b)$$

$$\begin{aligned} &= \int_{t_k:0}^{W^2} \int_{t_p:0}^{W^2} (\exp(-\gamma_p t_k - \gamma_k t_p) - \exp(-\gamma_k t_k - \gamma_p t_p)) \left(\frac{1}{W\sqrt{t_k}} - \frac{1}{W^2} \right) \left(\frac{1}{L\sqrt{t_p}} - \frac{1}{L^2} \right) dt_k dt_p \\ &\quad + \int_{t_k:0}^{W^2} \int_{t_p:W^2}^{L^2} (\exp(-\gamma_p t_k - \gamma_k t_p) - \exp(-\gamma_k t_k - \gamma_p t_p)) \left(\frac{1}{W\sqrt{t_k}} - \frac{1}{W^2} \right) \left(\frac{1}{L\sqrt{t_p}} - \frac{1}{L^2} \right) dt_k dt_p \end{aligned} \quad (49c)$$

$$= \int_{t_k:0}^{W^2} \int_{t_p:0}^{W^2} c(t_k, t_p) dt_k dt_p + \int_{t_k:0}^{W^2} \int_{t_p:W^2}^{L^2} c(t_k, t_p) dt_k dt_p. \quad (49d)$$

Note that for any member of set

$$\{(W, L, t_p, t_k, \gamma_p, \gamma_k) \mid 0 < W < L, 0 < \gamma_k < \gamma_p, 0 \leq t_k \leq W^2, W^2 \leq t_p \leq L^2\}, \quad (50)$$

we have

$$\left(\frac{1}{w\sqrt{t_k}} - \frac{1}{w^2} \right) \left(\frac{1}{L\sqrt{t_p}} - \frac{1}{L^2} \right) > 0, \quad (51)$$

and also

$$(-\gamma_p t_k - \gamma_k t_p) - (-\gamma_k t_k - \gamma_p t_p) = (\gamma_p - \gamma_k)(t_p - t_k) > 0, \quad (52)$$

where the latter results in

$$\exp(-\gamma_p t_k - \gamma_k t_p) - \exp(-\gamma_k t_k - \gamma_p t_p) > 0. \quad (53)$$

Therefore, by (51) and (53), the integrand $c(t_k, t_p)$ in (49d) is positive for any member of set (50), so is integral $\int_{t_k:0}^{W^2} \int_{t_p:W^2}^{L^2} c(t_k, t_p) dt_k dt_p$. Hence, to complete the proof we need to show integral $\int_{t_k:0}^{w^2} \int_{t_p:0}^{w^2} c(t_k, t_p) dt_k dt_p$ in (49d) is also positive. To show this, we expand the integral,

$$\int_{t_k:0}^{W^2} \int_{t_p:0}^{W^2} c(t_k, t_p) dt_k dt_p = \int_{t_k:0}^{W^2} \int_{t_p:t_k}^{W^2} c(t_k, t_p) dt_k dt_p + \int_{t_p:0}^{W^2} \int_{t_k:t_p}^{W^2} c(t_k, t_p) dt_p dt_k \quad (54a)$$

$$= \int_{t_k:0}^{W^2} \int_{t_p:t_k}^{W^2} c(t_k, t_p) dt_k dt_p + \int_{t_k:0}^{W^2} \int_{t_p:t_k}^{W^2} c(t_p, t_k) dt_k dt_p = \int_{t_k:0}^{W^2} \int_{t_p:t_k}^{W^2} (c(t_k, t_p) + c(t_p, t_k)) dt_k dt_p \quad (54b)$$

$$= \frac{1}{wL} \int_{t_k:0}^{W^2} \int_{t_p:t_k}^{W^2} (\exp(-\gamma_p t_k - \gamma_k t_p) - \exp(-\gamma_k t_k - \gamma_p t_p)) \left(\frac{1}{\sqrt{t_k}} - \frac{1}{\sqrt{t_p}} \right) \left(\frac{1}{W} - \frac{1}{L} \right) dt_k dt_p. \quad (54c)$$

Similar to (50)-(53), for any member of set

$$\{(W, L, t_p, t_k, \gamma_p, \gamma_k) \mid 0 < W < L, 0 < \gamma_k < \gamma_p, 0 \leq t_k \leq W^2, t_k \leq t_p \leq W^2\}, \quad (55)$$

we have $(\frac{1}{\sqrt{t_k}} - \frac{1}{\sqrt{t_p}}) > 0$, $(\frac{1}{W} - \frac{1}{L}) > 0$, and $(\exp(-\gamma_p t_k - \gamma_k t_p) - \exp(-\gamma_k t_k - \gamma_p t_p)) > 0$. Hence the integrand in (54c) is positive for any member of set (55), so is integral (54c), and the proof is complete. □

Appendix D A simulation study on the relation between expected error (15) and $\mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}'))$

Consider the squared exponential Gaussian kernel $\phi(x, x') = \exp(-\gamma(x - x')^2)$ with $\gamma > 0$ defined on

$$\Omega_s = \{x \in \mathbb{R} | a \leq x \leq a + b\} \quad (56)$$

with uniform sampling distribution

$$x \sim \mathcal{U}(a, a + b) \quad \forall x \in \Omega_s. \quad (57)$$

To have a general simulation study, we need the following lemma.

Lemma 5. $\mathbb{E}_{z_1, z_2} \left(\exp(-c(z_1 - z_2)^2) \right)$, where $z_1, z_2 \stackrel{i.i.d}{\sim} \mathcal{U}(a, a + b)$, is a monotonically decreasing function of c and b .

Proof. We need to show that $\nabla g(b, c) = [\frac{\partial g(b, c)}{\partial b}, \frac{\partial g(b, c)}{\partial c}]^T < 0$ for all $[b, c]^T > 0$, where

$$g(b, c) = \mathbb{E}_{z_1, z_2} \left(\exp(-c(z_1 - z_2)^2) \right) = \int_0^{b^2} \exp(-ct) \left(\frac{1}{b\sqrt{t}} - \frac{1}{b^2} \right) dt$$

by Lemma 4.

We can write $\frac{\partial g(b, c)}{\partial b}$ as

$$\frac{\partial g(b, c)}{\partial b} = \frac{1}{b^2} \int_0^{b^2} \exp(-ct) \left(\frac{2}{b} - \frac{1}{\sqrt{t}} \right) dt \quad (58a)$$

$$= \frac{1}{b^2} \left(\left[\exp(-ct) \left(\frac{2t}{b} - 2\sqrt{t} \right) \right]_0^{b^2} - \int_0^{b^2} -c \exp(-ct) \left(\frac{2t}{b} - 2\sqrt{t} \right) dt \right) \quad (58b)$$

$$= \frac{2c}{b^2} \int_0^{b^2} \exp(-ct) \left(\frac{t}{b} - \sqrt{t} \right) dt, \quad (58c)$$

where equalities (58a) and (58b) follow from the Leibniz integral differentiation and the integration by part rules, respectively. It is easy to check that integrand $\exp(-ct) \left(\frac{t}{b} - \sqrt{t} \right)$ is always negative for any member of set $\{(b, c, t) \mid 0 < b, 0 < c, 0 \leq t \leq b^2\}$; therefore, we always have $\frac{\partial g(b, c)}{\partial b} < 0$.

Moreover, for $\frac{\partial g(b,c)}{\partial c}$,

$$\frac{\partial g(b,c)}{\partial c} = \int_0^{b^2} -t \exp(-ct) \left(\frac{1}{b\sqrt{t}} - \frac{1}{b^2} \right) dt = \frac{-1}{b} \int_0^{b^2} t \exp(-ct) \left(\frac{1}{\sqrt{t}} - \frac{1}{b} \right) dt.$$

It is again easy to check that the integrand $t \exp(-ct) (\frac{1}{\sqrt{t}} - \frac{1}{b})$ is positive for any member of set $\{(b, c, t) \mid 0 < b, 0 < c, 0 \leq t \leq b^2\}$. Therefore, $\frac{\partial g(b,c)}{\partial c}$ is always negative. $\square$

By Lemma 5, expectation function

$$\mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}')) = \mathbb{E}_{x, x'}(\exp(-2\gamma(x - x')^2)) \quad (59)$$

is a monotonically decreasing function of γ and b . This means that there are only two ways to increase expectation $\mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}'))$, which are either decreasing γ or decreasing b . The approximation of expected error function (15) on domain (56) and sampling distribution (57) for varying values of γ and b and a fixed value of m_s using a heat map plot is shown in Figure 7. We observe that as the values of γ or b decrease, or equivalently, $\mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}'))$ increases, the approximation of the expected error function decreases.

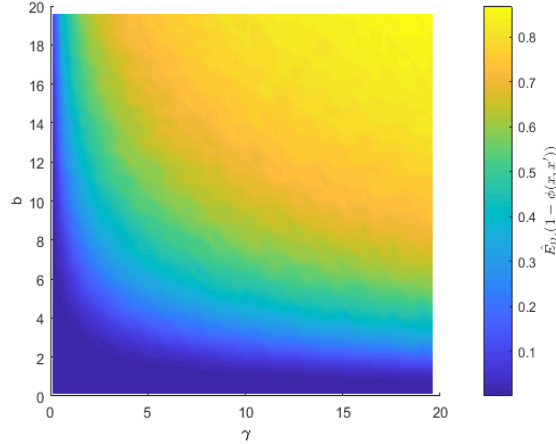


Figure 7: Heat map of the approximation of expected error function (15) on domain (56) and sampling distribution (57) for varying values of γ and b and a fixed value of m_s

Our simulation study can be used to infer a more general case. Consider the covariance function as $\phi(\mathbf{x}, \mathbf{x}') = \exp(-\sum_{i=1}^p \gamma_k(x_k - x'_k))$ defined on Ω_s as a p -dimensional hyper-rectangle with side lengths $b_1, \dots, b_p$ with a uniform sampling distribution, i.e., $x_k \sim \mathcal{U}(a_k, a_k + b_k) \quad \forall \mathbf{x} \in \Omega_s$. With

this setup, we can write

$$\mathbb{E}_{\Omega_s}(\phi^2(\mathbf{x}, \mathbf{x}')) = \prod_{i=1}^p \mathbb{E}_{x_k, x'_k}(\exp(-2\gamma_k(x_k - x'_k))), \quad (60)$$

which is a monotonic function in each b_i and γ_i by lemma 5. Therefore, our simulation results are valid for this generalized case as well.

Finally, we present some intuition behind the theoretical results in Section 3. The reason why the direction $\mathbf{a}$, found by solving optimization problem (20), results in a better covariance approximation in each subdomain can be visually perceived for a two-dimensional domain. Suppose we can partition the domain of two-dimensional function $f(\mathbf{x}) = \cos(0.05x_1 + 0.1x_2)$ by cutting orthogonal to either of three directions $[1, 0]$, $[0.43, 0.9]$, or $[0, 1]$, where direction $[0.43, 0.9]$ is the direction of the fastest covariance decay obtained by optimizing (20). Figure 8 shows the 3-D presentations of three local functions created by cutting orthogonal to each direction. We observe that the local functions created by cutting orthogonal to the desired direction have a less fluctuating behaviour compared to those of directions $[1, 0]$ and $[0, 1]$. That the function has less fluctuation allows a random point on the local functions of Figure 8b to have (on average) higher correlation to its neighboring data points. Therefore, we can obtain a better approximation of local covariance structures by using the same number of pseudo data points located in each subdomain.

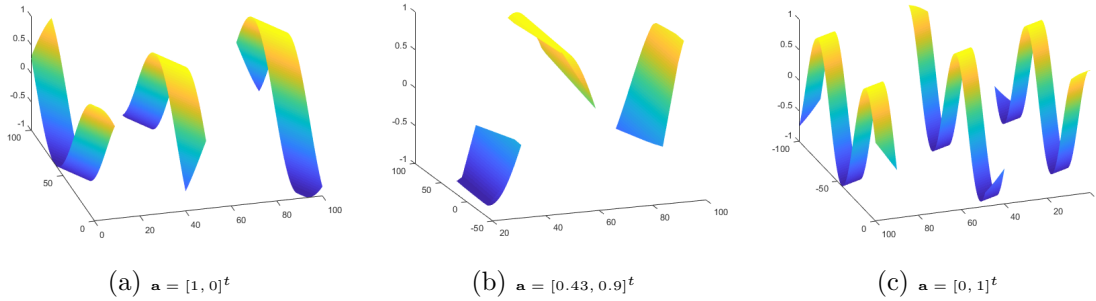


Figure 8: Local functions created by cutting orthogonal to directions $[1, 0]$, $[0.43, 0.9]$ (solution of (19)), and $[0, 1]$ on a synthetic dataset

Appendix E Solving optimization problem (20)

Let first write the partial derivatives of objective function in (20),

$$\frac{\partial \mathcal{L}(\bar{\mathbf{a}})}{\partial a_k} = -\mathbf{y}_n^T (\mathbf{K}_n^{\bar{\mathbf{a}}} + \sigma^2 \mathbf{I}_n)^{-1} \frac{\partial \mathbf{K}_n^{\bar{\mathbf{a}}}}{\partial a_k} (\mathbf{K}_n^{\bar{\mathbf{a}}} + \sigma^2 \mathbf{I}_n)^{-1} \mathbf{y}_n + \text{tr}((\mathbf{K}_n^{\bar{\mathbf{a}}} + \sigma^2 \mathbf{I}_n)^{-1} \frac{\partial \mathbf{K}_n^{\bar{\mathbf{a}}}}{\partial a_k}), \quad (61)$$

where $\frac{\partial \mathbf{K}_n^{\bar{\mathbf{a}}}}{\partial a_k}$ is the matrix of element-wise derivatives with respect to the k^{th} element of $\bar{\mathbf{a}}$. Note that each element of $\frac{\partial \mathbf{K}_n^{\bar{\mathbf{a}}}}{\partial a_k}$ involves the term $\frac{1}{\sqrt{1-\bar{\mathbf{a}}^T \bar{\mathbf{a}}}}$. Therefore, the gradient of the objective function in (20) does not exist on the boundary of the feasible region, i.e., $\nabla \mathcal{L}(\bar{\mathbf{a}}) \rightarrow \infty$ as $\bar{\mathbf{a}}^T \bar{\mathbf{a}} \rightarrow 1$. Therefore, to avoid an undefined gradient on the boundary, we modify the optimization by making the feasible region slightly tighter, i.e.,

$$\begin{aligned} \min_{\bar{\mathbf{a}}} \quad & \mathcal{L}(\bar{\mathbf{a}}) = \mathbf{y}_n^T (\mathbf{K}_n^{\bar{\mathbf{a}}} + \sigma^2 \mathbf{I}_n)^{-1} \mathbf{y}_n + \log |\mathbf{K}_n^{\bar{\mathbf{a}}} + \sigma^2 \mathbf{I}_n| \\ \text{subject to} \quad & \bar{\mathbf{a}}^T \bar{\mathbf{a}} \leq 1 - \epsilon, \end{aligned} \quad (62)$$

where ϵ is a very small number. In our experiments, we set $\epsilon = 0.001$.

Due to the simple convex structure of constraint $\bar{\mathbf{a}}^T \bar{\mathbf{a}} \leq 1 - \epsilon$, i.e., a $d - 1$ -dimensional hypersphere, optimization (62) can be solved by the Projected Gradient Descent algorithm (Nesterov and Nemirovskii, 1994). In this projection algorithm, the $(j + 1)^{\text{th}}$ decent step is defined by

$$\bar{\mathbf{a}}^{j+1} = \mathcal{P}(\bar{\mathbf{a}}^j - \frac{\alpha}{\|\nabla \mathcal{L}(\bar{\mathbf{a}}^j)\|} \nabla \mathcal{L}(\bar{\mathbf{a}}^j)), \quad (63)$$

where $\frac{\alpha}{\|\nabla \mathcal{L}(\bar{\mathbf{a}}^j)\|}$ is a normalized length step, and

$$\begin{aligned} \mathcal{P}(\mathbf{z}) = \quad & \text{argmin}_{\mathbf{w}} \|\mathbf{w} - \mathbf{z}\| \\ \text{subject to} \quad & \mathbf{w}^T \mathbf{w} \leq 1 - \epsilon. \end{aligned} \quad (64)$$

$\mathcal{P}(\mathbf{z}) = \mathbf{z}$ when $\mathbf{z}^T \mathbf{z} \leq 1 - \epsilon$, otherwise the solution to $\mathcal{P}(\mathbf{z})$ occurs at the point that the line defined by $\mathbf{z}$ and the center of the hypersphere, $(\mathbf{0})$, crosses the boundary of the hypersphere, i.e, intersection of $\frac{w_1}{z_1} = \frac{w_2}{z_2} = \dots = \frac{w_{p-1}}{z_{p-1}}$ and $\mathbf{w}^T \mathbf{w} = 1 - \epsilon$. Therefore, the solution to $\mathcal{P}(\mathbf{z})$ has the

closed form,

$$\mathcal{P}(\mathbf{z}) = \begin{cases} \mathbf{z} & \mathbf{z}^T \mathbf{z} \leq 1 - \epsilon \\ [\frac{z_1}{\sqrt{\mathbf{z}^T \mathbf{z}}}, \dots, \frac{z_{p-1}}{\sqrt{\mathbf{z}^T \mathbf{z}}}]^T & \mathbf{z}^T \mathbf{z} > 1 - \epsilon. \end{cases} \quad (65)$$

Appendix F Practical Considerations

F.1 Creating boundaries, control points, and boundary functions

The focus of this section is on the practical implementation of SPLK, and therefore, the characterization of cutting hyperplanes differs from the discussion in Section 3.2. Here, instead of using a vector of angles corresponding to primary axes of input space, we use a given direction, which can be the solution to optimization (20) or any other arbitrary direction, to define the cutting hyperplanes.

Recall that in our partitioning policy all the cutting hyperplanes are parallel to each other, and therefore, orthogonal to a unique direction, which is characterized by a vector $\mathbf{a} = [a_1, \dots, a_p]^T$. Let $\mathbf{Z} = \{\mathbf{x}_i^T \mathbf{a} \mid \mathbf{x}_i \in \mathbf{X}\}$ denote the projection of all the input vectors onto $\mathbf{a}$. Next, consider the ordered set $\{z_1, \dots, z_{S-1}\}$, where $\min \mathbf{Z} < z_1$ and $z_{S-1} < \max \mathbf{Z}$, and $z_\ell < z_{\ell+1}$, for $\ell \in [S-1]$.

Given the set $\{z_1, \dots, z_{S-1}\}$ and direction $\mathbf{a}$, which is in fact the normal vector of all of the cutting hyperplanes, we define the ℓ^{th} cutting hyperplane orthogonal to $\mathbf{a}$ as $H_{\ell, \mathbf{a}} = \{\mathbf{x} \in \Omega \mid a_1 x_1 + \dots + a_p x_p = z_\ell\}$ for $\ell \in [S-1]$. We use the data points close to $H_{\ell, \mathbf{a}}$ to locate the control points. To this end, we first define $\Delta_\ell = \{\mathbf{x}_i \in \mathbf{X} \mid |\mathbf{x}_i^T \mathbf{a} - z_\ell| < \delta\}$ as the set of training data points whose Euclidean distance to $\mathcal{H}_{\ell, \mathbf{a}}$ is less than a predefined constant δ . Then, calculate the maximum and minimum of the k^{th} dimension of the data points in Δ_ℓ , respectively,

$$\tau_{1,k,\ell} = \max_{\mathbf{x}_i \in \Delta_\ell} \mathbf{x}_i^T \mathbf{e}_k \quad \text{and} \quad \tau_{0,k,\ell} = \min_{\mathbf{x}_i \in \Delta_\ell} \mathbf{x}_i^T \mathbf{e}_k, \quad (66)$$

where $\mathbf{e}_k$ is the unit vector along the k^{th} primary axis of the space for $k \in [p]$. As such, the set $\mathbf{V}_\ell = \{[\tau_{b,1,\ell}, \dots, \tau_{b,p,\ell}]^T \mid b = 0, 1\}$ characterizes the vertices of the hyper-rectangle inscribing Δ_ℓ . Next, we uniformly sample $Q > 0$ points from $\mathbf{V}_\ell$ and denote the set of all these points as $\mathbf{U}_\ell$. We

obtain the set of control points on $H_{\ell,\mathbf{a}}$ denoted as $\mathbf{C}_\ell$ by projecting the points in $\mathbf{U}_\ell$ on $\mathcal{H}_{\ell,\mathbf{a}}$,

$$\mathbf{C}_\ell = \{(z_\ell - \mathbf{u}^T \mathbf{a})\mathbf{a} + \mathbf{u} \mid \forall \mathbf{u} \in \mathbf{U}_\ell\}. \quad (67)$$

There are several ways to choose the width of each subdomain, i.e., $z_{\ell+1} - z_\ell$ for $\ell \in [S - 1]$. One way is to choose a fixed width for the subdomains; however, this approach results in subdomains with different numbers of local data points depending on their distribution on the domain. Also *adaptive mesh generation* techniques (Becker and Rannacher, 2001) can be used to vary the widths to balance the error among the subdomains. In Section 4, we use varying widths for the subdomains to balance the numbers of local data points across the subdomains. This approach helps us to control the computation time of the algorithm, because it is evenly distributed among the subdomains.

Furthermore, to impose connectivity on the optimization procedure discussed in Section 3.1, we need to specify the boundary values for each control point $\mathbf{c} \in \mathbf{C}_\ell$. To this end, we fit a boundary GPR over the hyper-rectangle defined by $\mathbf{V}_\ell$ using the data points in Δ_ℓ . We then use the predictive mean function of this GPR to determine the boundary values. Letting $\mathcal{R}_\ell(\cdot)$ denote as the predictive mean function of the GPR constructed by Δ_ℓ , the boundary value for each $\mathbf{c} \in \mathbf{C}_\ell$ is

$$\mathcal{R}_\ell(\mathbf{c}) = \mathbf{k}_{\mathbf{c}\Delta_\ell}(\mathbf{K}_{\Delta_\ell\Delta_\ell} + \sigma_\ell^2 \mathbf{I}_\ell)^{-1} \mathbf{y}_{\Delta_\ell}, \quad (68)$$

where $\mathbf{k}_{\mathbf{c}\Delta_\ell}$ is the covariance vector between the control point $\mathbf{c} \in \mathbf{C}_\ell$ and the neighboring data points in Δ_ℓ , and $\mathbf{K}_{\Delta_\ell\Delta_\ell}$ is the covariance matrix between the neighboring data points in Δ_ℓ themselves. In Section 3.1, with a slight abuse of notation, we denote $\mathcal{R}(\cdot)$ as a function that takes a control point as an input and returns $\mathcal{R}_\ell(\cdot)$, depending on the location of the control point. Note that since the set of neighboring data points Δ_ℓ is a small set, we use a full GPR to obtain functions 68.

Dataset	q	Time	MSE	NLPD
TCO	3	145.50	12.18	2.61
	4	145.61	12.15	2.61
	5	146.06	11.98	2.60
Levitus	3	134.48	25.50	2.60
	4	134.47	25.44	2.59
	5	135.36	25.25	2.59
Dasilva	3	157.62	0.42	4.01
	4	159.98	0.38	3.30
	5	167.79	0.38	3.05
Protein	2.2	147.53	17.41	2.67
	2.5	202.09	17.39	2.66
	3	3651.33	17.38	2.65

Table 1: Effect of q on efficiency of SPLK. $S = 30$ and $\kappa = 4$ across all the datasets

F.2 Control points density

As discussed in Section 3.4, we use a density parameter and the dimension of the boundary space, i.e., q and $p - 1$, to determine the number of control points to be uniformly located on each boundary. Notably, our experiments show that setting q to small values usually results in satisfactory performance, while increasing it does not significantly affect the prediction accuracy, but increases the computation burden, particularly in higher dimensional domains. The results of testing SPLK on our four datasets with varying values of q and all other parameters fixed are reported in Table 1. An increase in the value of q slightly improves the prediction accuracy in terms of NLPD and MSE. Moreover, as the dimension of the domain of data increases, an increase in the value of q results in much longer computation time.

F.3 Hyperparameter learning

Maximizing the marginal likelihood of the training data, $p(\mathbf{y})$, is a popular method for learning the hyperparameters of a model (Rasmussen and Williams, 2006). In SPLK, instead of one global marginal likelihood function, there are S local functions $p(\mathbf{y}_s)$, each of which can be trained independently. Recall that our local predictors are in fact SPGP predictors that consider pseudo-inputs as parameters of the model. Therefore, we have two types of parameters: one is the location of local pseudo-inputs and the other is the hyperparameters of the underlying covariance function. Maximizing the logarithm of the local SPGP marginal likelihood functions using gradient descent with respect to local pseudo-inputs and hyperparameters provides local optimal locations. Specifically,

the logarithm of the marginal likelihood of SPLK's s^{th} local model is

$$\log(p(\mathbf{y}_s)) = -\frac{1}{2} \log |\mathbf{G}_s| - \frac{1}{2} \mathbf{y}_s^T \mathbf{G}_s^{-1} \mathbf{y}_s - \frac{n_s}{2} \log 2\pi, \quad (69)$$

where $\mathbf{G}_s$ is the same as that of Section 3.1.

Moreover, we use anti-isotropic squared exponential function as the choice of our local covariance functions,

$$\phi(\mathbf{x}, \mathbf{x}') = C \exp \left(-(\mathbf{x} - \mathbf{x}')^T \mathbf{\Gamma} (\mathbf{x} - \mathbf{x}') \right), \quad (70)$$

where $\mathbf{\Gamma}$ is a diagonal matrix with length-scale parameters $\gamma_1, \dots, \gamma_p$ on the diagonal. This covariance function automatically determines the significance of predictors after training its parameters by minimizing local likelihood function (69).

Appendix G A simulation study on the performance of SPLK

In this section, we conduct a simulation study to further investigate the performance of SPLK comparing to the other competing algorithms in terms of MSE. As mentioned in Section 4.2, when the rates of covariance decay highly vary in different directions (similar to the Dataset Dasilva), SPLK can perform better than the competing algorithms considered in this study. This is because SPLK partitions the domain of data orthogonal to the direction of the fastest rate of covariance decay, which potentially reduces the degree of mismatch on the boundaries compared with the other directions.

To test this claim we generate 10,000 samples from a Gaussian process with covariance function (17) and highly different length scale parameters $\gamma_1 = 50$, $\gamma_2 = 10$, and $\gamma_3 = 0.001$. To this end, we first generate 10,000 vectors, $\mathbf{x}_i$, uniformly from the cube $[0, 5] \times [0, 5] \times [0, 5]$ and form the covariance matrix $\mathbf{K}_{\mathbf{X}\mathbf{X}}$. Then we draw 10,000 responses, y_i , using $\mathbf{K}_{\mathbf{X}\mathbf{X}}$ and add a noise to each response from distribution $\mathcal{N}(0, 4)$. Finally, we use 9,000 of these samples for training and 1,000 for testing.

For this simulated dataset, SPLK partitions the domain of data from the first direction which has the largest associated length scale parameter. Figures 9a and 9b show the performance of all the

competing algorithms in terms of MSE and NLPD versus computation time. As expected, due to the designed covariance structure, i.e., highly varying rates of covariance decay, SPLK outperforms the other competing algorithms in terms of MSE, while performs as well as PWK and PIC in terms of NLPD.

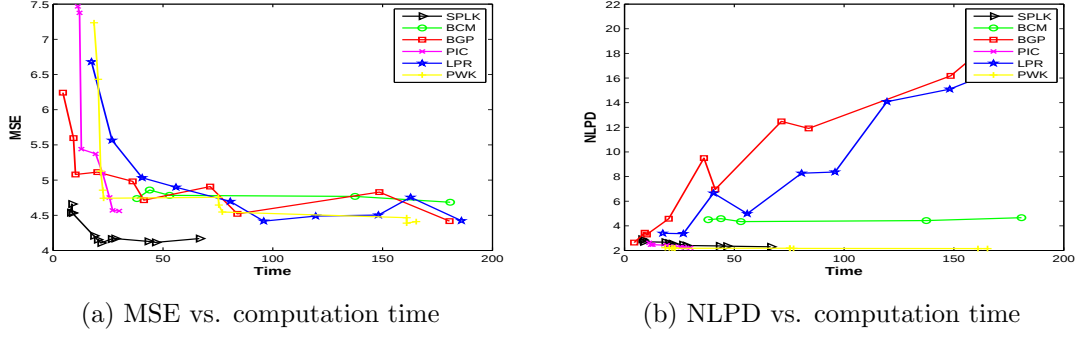


Figure 9: MSE and NLPD versus computation time. For SPLK, $q = 3$ and $k \in \{2, 4, 6, 8\}$. The value of parameter S is selected from the set $\{8, 16, 32, 64, 128, 256\}$.

References

- Becker, R. and R. Rannacher (2001). An optimal control approach to a posteriori error estimation in finite element methods. *Acta Numerica* 10, 1–102.
- Philip, J. (2007). The probability distribution of the distance between two random points in a box. *TRITA MAT 10(7)*.