

Supplement to: Fast Emulation, Modular Calibration, and Active Learning for Simulators with Functional Response

Grant Hutchings^{a,b,*}, Derek Bingham^b, Kellin Rumsey^a, Earl Lawrence^a

^a*Statistical Sciences, Los Alamos National Laboratory, NM, United States*

^b*Department of Statistics, Simon Fraser University, BC, Canada*

SM1. Detailed FlaGP algorithms

Detailed algorithms for the methods presented in the manuscript are provided here. The methods are implemented in the R package [FlaGP](#). Algorithm 1 in the manuscript outlines emulator fitting and prediction, while Algorithms 2 and 3 describe modular calibration without and with a discrepancy model. The supplemental algorithms below provide sufficient detail to reproduce the main computational steps.

Algorithm 1 Detailed FlaGP emulator fitting and prediction

Require: Input matrix $\mathbf{X} \in \mathbb{R}^{M \times d_x}$ and functional-output matrix $\mathbf{Z} \in \mathbb{R}^{d_y \times M}$.

Fit the FlaGP emulator

Precomputing 1: Data setup

- 1: Scale each column of $\mathbf{X}$ to the unit interval so that $\mathbf{X} \subset [0, 1]^{d_x}$.
- 2: Center each row of $\mathbf{Z}$ by subtracting its mean across the M simulator runs. Scale the centered matrix so that the total output variance is one, as in the manuscript.
- 3: Compute the singular value decomposition

$$\mathbf{Z} = \mathbf{U} \mathbf{D} \mathbf{V}^T.$$

- 4: Select the number of retained basis components p_η . If d_r is the r th singular value, the cumulative proportion of variance explained by the first p components is

$$\text{PVE}(p) = \frac{\sum_{r=1}^p d_r^2}{\sum_{r=1}^{\min(d_y, M)} d_r^2}.$$

A common choice is the smallest p for which $\text{PVE}(p) \geq 0.95$.

- 5: Let $\mathbf{U}_{p_\eta}$, $\mathbf{D}_{p_\eta}$, and $\mathbf{V}_{p_\eta}$ denote the retained SVD factors. Define

$$\mathbf{B} = \frac{\mathbf{U}_{p_\eta} \mathbf{D}_{p_\eta}}{\sqrt{M}} \quad \text{and} \quad \mathbf{W}(\mathbf{X}) = \sqrt{M} \mathbf{V}_{p_\eta}^T,$$

*Corresponding author. Phone: +1 505-667-2200. Present address: P.O. Box 1663, Los Alamos, NM 87545.

Email address: grant.hutchings@lanl.gov (Grant Hutchings)

so that $\mathbf{Z} \approx \mathbf{B}\mathbf{W}(\mathbf{X})$, with $\mathbf{B} \in \mathbb{R}^{d_y \times p_\eta}$ and $\mathbf{W}(\mathbf{X}) \in \mathbb{R}^{p_\eta \times M}$.

Precomputing 2: Estimate emulator hyperparameters

- 6: *Option 1 (BLHS)*: For each row $j = 1, \dots, p_\eta$ of $\mathbf{W}(\mathbf{X})$, use `laGP::blhs.loop` with b_{blhs} divisions per input dimension and r_{est} bootstrap replicates. For each replicate, estimate the lengthscale vector by empirical-Bayes MAP estimation. Use the componentwise median across replicates as $\hat{\mathbf{l}}_j$.
- 7: *Option 2 (stratified or random subsets)*: For each basis component, draw r_{est} subsets of size m_{est} using stratified or simple random sampling. Use `laGP::newGPsep` and `laGP::mleGPsep` to obtain an empirical-Bayes MAP lengthscale estimate for each subset, and use the componentwise median as $\hat{\mathbf{l}}_j$.
- 8: Increasing b_{blhs} in Option 1 or m_{est} in Option 2 increases the estimation subset size and computational cost. Subject to the available computational budget, choose the subset size and r_{est} large enough to obtain stable median estimates.

Precomputing 3: Transform the inputs

- 9: For each basis component $j = 1, \dots, p_\eta$, define

$$\tilde{\mathbf{X}}_j = \mathbf{X} \oslash \sqrt{\hat{\mathbf{l}}_j},$$

where $\oslash$ denotes elementwise division with column-wise recycling. Equivalently,

$$[\tilde{\mathbf{X}}_j]_{ik} = \frac{X_{ik}}{\sqrt{\hat{l}_{j,k}}}, \quad i = 1, \dots, M, \quad k = 1, \dots, d_x.$$

After this transformation, the lengthscales used for local prediction are fixed at $\mathbf{1}_{d_x}$.

Predict at a new input $\mathbf{x}^*$

- 10: For each basis component j , compute the transformed prediction input

$$\tilde{\mathbf{x}}_j^* = \mathbf{x}^* \oslash \sqrt{\hat{\mathbf{l}}_j}.$$

- 11: Use a fast nearest-neighbor search, such as `FNN::get.knnx`, to find the m nearest rows of $\tilde{\mathbf{X}}_j$ to $\tilde{\mathbf{x}}_j^*$. Denote the local transformed design and associated basis weights by

$$\tilde{\mathbf{X}}_{j,m}^* = \{\tilde{\mathbf{x}}_{j,1}^*, \dots, \tilde{\mathbf{x}}_{j,m}^*\} \quad \text{and} \quad \mathbf{w}_{j,m}^*.$$

- 12: Under the zero-mean GP specification in the manuscript, define the target-to-neighborhood correlation vector and local correlation matrix by

$$[\mathbf{p}_j^*]_l = \rho(\tilde{\mathbf{x}}_j^*, \tilde{\mathbf{x}}_{j,l}^*), \quad l = 1, \dots, m,$$

and

$$[\mathbf{P}_j^*]_{ll'} = \rho(\tilde{\mathbf{x}}_{j,l}^*, \tilde{\mathbf{x}}_{j,l'}^*) + g \mathbb{1}\{l = l'\}.$$

Let

$$\Psi_j = \mathbf{w}_{j,m}^{*T} (\mathbf{P}_j^*)^{-1} \mathbf{w}_{j,m}^*.$$

The predictive distribution of $w_j(\mathbf{x}^*)$ is Student- t with m degrees of freedom, predictive mean

$$\mu_j(\tilde{\mathbf{x}}_j^*) = \mathbf{p}_j^{*T} (\mathbf{P}_j^*)^{-1} \mathbf{w}_{j,m}^*,$$

and scale

$$s_j^2(\tilde{\mathbf{x}}_j^*) = \frac{\Psi_j}{m} [1 - \mathbf{p}_j^{*T} (\mathbf{P}_j^*)^{-1} \mathbf{p}_j^*].$$

13: Assemble

$$\boldsymbol{\mu}_w(\mathbf{x}^*) = [\mu_1(\tilde{\mathbf{x}}_1^*), \dots, \mu_{p_\eta}(\tilde{\mathbf{x}}_{p_\eta}^*)]^T$$

and

$$\boldsymbol{\Sigma}_w(\mathbf{x}^*) = \frac{m}{m-2} \text{diag}\{s_1^2(\tilde{\mathbf{x}}_1^*), \dots, s_{p_\eta}^2(\tilde{\mathbf{x}}_{p_\eta}^*)\}.$$

Then

$$\boldsymbol{\mu}_\eta(\mathbf{x}^*) = \mathbf{B} \boldsymbol{\mu}_w(\mathbf{x}^*), \quad \boldsymbol{\Sigma}_\eta(\mathbf{x}^*) = \mathbf{B} \boldsymbol{\Sigma}_w(\mathbf{x}^*) \mathbf{B}^T.$$

Alternatively, draw $w_j^{(s)}(\mathbf{x}^*)$ independently from the component-specific predictive distributions and form

$$\boldsymbol{\eta}^{(s)}(\mathbf{x}^*) = \sum_{j=1}^{p_\eta} \mathbf{b}_j w_j^{(s)}(\mathbf{x}^*).$$

If the simulator outputs were standardized before fitting, transform predictions back to the original output scale as needed.

Algorithm 2 Detailed modular calibration with the FlaGP emulator

Require: A fitted FlaGP emulator, field control inputs $\mathbf{X}_F$, field outputs $\mathbf{Y}^F = [\mathbf{y}^F(\mathbf{x}_1^F), \dots, \mathbf{y}^F(\mathbf{x}_n^F)]$, priors for $\boldsymbol{\theta}$ and σ_y^2 , and an optional discrepancy basis $\mathbf{K} \in \mathbb{R}^{d_y \times p_\delta}$.

- 1: The combined simulator input contains $d_x + d_t$ dimensions. For each emulator component j , let $\hat{\mathbf{l}}_j$ contain the lengthscales for both the control and calibration inputs. Use m_c neighbors for emulator predictions made during calibration.

Propose calibration parameters

- 2: Given the current transformed parameter vector

$$\boldsymbol{\psi}^{\text{curr}} = [\text{logit}(\boldsymbol{\theta}^{\text{curr}}), \log(\sigma_y^{2,\text{curr}})]^T,$$

propose

$$\boldsymbol{\psi}^{\text{new}} \sim N(\boldsymbol{\psi}^{\text{curr}}, \boldsymbol{\Sigma}_{\text{prop}}),$$

and transform back to obtain the proposed $\boldsymbol{\theta}$ and σ_y^2 .

Compute emulator predictive moments

- 3: For each field observation $i = 1, \dots, n$, form the combined input

$$\mathbf{u}_i = (\mathbf{x}_i^F, \boldsymbol{\theta}).$$

For each basis component $j = 1, \dots, p_\eta$, transform $\mathbf{u}_i$ using $\hat{\mathbf{l}}_j$, select its m_c nearest neighbors, and apply Algorithm 1 with m replaced by m_c to obtain the predictive mean $\mu_{w,j,i}$ and scale $s_{w,j,i}^2$.

- 4: Assemble the basis-weight predictive moments

$$\boldsymbol{\mu}_{w,i} = [\mu_{w,1,i}, \dots, \mu_{w,p_\eta,i}]^T,$$

$$\boldsymbol{\Sigma}_{w,i} = \frac{m_c}{m_c - 2} \text{diag}(s_{w,1,i}^2, \dots, s_{w,p_\eta,i}^2),$$

and project them to the functional-response space:

$$\boldsymbol{\mu}_{\eta,i} = \mathbf{B}\boldsymbol{\mu}_{w,i}, \quad \boldsymbol{\Sigma}_{\eta,i} = \mathbf{B}\boldsymbol{\Sigma}_{w,i}\mathbf{B}^T.$$

Fit the optional discrepancy model

- 5: **if** a discrepancy basis $\mathbf{K}$ is included **then**
- 6: Form the residual matrix using the emulator predictive means,

$$\mathbf{R}_y = \mathbf{Y}^F - [\boldsymbol{\mu}_{\eta,1}, \dots, \boldsymbol{\mu}_{\eta,n}].$$

- 7: Compute the projected discrepancy weights

$$\mathbf{V}^T = (\mathbf{K}^T \mathbf{K})^{-1} \mathbf{K}^T \mathbf{R}_y.$$

Fit p_δ independent full zero-mean GP models to the rows of $\mathbf{V}^T$, using $\mathbf{X}_F$ as the input matrix and, for example, `laGP::newGPsep` and `laGP::mleGPsep`.

- 8: At each observed control input $\mathbf{x}_i^F$, compute the discrepancy-weight predictive mean $\boldsymbol{\mu}_{v,i}$ and covariance $\boldsymbol{\Sigma}_{v,i}$. Project them to the functional-response space:

$$\boldsymbol{\mu}_{\delta,i} = \mathbf{K}\boldsymbol{\mu}_{v,i}, \quad \boldsymbol{\Sigma}_{\delta,i} = \mathbf{K}\boldsymbol{\Sigma}_{v,i}\mathbf{K}^T.$$

9: **else**

10: Set $\boldsymbol{\mu}_{\delta,i} = \mathbf{0}_{d_y}$ and $\boldsymbol{\Sigma}_{\delta,i} = \mathbf{0}_{d_y \times d_y}$ for all i .

11: **end if**

Evaluate the transformed-space log posterior

- 12: For computational efficiency, approximate the joint calibration likelihood by the product of the observation-specific marginal predictive densities. For each i , define

$$\boldsymbol{\mu}_i = \boldsymbol{\mu}_{\eta,i} + \boldsymbol{\mu}_{\delta,i},$$

$$\boldsymbol{\Sigma}_i = \boldsymbol{\Sigma}_{\eta,i} + \boldsymbol{\Sigma}_{\delta,i} + \sigma_y^2 \mathbf{I}_{d_y},$$

and evaluate

$$\mathbf{y}^F(\mathbf{x}_i^F) \sim N(\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i).$$

The efficient likelihood calculation in Appendix A of the manuscript reduces each observation-specific factorization to dimension p_η without discrepancy, or $p_\eta + p_\delta$ with discrepancy.

- 13: Evaluate

$$\log \pi(\boldsymbol{\psi} \mid \cdot) = \sum_{i=1}^n \log p\{\mathbf{y}^F(\mathbf{x}_i^F) \mid \boldsymbol{\theta}, \sigma_y^2\} + \log p(\boldsymbol{\theta}) + \log p(\sigma_y^2) + \log |J(\boldsymbol{\psi})|,$$

where the change-of-variables term is

$$\log |J(\boldsymbol{\psi})| = \sum_{k=1}^{d_t} \{\log \theta_k + \log(1 - \theta_k)\} + \log \sigma_y^2.$$

- 14: Apply the Metropolis–Hastings acceptance rule and repeat the proposal, prediction, optional discrepancy, and posterior-evaluation steps to obtain posterior samples.

SM2. Replacing points in a prediction neighborhood

When a candidate point replaces one member of a fixed-size nearest-neighbor prediction set, prediction error is not guaranteed to decrease, even when the candidate is closer to the prediction location than the point it replaces. This behavior is unsurprising in at least two settings:

- the replacement makes the prediction location an extrapolation point; or
- neighbors are selected under a distance metric that does not reflect the fitted anisotropic GP correlation structure.

The second setting can be illustrated with the scalar response

$$y(x_1, x_2) = x_1.$$

The response varies with x_1 but is constant in x_2 . Under the Gaussian correlation parameterization used in the manuscript, one therefore expects a shorter or finite lengthscale in the x_1 direction and a very long lengthscale in the x_2 direction. Input scaling consequently places greater importance on closeness in x_1 . Suppose prediction is required at $\mathbf{x}_p = (0.5, 0.5)$ using two neighbors, with current neighborhood points $\mathbf{x}^1 = (0.4, 0.4)$ and $\mathbf{x}^2 = (0.6, 1)$. Replacing $\mathbf{x}^2$ by the Euclidean-closer point $\mathbf{x}^3 = (0.7, 0.4)$ can increase prediction error because $\mathbf{x}^2$ is closer to the target in the important x_1 direction.

Prediction error can also increase when the neighborhood does not induce extrapolation and the distance is isotropic. Consider

$$y(x_1, x_2) = x_1 + x_2,$$

for which both inputs are equally important. Nearest neighbors are selected using Euclidean distance

$$d(\mathbf{x}, \mathbf{x}') = \left\{ \sum_{j=1}^2 (x_j - x'_j)^2 \right\}^{1/2}.$$

This ordering is equivalent to that induced by an isotropic Gaussian correlation function, which is monotone in squared Euclidean distance.

Figure [SM1](#) summarizes the example. The red $\times$ marks the prediction location, the black circles are the initial training points, and the blue points are the five current nearest neighbors. Candidate points closer to the prediction location than the most distant current neighbor can replace that neighbor. Green candidates decrease prediction error after replacement, whereas orange candidates increase it. All candidate additions remove the blue point at $(0.55, 0.24)$, yet there is no simple geometric rule separating the candidates that improve prediction from those that degrade it.

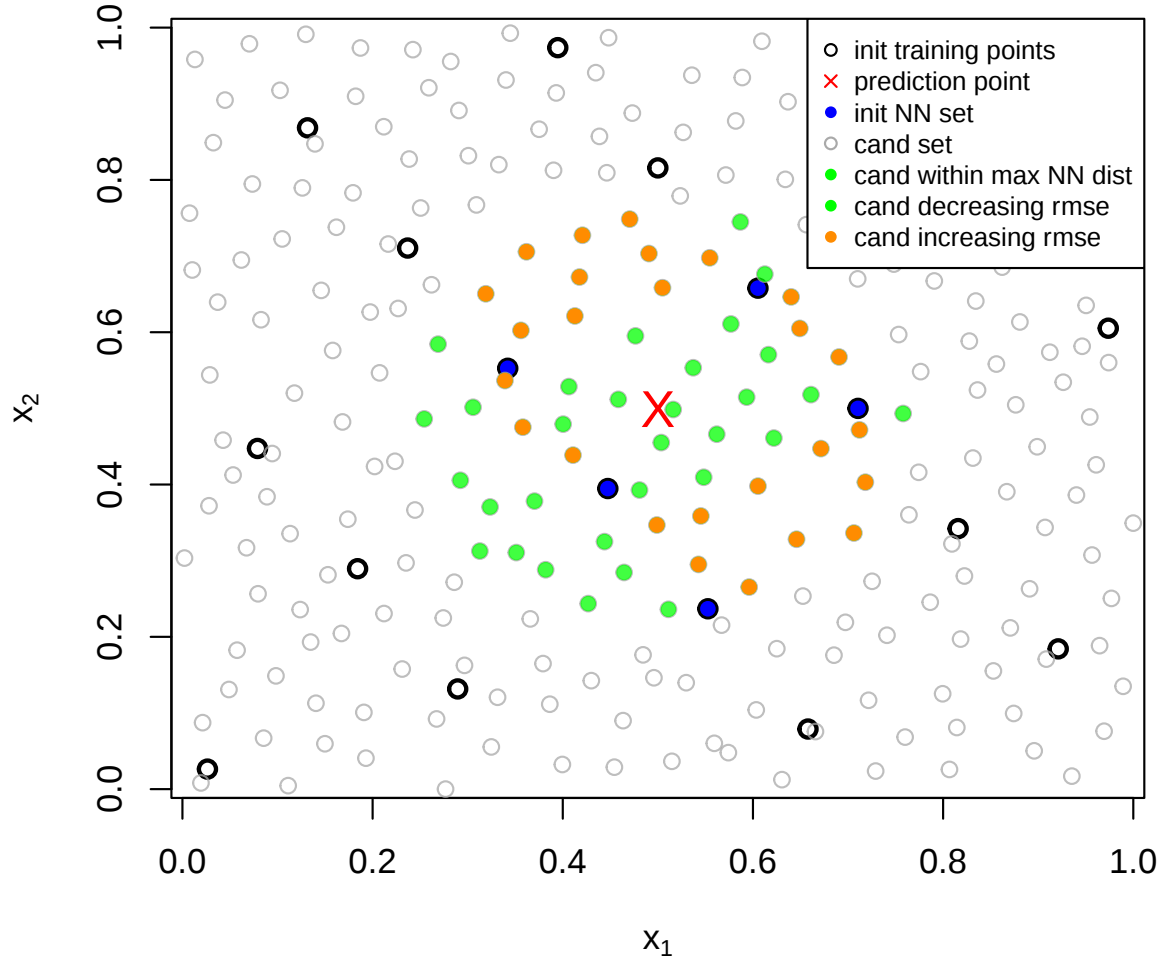


Figure SM1: Effect of replacing one point in a five-nearest-neighbor prediction set. The red $\times$ is the prediction location, black circles are the initial training points, and blue points form the current nearest-neighbor set. Candidate points inside the current maximum neighbor distance are colored green when the replacement decreases prediction error and orange when it increases prediction error.

SM3. Computational speedups from input scaling

A simulation study is used to illustrate the computational gains from input scaling and nearest-neighbor prediction relative to the default local approximate GP workflow. The manuscript identifies two benefits of input scaling. First, under an anisotropic correlation function, scaling makes ordinary nearest neighbors more closely aligned with the fitted GP correlation structure, avoiding repeated greedy neighborhood construction with the active-learning-Cohn (ALC) criterion. Second, global lengthscales are estimated once and reused rather than re-estimated at every prediction. This sacrifices some of laGP’s nonstationary flexibility, but the timing experiments below indicate that the combined speedup can approach two orders of magnitude for problem sizes similar to the application.

SM3.1. Neighborhood selection with laGP

The first experiment quantifies the additional cost of ALC neighborhood selection relative to simple nearest-neighbor selection. We report the wall-time ratio

$$\frac{t_{\text{ALC}}}{t_{\text{NN}}}.$$

We first fix $d_x = 10$ and generate 500 random datasets with ensemble sizes spanning $100 \leq M \leq 100000$. Input matrices and scalar outputs are generated independently and uniformly, and predictions are made sequentially at 10 locations using both ALC and nearest-neighbor neighborhood selection. Input scaling is not used, so local correlation parameters are estimated for each prediction.

A second set of 500 datasets uses $M = 20000$ and dimensions $1 \leq d_x \leq 100$, with five replicates at each dimension. Figure SM2 shows that ALC is generally three to five times slower than nearest-neighbor selection over a broad range of M when $d_x = 10$. The ratio is more variable as d_x changes and ranges from approximately 3 to 15. Vertical black lines indicate the application-scale setting $M = 20000$ and $d_x = 10$.

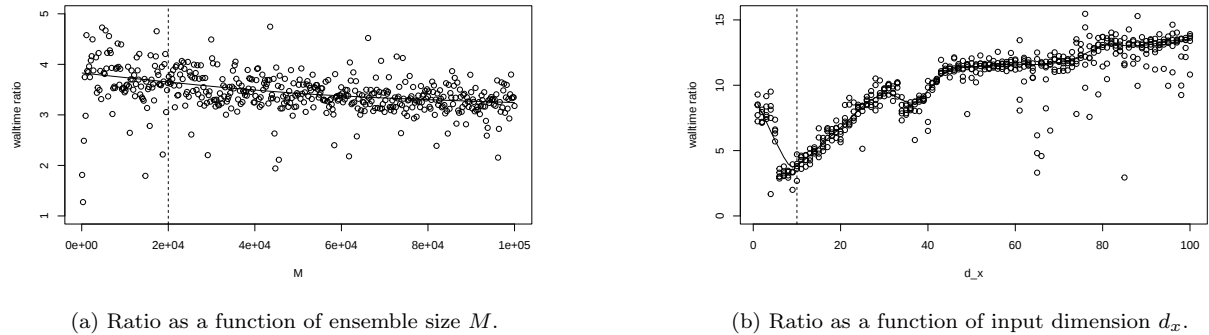


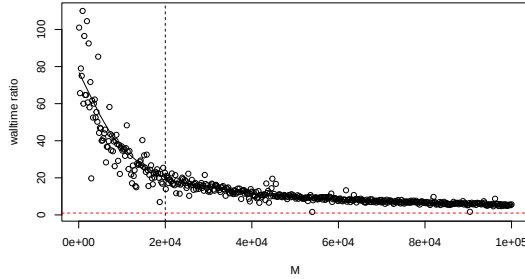
Figure SM2: Wall-time ratio $t_{\text{ALC}}/t_{\text{NN}}$ for ALC neighborhood selection relative to simple nearest-neighbor selection. For fixed d_x , the ratio decreases slightly with increasing M . For fixed M , it generally increases with d_x after approximately $d_x = 5$. Vertical black lines mark the approximate size of the application dataset.

SM3.2. Parameter estimation with laGP

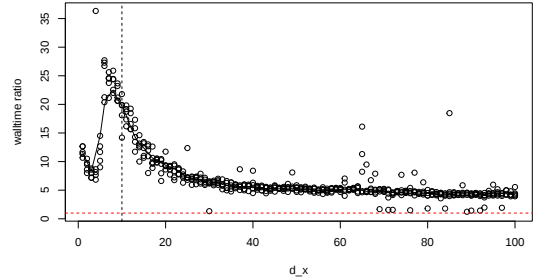
Repeated local parameter estimation can be costly when predictions are made thousands of times during MCMC. The second experiment compares nearest-neighbor prediction with local parameter estimation against nearest-neighbor prediction using pre-estimated global lengthscales. We report the wall-time ratio

$$\frac{t_{\text{NN+local estimation}}}{t_{\text{NN+fixed lengthscales}}}.$$

The results are shown in Figure SM3. As either M or d_x becomes large, the ratio approaches one because nearest-neighbor search increasingly dominates parameter-estimation time. At the application-scale setting $M = 20000$ and $d_x = 10$, however, pre-estimating the lengthscales still provides an approximately $20\times$ speedup. Combining this with an approximately $5\times$ speedup from nearest-neighbor rather than ALC neighborhood selection gives an approximate two-order-of-magnitude speedup. The one-time cost of global lengthscale estimation is not negligible, but for the application considered here it is small relative to the cost of repeated prediction during Bayesian calibration.



(a) Ratio as a function of ensemble size M .



(b) Ratio as a function of input dimension d_x .

Figure SM3: Wall-time ratio for nearest-neighbor prediction with repeated local parameter estimation relative to nearest-neighbor prediction using fixed, pre-estimated global lengthscales. The horizontal red line indicates a ratio of one, and vertical black lines mark the approximate size of the application dataset.