

Sequential Bayesian experimental design for calibration of expensive simulation models

Özge Sürer*

Department of Information Systems and Analytics, Miami University

and

Matthew Plumlee[†]

Department of Industrial Engineering and Management Sciences, Northwestern University

and

Stefan M. Wild[‡]

Applied Mathematics and Computational Research Division, Lawrence Berkeley National Laboratory

Abstract

Simulation models of critical systems often have parameters that need to be calibrated using observed data. For expensive simulation models, calibration is done using an emulator of the simulation model built on simulation output at different parameter settings. Using intelligent and adaptive selection of parameters to build the emulator can drastically improve the efficiency of the calibration process. The article proposes a sequential framework with a novel criterion for parameter selection that targets learning the posterior density of the parameters. The emergent behavior from this criterion is that exploration happens by selecting parameters in uncertain posterior regions while simultaneously exploitation happens by selecting parameters in regions of high posterior density. The advantages of the proposed method are illustrated using several simulation experiments and a nuclear physics reaction model.

Keywords: Acquisition, Active learning, Emulation, Gaussian process, Uncertainty quantification

*The authors gratefully acknowledge support from the NSF grant OAC 2004601

†The authors gratefully acknowledge support from the NSF grant DMS 1953111

‡This material was based upon work supported by the U.S. Department of Energy, Office of Science, Office of Advanced Scientific Computing Research, applied mathematics and SciDAC programs under Contract No. DE-AC02-05CH11231

1 Introduction

Across many engineering and science disciplines, simulation models are used to understand the properties and behaviors of complex systems. These simulation models often include uncertain or unknown input parameters that affect the mechanics of the simulation. Calibration, the topic of this article, leverages observed data from the system to infer parameters and create predictions of quantities of interest. Bayesian calibration is a particular brand of calibration that provides uncertainty quantification on parameters and predictions. Bayesian calibration involves constructing a posterior distribution using knowledge encapsulated in a prior distribution, which is typically a known, closed-form function of the parameters, and a likelihood, which requires the evaluation of the simulation model. The posterior distribution then reflects the relative probability of a set of parameters given both prior knowledge and the adherence to the observed data. In a standard Bayesian calibration process, Markov chain Monte Carlo (MCMC) computational techniques (Gelman et al., 2004) are used to generate samples from the posterior using the unnormalized posterior density which represents the posterior density up to some unknown scaling constant. MCMC requires many evaluations of the likelihood for convergence of the posterior distribution.

Bayesian calibration using direct MCMC becomes difficult when a simulation model is computationally expensive. For example, a single simulation evaluation can take hours or more in fields such as epidemiology (Yang et al., 2021) and nuclear physics (Sürer et al., 2022a). A popular solution is to build a computationally cheaper emulator to be used in place of the simulation model (Kennedy and O’Hagan, 2001; Higdon et al., 2004). An emulator is built via flexible statistical models such as Gaussian processes (GPs) (Rasmussen and Williams, 2005; Gramacy, 2020) or Bayesian additive trees (Chipman et al., 2010) using simulation outputs recorded at a design consisting of a set of parameters. After an emulator is constructed, it is used in place of the expensive simulation model throughout the parameter space.

Space-filling designs such as Latin Hypercube Sampling (LHS) or Sobol sequences are widely

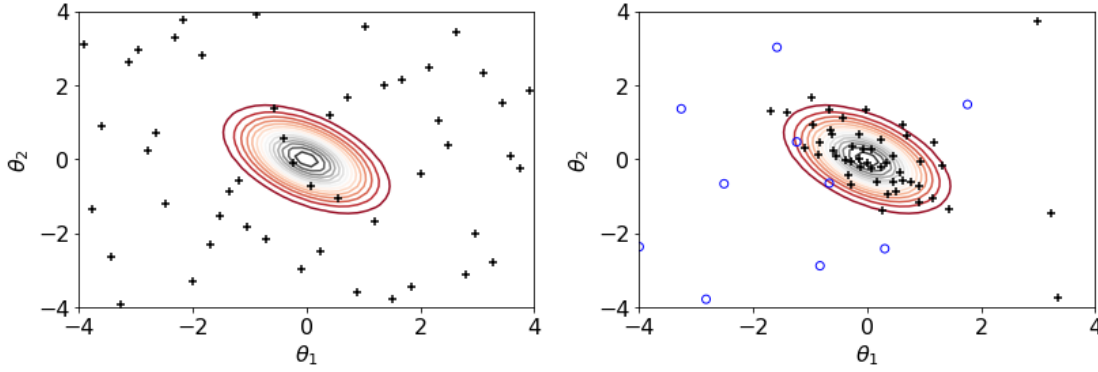


Figure 1: 50 samples (plus markers) using LHS (left) and the sequential sampling procedure introduced in this article (right). Blue circles illustrate 10 samples used to initiate the proposed procedure. The lines represent the level curves of the posterior.

used in the literature to generate designs for general-purpose emulator construction (see Santner et al. (2018) for a detailed survey). The downside of space-filling designs for calibration is that they sample the parameter space without regard to the agreement between the simulation output and the observed data. This means space-filling designs often query the simulation model in regions of the parameter space with ill-fitting simulation output. This problem is compounded in multidimensional parameter spaces where high posterior density regions can have very small volume relative to the support of the prior distribution. As an illustration, consider the unimodal density function illustrated in Figure 1 with two parameters in $[-4, 4] \times [-4, 4]$; see Section 8 for details. The left panel of Figure 1 shows 50 samples using LHS where most of the design parameters are placed in the regions with near-zero posterior density. Many of these parameters provide little information for the behavior of the simulation model near the region of interest.

Sequential design (also called active learning) can be used to provide more precise inference by leveraging information learned to select the most informative parameters to sample. In sequential design, previous simulation evaluations are used to build a criterion that assesses the value of evaluating the simulation model at a new set of parameters. The criterion, also called an acquisition function, is then optimized to select a new set of parameters. While previous research

has shown that active learning paired with an emulator results in global high prediction accuracy (Gramacy, 2020, Chapter 6), there is less study on criteria in the calibration setting where global prediction is not needed. There is also a great deal of activity in Bayesian optimization (Frazier, 2018), where the goal is to find an optimizer (e.g., maximum a-posterior estimate), but this often results in designs that are hyper-localized and thus do not service calibration inference.

This article proposes a novel sequential Bayesian experimental design procedure to accurately learn the unnormalized posterior density via an expected integrated variance (EIVAR) criterion. As a bit of foreshadowing, the right panel of Figure 1 illustrates the results from our proposed approach, which generated 50 samples sequentially after 10 initial samples. Our approach places the majority of parameters in the high probability region, which facilitates learning of the posterior distribution. EIVAR measures the uncertainty in the estimate of the unnormalized posterior density, and the sequential strategy selects the parameters to minimize the overall uncertainty on the estimate. The resulting design allocates more points in the high probability regions than in the low probability regions, and hence the unnormalized posterior density function is estimated more accurately with the proposed approach than with space-filling approaches. Moreover, we derive the proposed acquisition strategy considering both the high-dimensional simulation input and output since this is the case for most of the simulation models that require calibration.

Our EIVAR criterion is unique in how it leverages the natural emulation strategy used in modern calibration. The literature on sequential estimation of the posterior (or its analogs) often does not emulate the simulation model itself but instead emulate goodness-of-fit measures directly. For example, Joseph et al. (2015) and Joseph et al. (2019) propose an energy design criterion that aims to mimic a distribution while improving point diversity using an emulator over the distribution. Other examples of this perspective include Kandasamy et al. (2015), who predict the (unnormalized) posterior, and Järvenpää et al. (2019), who instead predict the distance measure for approximate Bayesian computation (Marin et al., 2012). In the related topic of Bayesian quadrature methods (O’Hagan, 1991), which aim to infer on the integral of a density function,

there are methods that also use the emulation of the density itself (e.g., [Fernández et al. \(2020\)](#)). Overall, these approaches seek to create good designs using the emulation of a goodness-of-fit measure itself such as the posterior. Building such emulators can be significantly more difficult compared to directly building emulators of the simulation model. There is untapped potential in directly using an emulator on the simulation model with its associated uncertainty quantification. One example of this perspective for the purposes of point estimation is [Damblin et al. \(2018\)](#), who found success using an emulator for the simulation model to build an expected improvement criterion on the sum of the squared residuals. In contrast to that work, our goal is to infer on the parameters’ posterior density to provide both point estimation and uncertainty quantification for these parameters. Recently, [Koermer et al. \(2023\)](#) propose an integrated mean-squared prediction error criterion for improved field prediction within the Kennedy and O’Hagan calibration framework. In this paper, our focus is on the parameters rather than on the field prediction.

The remainder of the paper is organized as follows. Section [2](#) overviews the main steps of the sequential design procedure. Section [3](#) presents the resulting predictive mean and the variance of the unnormalized posterior density used to build our new acquisition strategy in Section [4](#). We then extend the algorithm to the batched setting in Section [7](#) where several parameters are simultaneously evaluated in parallel to substantially reduce the total computation time. Finally, Section [8](#) illustrates the computational and predictive advantages of the proposed approach using results from several simulation experiments including synthetic problems and a nuclear physics simulation model.

2 Review and Overview

This section overviews the notation, Bayesian calibration methodology, and the proposed sequential algorithm.

2.1 Simulation Structure and Notation

Throughout this article, vectors and matrices are represented in boldface. Our derivations consider a setting where simulation outputs are obtained at all design points once the simulation model is evaluated with parameters $\boldsymbol{\theta}$. To establish some notation, let $(\mathbf{x}_1, \dots, \mathbf{x}_d)$ be the design points where the data $\mathbf{y} = (y(\mathbf{x}_1), \dots, y(\mathbf{x}_d))$ is collected. The simulation model, represented with $\boldsymbol{\eta}$, is a function that takes the parameters $\boldsymbol{\theta} = (\theta_1, \dots, \theta_p)$ in a space $\Theta \subset \mathbb{R}^p$ as inputs and returns a d -dimensional output $\boldsymbol{\eta}(\boldsymbol{\theta}) = (\eta(\mathbf{x}_1, \boldsymbol{\theta}), \dots, \eta(\mathbf{x}_d, \boldsymbol{\theta}))$ at design points $(\mathbf{x}_1, \dots, \mathbf{x}_d)$. We consider the statistical model

$$\mathbf{y} = \boldsymbol{\eta}(\boldsymbol{\theta}) + \boldsymbol{\varepsilon}, \quad (1)$$

where $\boldsymbol{\varepsilon} \sim \text{MVN}(\mathbf{0}, \boldsymbol{\Sigma})$ denotes the residual error and we presume $\boldsymbol{\Sigma}$ is a known matrix. While the precise form of the residual error model is debated (for example, if a discrepancy term is included; see the articles of [Kennedy and O'Hagan \(2001\)](#), [Bayarri et al. \(2007\)](#), [Plumlee \(2017\)](#), and [Gu and Wang \(2018\)](#) for discussions on this in a Bayesian context), for now we consider a known $\boldsymbol{\Sigma}$ accounts for the uncertainty in the difference between the data and the model. We extend our EIVAR criterion for the case of unknown $\boldsymbol{\Sigma}$ with a discrepancy term in [Section 6](#).

In the Bayesian calibration framework, the model parameters are viewed as random variables, and $p(\boldsymbol{\theta}|\mathbf{y})$ indicates the posterior probability density of the parameters $\boldsymbol{\theta}$ given the observed data $\mathbf{y}$. The prior knowledge about the unknown parameters $\boldsymbol{\theta}$ is represented by the prior probability density $p(\boldsymbol{\theta})$. Based on Bayes' rule, the posterior density has the form

$$p(\boldsymbol{\theta}|\mathbf{y}) = \frac{p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{\int_{\Theta} p(\mathbf{y}|\boldsymbol{\theta}')p(\boldsymbol{\theta}')d\boldsymbol{\theta}'} \propto \tilde{p}(\boldsymbol{\theta}|\mathbf{y}) = p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta}), \quad (2)$$

where $p(\mathbf{y}|\boldsymbol{\theta})$ represents the likelihood and $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$ represents the unnormalized posterior. Using the statistical model in [\(1\)](#),

$$p(\mathbf{y}|\boldsymbol{\theta}) = (2\pi)^{-d/2} |\boldsymbol{\Sigma}|^{-1/2} \exp \left(-\frac{1}{2} (\mathbf{y} - \boldsymbol{\eta}(\boldsymbol{\theta}))^\top \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \boldsymbol{\eta}(\boldsymbol{\theta})) \right), \quad (3)$$

where it becomes clear that evaluation of $p(\mathbf{y}|\boldsymbol{\theta})$ requires evaluating the computationally expensive simulation output $\boldsymbol{\eta}(\boldsymbol{\theta})$ at any desired $\boldsymbol{\theta}$.

Our goal is to study estimation of the unnormalized posterior density, $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$, as a function of $\boldsymbol{\theta}$. The unnormalized posterior represents the functional form of the posterior and is used to feed MCMC schemes to get posterior samples. Throughout the remainder of this paper, $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$ is referred to as the posterior for brevity.

2.2 Emulation and Calibration

Emulators are used to accelerate Bayesian inference by predicting the simulation output at novel parameter combinations. A particularly powerful class of emulators has been GPs (Santner et al., 2018; Gramacy, 2020). As noted in the introduction, emulation of a simulation model is often easier compared to emulation of the posterior directly. Typically this is because transferring the simulation model’s output $\eta(\cdot)$ to obtain the likelihood in Equation (3) adds more complexity to the inference problem.

To understand the advantage of emulation of the simulation output, consider the one-dimensional example illustrated in Figure 2. Note that we remove the boldface characters from a parameter θ , a model’s output $\eta(\theta)$, and observed data y in this example since inputs and outputs are both scalars. The left panel of Figure 2 shows an emulator built to predict the simulation output. The emulator can directly predict the likelihood as well and get a variance of that prediction (more details are provided in Section 3) as in the middle panel. In contrast, an emulator designed directly to predict the likelihood $p(y|\theta)$ is shown in the right panel. Figure 2 shows that the prediction coming from the emulator of the simulation output is better at representing the complex posterior surface.

2.3 Overview of the Sequential Bayesian Experimental Design

Our goal is to sequentially acquire simulation outputs from n parameters to accurately estimate the posterior $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$. We first outline an approach where a single parameter and its simulation

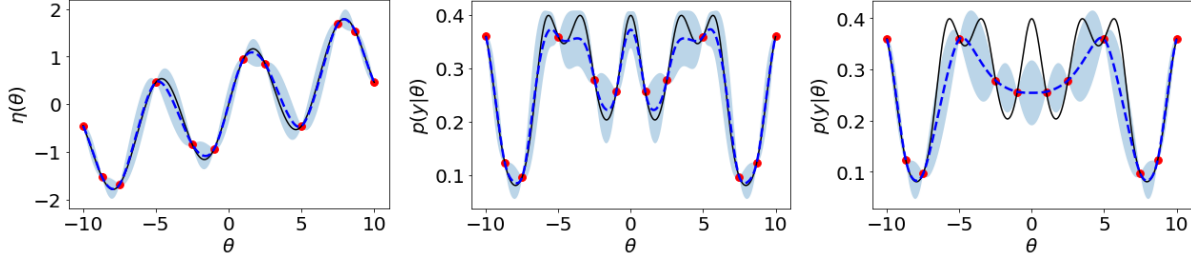


Figure 2: One-dimensional example of calibration with a single parameter θ and a simulation model $\eta(\theta) = \sin(\theta) + 0.1\theta$ and observation $y = 0$ with observation variance one. Red dots indicate the simulation data used to build the emulators and the black line illustrates the simulation model or the likelihood. The blue dashed-line shows the prediction mean and the shaded area illustrates one predictive standard deviation from the mean. Shown are a GP emulator for $\eta(\cdot)$ (left panel); estimated $p(y|\cdot)$ using a GP emulator in left panel suggested in this article (middle panel); and a GP emulator directly for the likelihood $p(y|\cdot)$ (right panel).

output are added to the existing simulation data at each stage indexed by t as described in Algorithm 1. This is extended to batched parameter collection in Section 7.

Algorithm 1: Sequential Bayesian experimental design

- 1 Initialize $\mathcal{D}_1 = \{(\theta_i, \eta(\theta_i)) : i = 1, \dots, n_0\}$
- 2 **for** $t = 1, \dots, n$ **do**
- 3 Fit an emulator $\eta(\theta) | \mathcal{D}_t$
- 4 Generate candidate solutions $\mathcal{L}_t$
- 5 Select $\theta^{\text{new}} \in \arg \min_{\theta^* \in \mathcal{L}_t} \mathcal{A}_t(\theta^*)$
- 6 Evaluate $\eta(\theta^{\text{new}})$
- 7 Update $\mathcal{D}_{t+1} \leftarrow \mathcal{D}_t \cup (\theta^{\text{new}}, \eta(\theta^{\text{new}}))$

Sequential Bayesian experimental design has two main components: a Bayesian statistical model (in the form of a GP-based emulator on the simulation in this paper) and an acquisition function (which is often based on said model). Algorithm 1 presents an overview of the process. It begins by evaluating the simulation model at an initial experimental design of size n_0 . During each stage t , a new parameter is acquired for evaluation with the simulation model. The simulation data set $\mathcal{D}_{t+1} = \{(\theta_i, \eta(\theta_i)) : i = 1, \dots, n_t\}$, where $n_t = n_0 + t$, includes all the simulation data obtained by the end of stage t . The simulation data, $\mathcal{D}_t$, is used to build the

acquisition function $\mathcal{A}_t$. At each stage t , the next parameter is chosen to minimize (approximately) the acquisition function $\mathcal{A}_t$. The acquisition function is minimized over a candidate set of parameters $\mathcal{L}_t$ to avoid difficult numerical optimization of the proposed EIVAR acquisition function (see Section 4). This new parameter and the simulation output is collected into $\mathcal{D}_{t+1}$ and the process repeats. Therefore, at each stage the most informative parameter is selected by leveraging information learned via simulation data $\mathcal{D}_t$.

For the experiments presented in this paper, the sequential procedure is terminated once the simulation outputs are acquired from n parameters. However, one can use alternative stopping criteria. One alternative is to decide on the best n depending on the available computational resources (i.e., available computational hours, number of workers) especially if the simulation model is computationally expensive. If the computational budget allows or the simulation model is less expensive, one can create an hold-out data set (via some sampling techniques such as LHS), and then terminate the approach once a certain level of accuracy is achieved in predicting the posterior of the hold-out sample. Alternatively, one can obtain k -fold cross validation error using the simulation data $\mathcal{D}_t$, and the algorithm can be terminated if the decrease in the CV error stops improving.

2.4 Gaussian Process Model for the High-Dimensional Simulation Output

A GP emulator is constructed at each stage of Algorithm 1. In this article, we show how this emulator can be used to construct an estimator for the posterior (see Section 3) and to build an acquisition function for the selection of a new parameter (see Section 4).

Since the emulation is over a d -dimensional output, where d is the number of observations, we leverage a GP-based emulator that employs the basis vector approach that is now standard practice (e.g., Higdon et al. (2008)). Similarly, for calibration, Huang et al. (2020) train separate independent emulators focused on each design point $\mathbf{x}_i$ instead of building one big emulator for the entire d -dimensional output space. We note that one can replace the basis vector approach

presented in this section with the approach presented by [Huang et al. \(2020\)](#), and our results still hold with minor adjustments to the resulting statements and proofs. Let the $p \times n_t$ parameter matrix $\boldsymbol{\theta}_{1:n_t} = [\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_{n_t}]$ represent the parameters where the simulation model has been evaluated. The standardized outputs are stored in a $d \times n_t$ matrix Ξ_t where the i th column is $(\boldsymbol{\eta}(\boldsymbol{\theta}_i) - \mathbf{h})/\mathbf{s}$ (computed elementwise), and $\mathbf{h}$ and $\mathbf{s}$ are the centering and scaling vectors for the simulation. The $d \times q$ matrix $\mathbf{B} = [\mathbf{b}_1, \dots, \mathbf{b}_q]$ stores the orthonormal basis vectors from principal component analysis of Ξ_t . Note that the dependence on t is dropped from $\mathbf{B}$, $\mathbf{h}$, and $\mathbf{s}$ for simplicity. Then, after transforming Ξ_t with $\mathbf{W}_t = [\mathbf{w}_{t,1}, \dots, \mathbf{w}_{t,q}] = \mathbf{B}^\top \Xi_t$ and letting $\mathbf{G} := \text{diag}(\mathbf{s})$, an independent GP is used for each latent output $w_j(\cdot) = \mathbf{b}_j^\top \mathbf{G}^{-1} (\boldsymbol{\eta}(\cdot) - \mathbf{h})$ for $j = 1, \dots, q$. The GP model results in a prediction with mean $m_{t,j}(\cdot)$ and variance $\varsigma_{t,j}^2(\cdot)$ such that

$$w_j(\boldsymbol{\theta}) | \mathbf{w}_{t,j} \sim \mathcal{N}(m_{t,j}(\boldsymbol{\theta}), \varsigma_{t,j}^2(\boldsymbol{\theta})), \quad j = 1, \dots, q. \quad (4)$$

Following the properties of GPs, the mean $m_{t,j}(\boldsymbol{\theta})$ and variance $\varsigma_{t,j}^2(\boldsymbol{\theta})$ are given by

$$\begin{aligned} m_{t,j}(\boldsymbol{\theta}) &= \mathbf{k}_j(\boldsymbol{\theta}, \boldsymbol{\theta}_{1:n_t}) \mathbf{K}_j(\boldsymbol{\theta}_{1:n_t})^{-1} \mathbf{w}_{t,j} \\ \varsigma_{t,j}^2(\boldsymbol{\theta}) &= \mathbf{k}_j(\boldsymbol{\theta}, \boldsymbol{\theta}) - \mathbf{k}_j(\boldsymbol{\theta}, \boldsymbol{\theta}_{1:n_t}) \mathbf{K}_j(\boldsymbol{\theta}_{1:n_t})^{-1} \mathbf{k}_j(\boldsymbol{\theta}_{1:n_t}, \boldsymbol{\theta}), \end{aligned} \quad (5)$$

where

$$k_j(\boldsymbol{\theta}, \boldsymbol{\theta}') = \tau_j^2 c(\boldsymbol{\theta}, \boldsymbol{\theta}'; \boldsymbol{\zeta}_j), \quad (6)$$

$$\mathbf{K}_j(\boldsymbol{\theta}_{1:n_t}) = \mathbf{k}_j(\boldsymbol{\theta}_{1:n_t}, \boldsymbol{\theta}_{1:n_t}') + v_j \mathbf{I}.$$

Here, τ_j^2 represents a scaling parameter, $c(\boldsymbol{\theta}, \boldsymbol{\theta}'; \boldsymbol{\zeta}_j)$ is a correlation function, $v_j > 0$ is a nugget parameter, and $\boldsymbol{\zeta}_j = (\zeta_{j,1}, \dots, \zeta_{j,p})$ is a lengthscale parameter. For the correlation function c , in this work we use the separable version of the Matérn correlation function with smoothness parameter 1.5 ([Rasmussen and Williams, 2005](#)) such that

$$c(\boldsymbol{\theta}, \boldsymbol{\theta}'; \boldsymbol{\zeta}_j) = \prod_{l=1}^p [(1 + |(\theta_l - \theta'_l) \exp(\zeta_{j,l})|) \exp(-\exp(\zeta_{j,l})|\theta_l - \theta'_l|)]. \quad (7)$$

We define the q -dimensional vector $\mathbf{m}_t(\boldsymbol{\theta}) = (m_{t,1}(\boldsymbol{\theta}), \dots, m_{t,q}(\boldsymbol{\theta}))$ and $q \times q$ diagonal matrix $\mathbf{C}_t(\boldsymbol{\theta})$ with diagonal elements $\varsigma_{t,j}^2(\boldsymbol{\theta})$ for $j = 1, \dots, q$. Then, due to the pseudo-inverse of the

orthonormal $\mathbf{B}$, the predictive distribution on the emulator output is

$$\boldsymbol{\eta}(\boldsymbol{\theta})|\mathcal{D}_t \sim \text{MVN}(\boldsymbol{\mu}_t(\boldsymbol{\theta}), \mathbf{S}_t(\boldsymbol{\theta})), \quad (8)$$

where $\boldsymbol{\mu}_t(\boldsymbol{\theta}) := \mathbf{h} + \mathbf{G}\mathbf{B}\mathbf{m}_t(\boldsymbol{\theta})$ is the emulator predictive mean and $\mathbf{S}_t(\boldsymbol{\theta}) := \mathbf{G}\mathbf{B}\mathbf{C}_t(\boldsymbol{\theta})\mathbf{B}^\top\mathbf{G}$ is the covariance matrix.

3 Posterior Inference

This section describes the expected value and variance of $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$ at each stage t using the emulator described in Section 2.4. Given a parameter $\boldsymbol{\theta}$, the emulator returns a predictive mean $\boldsymbol{\mu}_t(\boldsymbol{\theta})$ and covariance matrix $\mathbf{S}_t(\boldsymbol{\theta})$ based on data $\mathcal{D}_t$. The following result allows us to compute the expectation $\mathbb{E}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t]$ and the variance $\mathbb{V}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t]$ of the posterior given $\mathcal{D}_t$. Throughout the paper, $f_{\mathcal{N}}(\mathbf{a}; \mathbf{b}, \mathbf{C})$ denotes the probability density function of the normal distribution with mean $\mathbf{b}$ and covariance $\mathbf{C}$, evaluated at the value $\mathbf{a}$.

Lemma 3.1. *Assuming that the covariance matrices $\boldsymbol{\Sigma}$ and $\mathbf{S}_t(\boldsymbol{\theta})$ are positive definite, under the model given by Equations (2), (3), and (8),*

$$\mathbb{E}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t] = f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta})) p(\boldsymbol{\theta}), \quad (9)$$

$$\begin{aligned} \mathbb{V}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t] &= \left(\frac{1}{2^{d_{\pi}d/2}|\boldsymbol{\Sigma}|^{1/2}} f_{\mathcal{N}}\left(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta})\right) \right. \\ &\quad \left. - (f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta})))^2 \right) p(\boldsymbol{\theta})^2. \end{aligned} \quad (10)$$

Proof. For the sake of brevity, we drop $\boldsymbol{\theta}$ from $\boldsymbol{\eta}(\boldsymbol{\theta})$, $\boldsymbol{\mu}_t(\boldsymbol{\theta})$, and $\mathbf{S}_t(\boldsymbol{\theta})$. By Equation (2), since

$$\mathbb{E}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t] = \mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})|\mathcal{D}_t] = \mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t]p(\boldsymbol{\theta}),$$

it suffices to show that

$$\mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t] = f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t, \boldsymbol{\Sigma} + \mathbf{S}_t). \quad (11)$$

Using Equations (3) and (8),

$$\mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t] = \int f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\eta}, \boldsymbol{\Sigma}) f_{\mathcal{N}}(\boldsymbol{\eta}; \boldsymbol{\mu}_t, \mathbf{S}_t) d\boldsymbol{\eta},$$

which can be equivalently written as

$$(2\pi)^{-d} |\boldsymbol{\Sigma}|^{-1/2} |\mathbf{S}_t|^{-1/2} \int \exp \left\{ -\frac{1}{2} (\mathbf{y} - \boldsymbol{\eta})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \boldsymbol{\eta}) - \frac{1}{2} (\boldsymbol{\eta} - \boldsymbol{\mu}_t)^\top \mathbf{S}_t^{-1} (\boldsymbol{\eta} - \boldsymbol{\mu}_t) \right\} d\boldsymbol{\eta}. \quad (12)$$

Defining $\mathbf{v} = \boldsymbol{\mu}_t - \boldsymbol{\eta}$ and $\mathbf{z} = \mathbf{y} - \boldsymbol{\mu}_t$, Equation (12) becomes

$$(2\pi)^{-d} |\boldsymbol{\Sigma} \mathbf{S}_t|^{-1/2} \int \exp \left\{ -\frac{1}{2} (\mathbf{v} + \mathbf{z})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{v} + \mathbf{z}) - \frac{1}{2} \mathbf{v}^\top \mathbf{S}_t^{-1} \mathbf{v} \right\} d\mathbf{v}. \quad (13)$$

Writing Equation (13) in matrix notation,

$$\begin{aligned} \mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t] &= (2\pi)^{-d} |\boldsymbol{\Sigma} \mathbf{S}_t|^{-1/2} \int \exp \left\{ -\frac{1}{2} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}^\top \begin{bmatrix} \boldsymbol{\Sigma}^{-1} + \mathbf{S}_t^{-1} & \boldsymbol{\Sigma}^{-1} \\ \boldsymbol{\Sigma}^{-1} & \boldsymbol{\Sigma}^{-1} \end{bmatrix} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix} \right\} d\mathbf{v} \\ &= \int f_{\mathcal{N}} \left(\begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}; \mathbf{0}, \begin{bmatrix} \mathbf{S}_t & -\mathbf{S}_t \\ -\mathbf{S}_t & \boldsymbol{\Sigma} + \mathbf{S}_t \end{bmatrix} \right) d\mathbf{v}, \end{aligned}$$

which corresponds to marginalizing over $\mathbf{v}$. Equation (11) then follows from Gaussian identities

in the appendix of (Rasmussen and Williams, 2005, Equation (A.6)).

We can use similar logic to compute the variance $\mathbb{V}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t]$ such that

$$\begin{aligned} \mathbb{V}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t] &= \mathbb{V}[p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})|\mathcal{D}_t] = \mathbb{V}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t]p(\boldsymbol{\theta})^2 \\ &= (\mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})^2|\mathcal{D}_t] - \mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_t]^2) p(\boldsymbol{\theta})^2. \end{aligned} \quad (14)$$

We first obtain

$$\begin{aligned} \mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})^2|\mathcal{D}_t] &= \int (f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\eta}, \boldsymbol{\Sigma}))^2 f_{\mathcal{N}}(\boldsymbol{\eta}; \boldsymbol{\mu}_t, \mathbf{S}_t) d\boldsymbol{\eta} \\ &= \frac{1}{(2\pi)^{3d/2} |\boldsymbol{\Sigma} \mathbf{S}_t \boldsymbol{\Sigma}|^{1/2}} \int \exp \left\{ -\frac{1}{2} (2(\mathbf{y} - \boldsymbol{\eta})^\top \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \boldsymbol{\eta}) + (\boldsymbol{\eta} - \boldsymbol{\mu}_t)^\top \mathbf{S}_t^{-1} (\boldsymbol{\eta} - \boldsymbol{\mu}_t)) \right\} d\boldsymbol{\eta}. \end{aligned} \quad (15)$$

Using again $\mathbf{v} = \boldsymbol{\mu}_t - \boldsymbol{\eta}$ and $\mathbf{z} = \mathbf{y} - \boldsymbol{\mu}_t$ and writing Equation (15) in matrix notation,

$$\begin{aligned}
&= \frac{1}{(2\pi)^{3d/2} |\boldsymbol{\Sigma}|^{1/2} 2^{d/2} |\frac{1}{2}\boldsymbol{\Sigma}\mathbf{S}_t|^{1/2}} \int \exp \left\{ -\frac{1}{2} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}^\top \begin{bmatrix} 2\boldsymbol{\Sigma}^{-1} + \mathbf{S}_t^{-1} & 2\boldsymbol{\Sigma}^{-1} \\ 2\boldsymbol{\Sigma}^{-1} & 2\boldsymbol{\Sigma}^{-1} \end{bmatrix} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix} \right\} d\mathbf{v} \\
&= \frac{1}{2^d \pi^{d/2} |\boldsymbol{\Sigma}|^{1/2}} \int f_{\mathcal{N}} \left(\begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}; \mathbf{0}, \begin{bmatrix} \mathbf{S}_t & -\mathbf{S}_t \\ -\mathbf{S}_t & \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t \end{bmatrix} \right) d\mathbf{v}.
\end{aligned}$$

The last equality holds since the determinant of the covariance matrix $\begin{bmatrix} \mathbf{S}_t & -\mathbf{S}_t \\ -\mathbf{S}_t & \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t \end{bmatrix}$ is $|\frac{1}{2}\boldsymbol{\Sigma}\mathbf{S}_t|$. The integral corresponds to marginalizing over $\mathbf{v}$, which proves that

$$\mathbb{E}[p(\mathbf{y}|\boldsymbol{\theta})^2|\mathcal{D}_t] = \frac{1}{2^d \pi^{d/2} |\boldsymbol{\Sigma}|^{1/2}} f_{\mathcal{N}} \left(\mathbf{y}; \boldsymbol{\mu}_t, \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t \right). \quad (16)$$

The substitution of Equation (16) into Equation (14) yields Equation (10). $\square$

The expected value in Equation (9) is the estimate of the posterior density and the variance in Equation (10) measures the uncertainty on the posterior. At the extreme, note that as $\mathbf{S}_t(\boldsymbol{\theta})$ gets small and $\boldsymbol{\mu}_t(\boldsymbol{\theta})$ gets close to $\boldsymbol{\eta}(\boldsymbol{\theta})$, meaning the emulator becomes more accurate, the posterior prediction will approach the actual posterior and the variance will approach zero. Our goal is to have this happen with as few samples as possible using smart selection of parameters to evaluate. The prediction and the variance can now be used to construct our EIVAR acquisition function described in the next section.

4 Expected Integrated Variance for Calibration

The acquisition function $\mathcal{A}_t(\boldsymbol{\theta})$ is used to measure the value of evaluating the simulation model $\boldsymbol{\eta}(\boldsymbol{\theta})$ at a set of parameters $\boldsymbol{\theta}$ given data $\mathcal{D}_t$. We suggest using an acquisition function of the aggregated variance of the posterior over the parameter space. The proposed acquisition is built using the mean and variance of posterior prediction from Lemma 3.1. The EIVAR criterion is

inspired by other criteria that aggregate variance under different statistical models (Seo et al., 2000; Binois et al., 2019; Järvenpää et al., 2019; Sauer et al., 2022), but this functional form is novel because we leverage a statistical emulation of the simulation output with a calibration objective. Specifically, the EIVAR criterion is calculated by

$$\mathcal{A}_t(\boldsymbol{\theta}^*) = \int_{\Theta} \mathbb{E}_{\boldsymbol{\eta}^* | \mathcal{D}_t} (\mathbb{V}[p(\mathbf{y} | \boldsymbol{\theta}) | (\boldsymbol{\theta}^*, \boldsymbol{\eta}^*) \cup \mathcal{D}_t]) p(\boldsymbol{\theta})^2 d\boldsymbol{\theta}, \quad (17)$$

where $\boldsymbol{\eta}^* := \boldsymbol{\eta}(\boldsymbol{\theta}^*)$ represents the new simulation output at $\boldsymbol{\theta}^*$. Notice that the expectation is taken over the hypothetical simulation output $\boldsymbol{\eta}^*$, which is a random variable because $\boldsymbol{\eta}^*$ is unknown under data $\mathcal{D}_t$.

To unpack Equation (17), consider Figure 3, which illustrates the three consecutive stages of sequential design using the example presented in Figure 2. For the left panel, using $n_0 = 12$ points and the emulator from Figure 2, the EIVAR criterion is calculated. Since the minima of EIVAR is approximately -6.4 , the simulation model is evaluated at this point. After the statistical emulator of the simulation is refitted with this evaluation included, the EIVAR criterion is computed again at the next stage. The EIVAR values around -6.4 are no longer optimal (those points have the largest EIVAR values) compared to other candidate points, which discourages the reselection of these points. In general, exploration under this criterion happens as previously evaluated parameters and surrounding areas have high EIVAR values.

To use Equation (17), one needs to take the results from Lemma 3.1 and then integrate over all possible values of $\boldsymbol{\eta}^*$ given data $\mathcal{D}_t$. To build closed-form estimates for this, we first establish general results for the GPs needed for the EIVAR criterion. Recall that $m_{t,j}(\boldsymbol{\theta})$ and $\varsigma_{t,j}^2(\boldsymbol{\theta})$ are the emulator mean and variance at stage t for latent output $j = 1, \dots, q$. After observing $(\boldsymbol{\theta}_{n_t+1}, \boldsymbol{\eta}(\boldsymbol{\theta}_{n_t+1})) = (\boldsymbol{\theta}^*, \boldsymbol{\eta}^*)$,

$$\begin{aligned} m_{t+1,j}(\boldsymbol{\theta}) &= [\mathbf{k}_j(\boldsymbol{\theta}, \boldsymbol{\theta}_{1:n_t}), k_j(\boldsymbol{\theta}, \boldsymbol{\theta}^*)] \begin{bmatrix} \mathbf{K}_j(\boldsymbol{\theta}_{1:n_t}) & \mathbf{k}_j(\boldsymbol{\theta}_{1:n_t}, \boldsymbol{\theta}^*) \\ \mathbf{k}_j(\boldsymbol{\theta}^*, \boldsymbol{\theta}_{1:n_t}) & k_j(\boldsymbol{\theta}^*, \boldsymbol{\theta}^*) + v_j \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{w}_{t,j} \\ w_j^* \end{bmatrix} \\ &= m_{t,j}(\boldsymbol{\theta}) + \frac{\text{cov}_{t,j}(\boldsymbol{\theta}, \boldsymbol{\theta}^*)}{\varsigma_{t,j}^2(\boldsymbol{\theta}^*) + v_j} (w_j^* - m_{t,j}(\boldsymbol{\theta}^*)), \end{aligned} \quad (18)$$

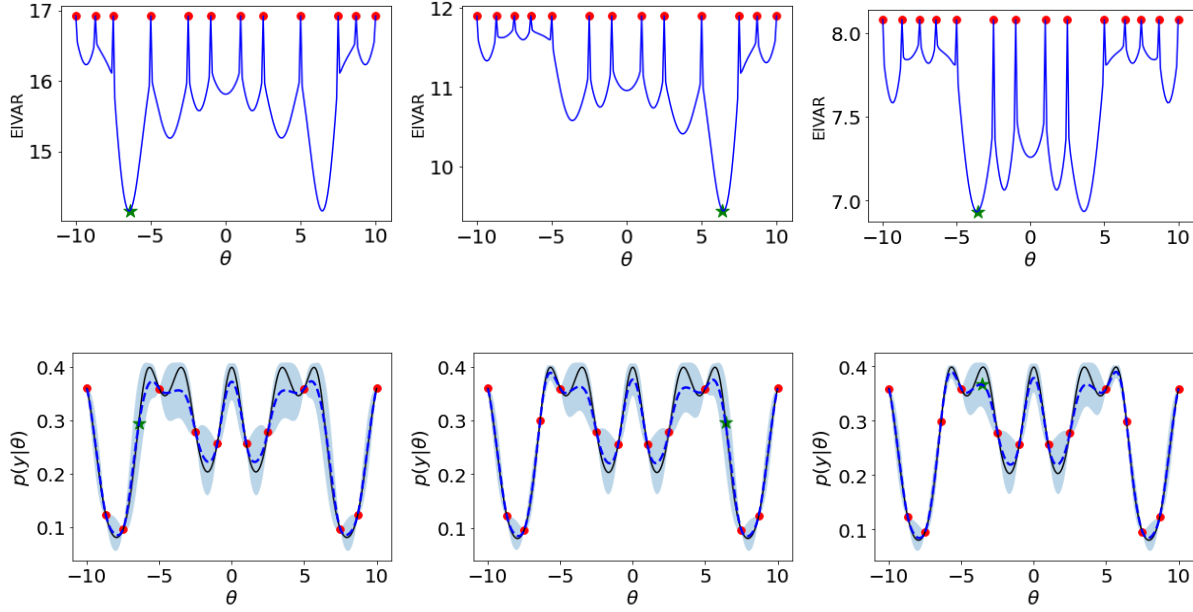


Figure 3: One-dimensional example introduced in Figure 2 to illustrate the EIVAR criterion in Equation (17) (first row). The blue line shows the EIVAR values across the range $\theta \in [-10, 10]$. Red dots illustrate the simulation data used to fit a GP emulator at each stage. The green star is the acquired point, which is added to the simulation data in the subsequent stage. Estimated and true $p(y|\theta)$ are illustrated in the second row.

where $\text{cov}_{t,j}(\theta, \theta^*) := k_j(\theta, \theta^*) - \mathbf{k}_j(\theta, \theta_{1:n_t})\mathbf{K}_j^{-1}(\theta_{1:n_t})\mathbf{k}_j(\theta_{1:n_t}, \theta^*)$, $w_j^* = \mathbf{b}_j^\top \mathbf{G}^{-1}(\boldsymbol{\eta}^* - \mathbf{h})$ is the value of the j th latent output for θ^* , and v_j is a nugget parameter introduced in Section 2.4. Following similar logic for the variance function yields

$$\varsigma_{t+1,j}^2(\theta) = \varsigma_{t,j}^2(\theta) - \frac{\text{cov}_{t,j}(\theta, \theta^*)^2}{\varsigma_{t,j}^2(\theta^*) + v_j}. \quad (19)$$

The variance depends only on the chosen parameter θ^* , not the output $\boldsymbol{\eta}(\theta^*)$. By taking the expected value and the variance of Equation (18),

$$\mathbb{E}_{\boldsymbol{\eta}^*|\mathcal{D}_t}[m_{t+1,j}(\theta)] = m_{t,j}(\theta) \quad \text{and} \quad \mathbb{V}_{\boldsymbol{\eta}^*|\mathcal{D}_t}[m_{t+1,j}(\theta)] = \tau_{t,j}^2(\theta, \theta^*), \quad (20)$$

where $\tau_{t,j}^2(\theta, \theta^*) = \text{cov}_{t,j}(\theta, \theta^*)^2 / (\varsigma_{t,j}^2(\theta^*) + v_j)$. Thus, since $w_j^*|\mathbf{w}_{t,j} \sim \mathcal{N}(m_{t,j}(\theta^*), \varsigma_{t,j}^2(\theta^*) + v_j)$

and by the transformation in Equation (18), we conclude that

$$m_{t+1,j}(\boldsymbol{\theta})|\mathcal{D}_t \sim \mathcal{N}(m_{t,j}(\boldsymbol{\theta}), \tau_{t,j}^2(\boldsymbol{\theta}, \boldsymbol{\theta}^*)). \quad (21)$$

For multidimensional simulation outputs, let $\boldsymbol{\tau}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ be the $q \times q$ diagonal matrix with diagonal elements $\tau_{t,j}^2(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ and let $\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) := \mathbf{B}\boldsymbol{\tau}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)\mathbf{B}^\top$. Equation (21) implies

$$\boldsymbol{\mu}_{t+1}(\boldsymbol{\theta})|\mathcal{D}_t \sim \text{MVN}(\boldsymbol{\mu}_t(\boldsymbol{\theta}), \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)). \quad (22)$$

In addition, Equation (19) implies $\mathbf{C}_{t+1}(\boldsymbol{\theta}) = \mathbf{C}_t(\boldsymbol{\theta}) - \boldsymbol{\tau}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ and $\mathbf{S}_{t+1}(\boldsymbol{\theta}) = \mathbf{S}_t(\boldsymbol{\theta}) - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$.

The next result is useful for computing the EIVAR criterion.

Lemma 4.1. *Assuming that $\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ is positive definite, under the conditions of Lemma 3.1*

$$\begin{aligned} & \mathbb{E}_{\boldsymbol{\eta}^*|\mathcal{D}_t} (\mathbb{V}[p(\mathbf{y}|\boldsymbol{\theta}) | (\boldsymbol{\theta}^*, \boldsymbol{\eta}^*) \cup \mathcal{D}_t]) \\ &= \frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}))}{2^{d\pi d/2}|\boldsymbol{\Sigma}|^{1/2}} - \frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}(\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)))}{2^{d\pi d/2}|\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}}. \end{aligned} \quad (23)$$

Proof. As in the previous results, we drop $\boldsymbol{\theta}$ from the notation and focus on

$$g_{t+1}^2(\boldsymbol{\theta}^*) = \mathbb{E}_{\boldsymbol{\eta}^*|\mathcal{D}_t} (\mathbb{V}[p(\mathbf{y}|\boldsymbol{\theta}) | (\boldsymbol{\theta}^*, \boldsymbol{\eta}^*) \cup \mathcal{D}_t]).$$

Let $\boldsymbol{\mu}_{t+1}^*$ be the mean vector at $\boldsymbol{\theta}$ after seeing data $\mathcal{D}_{t+1}$ that includes $(\boldsymbol{\theta}_{n_{t+1}}, \boldsymbol{\eta}(\boldsymbol{\theta}_{n_{t+1}})) = (\boldsymbol{\theta}^*, \boldsymbol{\eta}^*)$.

Lemma 3.1 shows

$$\mathbb{V}[p(\mathbf{y}|\boldsymbol{\theta})|\mathcal{D}_{t+1}] = \frac{1}{2^{d\pi d/2}|\boldsymbol{\Sigma}|^{1/2}} f_{\mathcal{N}}\left(\mathbf{y}; \boldsymbol{\mu}_{t+1}^*, \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_{t+1}\right) - (f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_{t+1}^*, \boldsymbol{\Sigma} + \mathbf{S}_{t+1}))^2.$$

Given that $\mathbf{S}_{t+1} = \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ does not depend on $\boldsymbol{\eta}^*$, from Equation (22)

$$\begin{aligned} g_{t+1}^2(\boldsymbol{\theta}, \boldsymbol{\theta}^*) &= \int \frac{1}{2^{d\pi d/2}|\boldsymbol{\Sigma}|^{1/2}} f_{\mathcal{N}}\left(\mathbf{y}; \boldsymbol{\mu}_{t+1}^*, \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)\right) f_{\mathcal{N}}(\boldsymbol{\mu}_{t+1}^*; \boldsymbol{\mu}_t, \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)) d\boldsymbol{\mu}_{t+1}^* \\ &\quad - \int (f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_{t+1}^*, \boldsymbol{\Sigma} + \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)))^2 f_{\mathcal{N}}(\boldsymbol{\mu}_{t+1}^*; \boldsymbol{\mu}_t, \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)) d\boldsymbol{\mu}_{t+1}^*. \end{aligned}$$

Letting $\mathbf{L} = \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$, $\mathbf{M} = \boldsymbol{\Sigma} + \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$, and $a_1 = \frac{2^{-d\pi d/2}|\boldsymbol{\Sigma}|^{-1/2}}{(2\pi)^d|\mathbf{L}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}}$, $a_2 =$

$\frac{(2\pi)^{-3d/2}}{|\mathbf{M}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)\mathbf{M}|^{1/2}}$, and assuming $\mathbf{L}$ and $\mathbf{M}$ are invertible, $g_{t+1}^2(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ is equivalently written as

$$\begin{aligned} & a_1 \int \exp\left\{-\frac{1}{2}\left((\mathbf{y} - \boldsymbol{\mu}_{t+1}^*)^\top \mathbf{L}^{-1}(\mathbf{y} - \boldsymbol{\mu}_{t+1}^*) + (\boldsymbol{\mu}_{t+1}^* - \boldsymbol{\mu}_t)^\top \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)^{-1}(\boldsymbol{\mu}_{t+1}^* - \boldsymbol{\mu}_t)\right)\right\} d\boldsymbol{\mu}_{t+1}^* \\ & - a_2 \int \exp\left\{-\frac{1}{2}\left(2(\mathbf{y} - \boldsymbol{\mu}_{t+1}^*)^\top \mathbf{M}^{-1}(\mathbf{y} - \boldsymbol{\mu}_{t+1}^*) + (\boldsymbol{\mu}_{t+1}^* - \boldsymbol{\mu}_t)^\top \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)^{-1}(\boldsymbol{\mu}_{t+1}^* - \boldsymbol{\mu}_t)\right)\right\} d\boldsymbol{\mu}_{t+1}^*. \end{aligned} \quad (24)$$

Defining $\mathbf{v} = \boldsymbol{\mu}_t - \boldsymbol{\mu}_{t+1}^*$ and $\mathbf{z} = \mathbf{y} - \boldsymbol{\mu}_t$, and writing Equation (24) in matrix notation

$$\begin{aligned}
&= \frac{2^{-d} \pi^{-d/2} |\boldsymbol{\Sigma}|^{-1/2}}{(2\pi)^d |\mathbf{L}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}} \int \exp \left\{ -\frac{1}{2} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}^\top \begin{bmatrix} \mathbf{L}^{-1} + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)^{-1} & \mathbf{L}^{-1} \\ & \mathbf{L}^{-1} \end{bmatrix} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix} \right\} d\mathbf{v} \\
&\quad - \frac{(2\pi)^{-3d/2}}{2^{d/2} |\mathbf{M}|^{1/2} |\frac{1}{2}\mathbf{M}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}} \int \exp \left\{ -\frac{1}{2} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}^\top \begin{bmatrix} 2\mathbf{M}^{-1} + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)^{-1} & 2\mathbf{M}^{-1} \\ & 2\mathbf{M}^{-1} \end{bmatrix} \begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix} \right\} d\mathbf{v} \\
&= \frac{1}{2^d \pi^{d/2} |\boldsymbol{\Sigma}|^{1/2}} \int f_{\mathcal{N}} \left(\begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}; \mathbf{0}, \begin{bmatrix} \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) & -\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \\ -\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) & \mathbf{L} + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \end{bmatrix} \right) d\mathbf{v} \\
&\quad - \frac{1}{2^d \pi^{d/2} |\mathbf{M}|^{1/2}} \int f_{\mathcal{N}} \left(\begin{bmatrix} \mathbf{v} \\ \mathbf{z} \end{bmatrix}; \mathbf{0}, \begin{bmatrix} \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) & -\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \\ -\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) & \frac{1}{2}\mathbf{M} + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \end{bmatrix} \right) d\mathbf{v}.
\end{aligned} \tag{25}$$

The last equality holds since the determinants of the covariance matrices are $|\mathbf{L}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|$ and $|\frac{1}{2}\mathbf{M}\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|$, respectively. Marginalizing over $\mathbf{v}$ and using the definition of $\mathbf{L}$ and $\mathbf{M}$ in Equation (25), we find that

$$g_{t+1}^2(\boldsymbol{\theta}, \boldsymbol{\theta}^*) = \frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t, \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t)}{2^d \pi^{d/2} |\boldsymbol{\Sigma}|^{1/2}} - \frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t, \frac{1}{2}(\boldsymbol{\Sigma} + \mathbf{S}_t + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)))}{2^d \pi^{d/2} |\boldsymbol{\Sigma} + \mathbf{S}_t - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}}.$$

□

At the t th stage, the suggestion is then to evaluate the simulation model at the parameter that minimizes Equation (17). In practice, the integral in Equation (17) is approximated with a sum over a uniformly distributed reference set Θ_{ref} within the sample space, and thus the approximate EIVAR is given as

$$\begin{aligned}
&\frac{1}{|\Theta_{\text{ref}}|} \sum_{\boldsymbol{\theta} \in \Theta_{\text{ref}}} p(\boldsymbol{\theta})^2 \left(\frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}))}{2^d \pi^{d/2} |\boldsymbol{\Sigma}|^{1/2}} \right. \\
&\quad \left. - \frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}(\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)))}{2^d \pi^{d/2} |\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}} \right).
\end{aligned} \tag{26}$$

The EIVAR criterion can be used to decide on the next point to be evaluated and it should be the minimizer from the candidate set $\mathcal{L}_t$. Since the first term in Equation (26) does not depend on $\boldsymbol{\theta}^*$,

minimizing Equation (26) is equivalent to maximizing

$$\frac{1}{|\Theta_{\text{ref}}|} \sum_{\boldsymbol{\theta} \in \Theta_{\text{ref}}} p(\boldsymbol{\theta})^2 \left(\frac{f_{\mathcal{N}}(\mathbf{y}; \boldsymbol{\mu}_t(\boldsymbol{\theta}), \frac{1}{2}(\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) + \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)))}{2^d \pi^{d/2} |\boldsymbol{\Sigma} + \mathbf{S}_t(\boldsymbol{\theta}) - \boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)|^{1/2}} \right),$$

which is used to efficiently compute EIVAR in the experiments. In this paper, the sequential design procedure with a closed-form EIVAR criterion is proposed to better estimate the parameters through estimating the posterior density. Additionally, since the posterior is well-estimated with our approach, one can use the emulator at the last stage of our sequential procedure (i.e., the emulator fitted with all acquired parameters) within the Bayesian calibration process to generate samples from the posterior of the parameters.

In our examples and implementation, the matrices $\mathbf{S}_t(\boldsymbol{\theta})$ and $\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ are rank at most q due to projection with $\mathbf{B}$. If $q < d$, this is a violation of the positive definite conditions for Lemmas 3.1 and 4.1 which are adopted to streamline the proofs. However, the final forms in the expressions for lemmas can be computed even without the positive definite condition on $\mathbf{S}_t(\boldsymbol{\theta})$ and $\boldsymbol{\phi}_t(\boldsymbol{\theta}, \boldsymbol{\theta}^*)$ if $\boldsymbol{\Sigma}$ is positive definite. In the case that $\boldsymbol{\Sigma}$ is not positive definite, we suggest adding a small positive definite matrix to $\boldsymbol{\Sigma}$, which is similar to what is suggested by Higdon et al. (2008).

5 Additional Acquisition Functions

While Section 4 describes our core proposed acquisition function, this section introduces alternative acquisition functions used in our empirical study in Section 8. Alternative acquisition functions are inspired from prior literature under different statistical models and represent reasonable alternatives. One can consider them competitors from the literature, but they are morphed to fit our emulator model and calibration objective to obtain a fair benchmarking.

The first alternative acquisition strategy is simply to choose the next point where the uncertainty of the posterior is highest. This is inspired from Seo et al. (2000) who propose the selection of the parameter with the highest emulator variance for fitting a global emulator. In the calibration context, Damblin et al. (2018) use the sum of squared errors of the residuals to measure the

uncertainty on the simulation output prediction. The analogous method called MAXVAR in the empirical study considers the uncertainty on the posterior prediction, and it replaces line 5 in Algorithm 1 by

$$\boldsymbol{\theta}^{\text{new}} \in \arg \max_{\boldsymbol{\theta} \in \mathcal{L}_t} \mathbb{V}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t]. \quad (27)$$

In addition to the posterior variance, one can utilize the expectation $\mathbb{E}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t]$ to build an acquisition function. In Bayesian quadrature scheme, Fernández et al. (2020) propose using expectation with the diversity term $\min_{\tilde{\boldsymbol{\theta}} \in \mathcal{D}_t} \|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|_2$ to encourage exploration by fitting a GP for the posterior. In our proposed framework, the corresponding method abbreviated by MAXEXP uses

$$\boldsymbol{\theta}^{\text{new}} \in \arg \max_{\boldsymbol{\theta} \in \mathcal{L}_t} \mathbb{E}[\tilde{p}(\boldsymbol{\theta}|\mathbf{y})|\mathcal{D}_t] \min_{\tilde{\boldsymbol{\theta}} \in \mathcal{D}_t} \|\boldsymbol{\theta} - \tilde{\boldsymbol{\theta}}\|_2 \quad (28)$$

in place of line 5 in Algorithm 1.

In addition to the competitors introduced above, the two common acquisition functions are the expected improvement (EI) (Jones et al., 1998) and the integrated mean squared error (IMSE) (Sacks et al., 1989) in the active learning literature. EI is used to optimize a function in Bayesian optimization, whereas IMSE is used to build globally accurate emulators. In the calibration setting, one can consider maximizing the posterior $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$ to find the maximum a-posterior estimate of the parameter via EI criterion. In such a case, EI tends to acquire parameters that are hyper-localized in the high posterior region and it provides minimal information about the entire posterior surface. In contrast, IMSE acquires many parameters at near-zero posterior region to better predict the entire simulation model and it does not bring any additional value for estimating the posterior. In Appendix A.1, an illustration is provided to explain the details between EI, IMSE, and EIVAR.

6 EIVAR with Unknown Covariance

The criterion introduced in Section 4 considers a known covariance matrix of errors between the data and the simulation model. This covariance matrix can incorporate observation or sensor

noise, but many modern analyses additionally account for the uncertainty in the model itself in the form of a discrepancy. A general representation is $\mathbf{y} = \boldsymbol{\eta}(\boldsymbol{\theta}) + \mathbf{b} + \boldsymbol{\varepsilon}$, where $\mathbf{b} = (b(\mathbf{x}_1), \dots, b(\mathbf{x}_d))$ represents the discrepancy or the bias term at design points $(\mathbf{x}_1, \dots, \mathbf{x}_d)$. One model, commonly attributed to [Kennedy and O’Hagan \(2001\)](#), utilized a GP prior specification for both the simulation model $\boldsymbol{\eta}(\cdot)$ and the discrepancy function $b(\cdot)$. In our setting, this amounts to a multivariate normal distribution on $\mathbf{b}$. The residual error $\boldsymbol{\varepsilon}$ is often assumed sampled from a multivariate normal distribution with mean zero and covariance $\sigma_\varepsilon^2 \mathbf{I}$. Let $\boldsymbol{\Sigma}^e$ denote the $d \times d$ covariance matrix of the discrepancy term $\mathbf{b}$ plus the noise term $\boldsymbol{\varepsilon}$. Let $\boldsymbol{\theta}^e$ represent all ancillary parameters to define this matrix, which includes GP hyperparameters for the discrepancy term as well as σ_ε^2 . In the case when there is no discrepancy between the simulation model and the observed data, one can consider that σ_ε^2 is $\boldsymbol{\theta}^e$ and set $\boldsymbol{\Sigma}^e = \sigma_\varepsilon^2 \mathbf{I}$. As the ancillary parameters are unknown, and upon giving them the Bayesian treatment, we find joint posterior distribution is

$$p(\boldsymbol{\theta}, \boldsymbol{\theta}^e | \mathbf{y}) \propto p(\mathbf{y} | \boldsymbol{\theta}, \boldsymbol{\theta}^e) p(\boldsymbol{\theta}) p(\boldsymbol{\theta}^e),$$

where the likelihood, $p(\mathbf{y} | \boldsymbol{\theta}, \boldsymbol{\theta}^e)$, matches the right hand side of Equation [\(3\)](#) with $\boldsymbol{\Sigma}$ replaced by $\boldsymbol{\Sigma}^e$.

Following modularization suggestions by [Bayarri et al. \(2007, 2009\)](#), we propose to obtain plausible estimates of $\boldsymbol{\theta}^e$ at each stage t , and then to act as if these were fixed during acquisition. To do this, first, the simulation data is used to build an emulator for the simulation model as described in Section [2.4](#). Then, the unknown ancillary parameters $\boldsymbol{\theta}^e$ are estimated by maximizing the likelihood, i.e.,

$$\hat{\boldsymbol{\theta}}_t^e, \hat{\boldsymbol{\theta}}_t = \arg \max_{\boldsymbol{\theta}^e \in \Theta^e, \boldsymbol{\theta} \in \Theta} |\boldsymbol{\Sigma}^e(\boldsymbol{\theta}^e) + \mathbf{S}_t(\boldsymbol{\theta})|^{-1/2} \exp \left(-\frac{1}{2} (\mathbf{y} - \boldsymbol{\mu}_t(\boldsymbol{\theta}))^\top (\boldsymbol{\Sigma}^e(\boldsymbol{\theta}^e) + \mathbf{S}_t(\boldsymbol{\theta}))^{-1} (\mathbf{y} - \boldsymbol{\mu}_t(\boldsymbol{\theta})) \right), \quad (29)$$

where Θ^e represents the space of ancillary parameters. These estimates are used to approximate the posterior distribution of the calibration parameters, $\tilde{p}(\boldsymbol{\theta} | \mathbf{y})$, by $\tilde{p}(\boldsymbol{\theta} | \mathbf{y}, \hat{\boldsymbol{\theta}}_t^e)$. Then, the EIVAR in Equation [\(26\)](#) can be computed by replacing $\boldsymbol{\Sigma}$ with the covariance estimate $\boldsymbol{\Sigma}^e(\hat{\boldsymbol{\theta}}_t^e)$. Although

our focus is not on the prediction of new, unseen field data, one can use the estimates of the ancillary parameters θ^e and alongside the posterior distribution of θ to predict new, unseen field data in the case of the unknown discrepancy.

7 Batch-Sequential Bayesian Experimental Design

Section 2.1 to this point focuses on a sequential strategy where only one simulation is evaluated at a time. This section extends those ideas to another algorithm that can run multiple expensive simulations in parallel in a batched setting to reduce the wall-clock time needed to estimate the posterior. This batched setting presumes access to a fixed set of workers to perform simulation evaluations in parallel and an additional worker to acquire new parameters at each stage t .

A batch synchronous procedure acquires a batch of size b parameters at each stage, and waits until b simulation evaluations are completed to acquire the next batch of b parameters. Batch synchronous Bayesian optimization methods have been widely used (see, e.g., Ginsbourger et al. (2010) and Azimi et al. (2010)). Figure 4 illustrates an example batch synchronous procedure with the number of workers to perform simulation evaluations equal to the batch size. In this example, once the simulation evaluations of the jobs with indices $\{1, 2, 3, 4\}$ and $\{5, 6, 7, 8\}$ are received from the workers, the acquisitions $\mathcal{A}_1$ and $\mathcal{A}_2$ are initiated on worker with index 0, respectively (see the same color jobs and acquisitions in the left panel of Figure 4). Each of the acquisitions $\mathcal{A}_1$ and $\mathcal{A}_2$ generate four parameters, and their associated jobs $\{5, 6, 7, 8\}$ and $\{9, 10, 11, 12\}$ and simulation evaluations are allocated to the workers (see the same color jobs and acquisitions in the right panel of Figure 4).

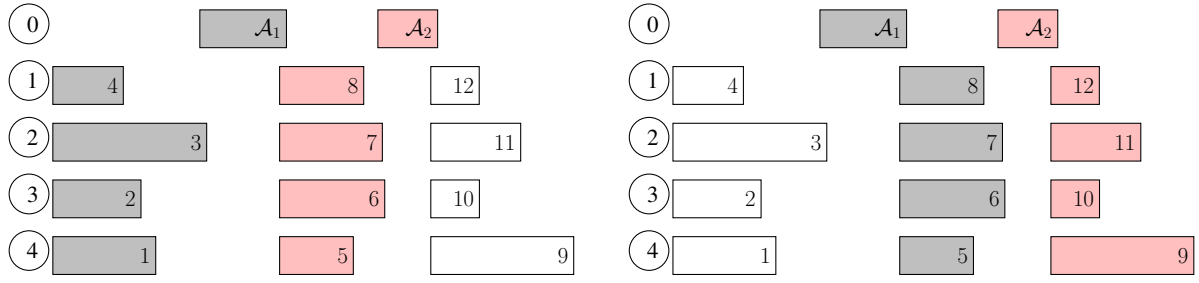


Figure 4: Illustration of batch synchronous procedure with a batch of four parameters. Circles and rectangles represent workers and jobs, respectively, and the numbers inside denote the indices of the corresponding workers and jobs. The left panel illustrates the jobs initiating the same-color acquisition. Once the simulation evaluations of the jobs with indices $\{1, 2, 3, 4\}$ and $\{5, 6, 7, 8\}$ are received, the acquisitions $\mathcal{A}_1$ and $\mathcal{A}_2$ are initiated, respectively. The right panel shows the jobs resulting from the same-color acquisition. The acquisitions $\mathcal{A}_1$ and $\mathcal{A}_2$ generate the jobs with indices $\{5, 6, 7, 8\}$ and $\{9, 10, 11, 12\}$ that are allocated to the workers.

Algorithm 2: Batch-sequential Bayesian experimental design

```

1 Initialize  $\mathcal{D}_1 = \{(\theta_i, \eta(\theta_i)) : i = 1, \dots, n_0\}$ 
2 for  $t = 1, \dots, n/b$  do
3    $\hat{\mathcal{D}}_t \leftarrow \mathcal{D}_t$ ; fit an emulator  $\eta(\theta) | \hat{\mathcal{D}}_t$ 
4   for  $n_a = 1, \dots, b$  do
5     Generate candidate solutions  $\mathcal{L}_t$ 
6     Select  $\theta^{\text{new}} \in \arg \min_{\theta^* \in \mathcal{L}_t} \mathcal{A}_{t, \hat{\mathcal{D}}_t}(\theta^*)$ 
7     Update  $\hat{\mathcal{D}}_t \leftarrow \hat{\mathcal{D}}_t \cup (\theta^{\text{new}}, \mu_t(\theta^{\text{new}}))$ 
8     Update the emulator fitted in line 3 with  $\hat{\mathcal{D}}_t$ 
9    $\mathcal{D}_{t+1} \leftarrow \mathcal{D}_t$ 
10  Update  $\mathcal{D}_{t+1}$  with  $b$  new parameters and their simulation outputs

```

Algorithm 2 describes the batch synchronous procedure where one has access to b workers to run b simulations in parallel, whereas Algorithm 1 allows to run only one simulation at a time. At each stage t , we receive b completed jobs from workers (lines 10) and then acquire b new parameters iteratively (lines 4–8). This process is repeated n/b times (assuming that n/b is an integer value) to acquire n parameters. Once b new parameters are acquired at each stage, the corresponding jobs to obtain their simulation outputs are allocated to the workers. In Algo-

Algorithm 2, a total of n/b emulators are fitted from scratch. Additionally, as the simulation output linked to the parameters is not yet obtained while constructing a batch, the emulator is updated using $\mu_t(\theta^{\text{new}})$ as emulation input (see lines 7-8), which is fitted at the start of each stage (see line 3). Imputing the current prediction as emulation input when building a batch is termed the kriging believer strategy in Ginsbourger et al. (2010) for Bayesian optimization. We note that the kriging believer strategy is used to illustrate the proposed batch-sequential procedure, and one can replace the kriging believer strategy with an alternative approach such as constant liar strategy. In Algorithm 2, the hyperparameters for the emulator are not altered on line 8; only the covariance matrix is updated with data $\hat{\mathcal{D}}_t$ and the prediction mean remains unchanged as the emulator is updated using $\mu_t(\theta^{\text{new}})$. The hyperparameters of the emulator are relearned at the beginning of each stage when fitting an emulator from scratch (line 3) after b simulation outputs are received from workers during the previous stage. In Algorithm 2, in addition to stage index t , $\hat{\mathcal{D}}_t$ is a subscript of the acquisition function $\mathcal{A}_{t, \hat{\mathcal{D}}_t}$ on line 6 to indicate that each time a new parameter is selected the acquisition function is minimized using the emulator based on data $\hat{\mathcal{D}}_t$.

8 Results

This section compares the performance of sequential and batch-sequential approaches with various acquisition functions described in Sections 4 and 5. Section 8.1 considers the one-at-a-time sequential procedure on a variety of test functions. Section 8.3 focuses on a computational physics problem and examines both sequential and batch-sequential approaches with batch sizes of 2, 4, 8, 16, and 32. We provide an open source implementation of different selection strategies and their parallel implementations in the PUQ Python package (Sürer et al., 2022b). All the test functions and the physics examples are provided in the `examples` directory within the PUQ Python package.

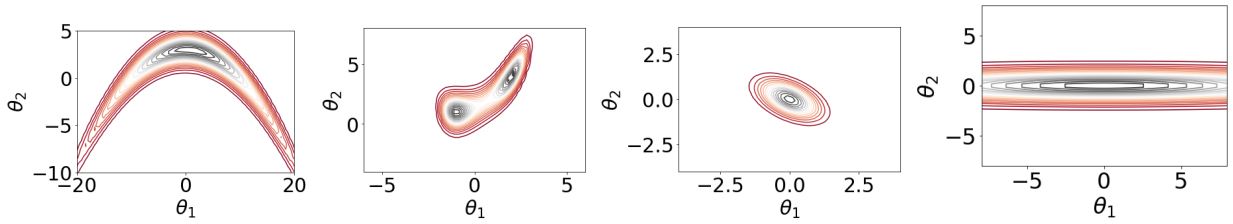


Figure 5: Posterior probability density illustration for four synthetic test functions. Labels of the test functions from left to right: banana, bimodal, unimodal, and unidentifiable.

8.1 Simulations from Synthetic Probability Densities

We investigate the performance of the sequential Bayesian experimental design with the proposed EIVAR acquisition function. As a baseline, we sample parameters from a uniform prior, and this approach is denoted as RND.

To achieve a baseline outside of our framework, the results are also compared with the minimum energy design (MINED) (Joseph et al., 2015, 2019). Instead of minimizing the acquisition function over a candidate set, MINED uses a simulated annealing algorithm to minimize its acquisition function, and the total number of simulation evaluations to acquire n parameters is larger than n , which makes it difficult to directly compare resulting designs; thus, the results are included in Appendix A.3. Results indicate that EIVAR performs better than MINED, and MINED has larger variability than EIVAR. We also note that MINED requires evaluating the simulation model at more than 1500 parameters to acquire 200 of them, whereas the proposed approach only requires $n + n_0$ evaluations (which is 210 in these experiments). Our article presumes that the simulation model is expensive to evaluate, meaning MINED is a potentially wasteful strategy relative to our proposed approach.

More reasonable competitors for a comparison use the sequential framework from Section 2.3 with the MAXVAR and MAXEXP acquisition functions described in Section 5. When acquiring a set of parameters via MAXEXP criterion, since the diversity term depends on the distance between two sets of parameters, it is computed after the parameters are transformed into a unit

hypercube. We examine these different acquisition procedures using four common synthetic test functions (see examples in [Joseph et al. \(2019\)](#); [Järvenpää et al. \(2019\)](#)) with different shapes of density illustrated in Figure [5](#). In those examples, the parameter space is two-dimensional ($p = 2$) and the simulation model returns a two-dimensional output ($d = 2$) for the banana, bimodal, and unidentifiable functions, and a one-dimensional output ($d = 1$) for the unimodal function. We also examine performance with higher dimensional input parameters and simulation outputs, namely three-dimensional ($p = d = 3$), six-dimensional ($p = d = 6$), and ten-dimensional ($p = d = 10$) inputs and outputs. The details of each simulation model and resulting likelihood are provided in Appendix [A.2](#).

A uniform prior is assumed for all cases and the initial sample used to build the first emulator is taken from this prior with ranges given in Appendix [A.2](#). It is important to note that the prior probability density $p(\boldsymbol{\theta})$ plays a critical role in EIVAR as in Equation [\(26\)](#) similar to the other Bayesian calibration procedures, and we provide an example in Appendix [A.4](#) to test the sensitivity of EIVAR for various strengths of the prior. We set the initial sample size $n_0 = 10$ for the two-dimensional functions in Figure [5](#) and set $n_0 = 20$, $n_0 = 40$, and $n_0 = 60$ for three-, six-, and ten-dimensional examples, respectively. The size of initial sample is chosen to allow the emulator to learn the response surface well enough during the initial stages and at the same time to allow improvements in the predictions through acquisitions. At each stage t , the candidate list $\mathcal{L}_t$ is constructed as a random sample from the uniform prior with a size of 100 for all the synthetic functions. For the banana, bimodal, unidentifiable, and unimodal functions, 200 total parameters are acquired. For higher dimensional functions, 400 total parameters are acquired. To compare the performance of different acquisition functions, at each stage we report the mean absolute difference (MAD) between the estimated posterior and the ground truth posterior. In the experiments, we generate a set of reference parameters Θ_{ref} and obtain the unnormalized posterior density for each parameter in this set. We then compute MAD at each stage t as $\text{MAD}_t(\tilde{p}, \hat{p}) = \frac{1}{|\Theta_{\text{ref}}|} \sum_{\boldsymbol{\theta} \in \Theta_{\text{ref}}} |\tilde{p}(\boldsymbol{\theta}|\mathbf{y}) - \hat{p}_t(\boldsymbol{\theta}|\mathbf{y})|$, where $\tilde{p}(\boldsymbol{\theta}|\mathbf{y})$ is the unnormalized poste-

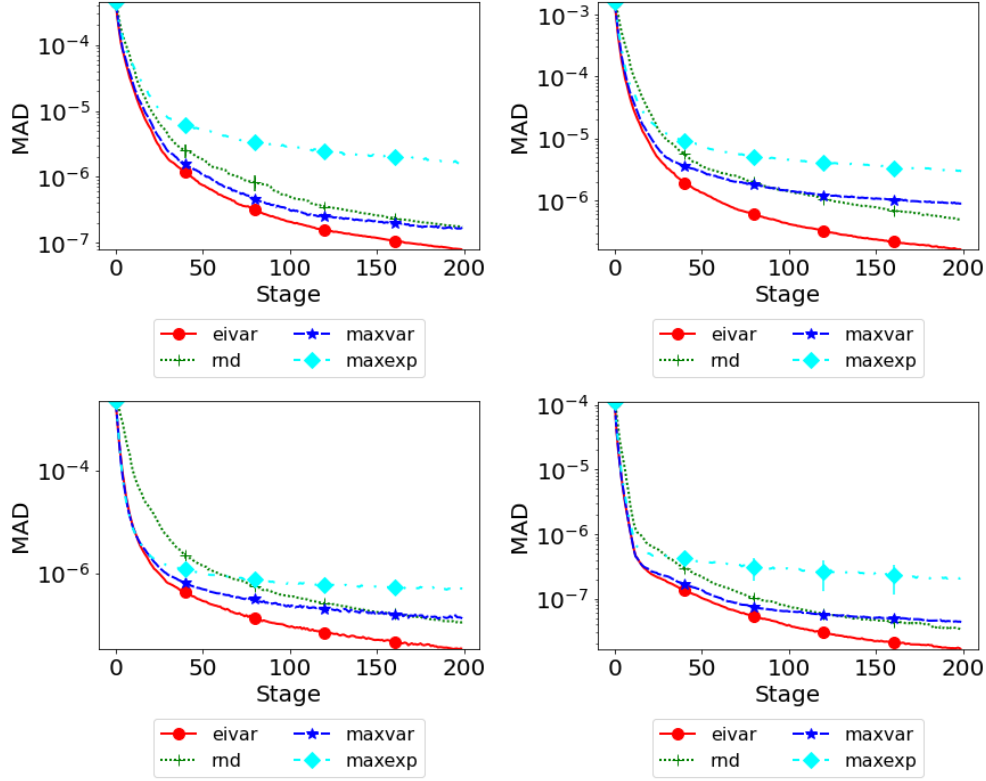


Figure 6: Comparison of different acquisition functions using the test functions illustrated in Figure 5 containing banana (upper-left), bimodal (upper-right), unimodal (lower-left), and unidentifiable (lower-right). Error bars denote 95% confidence intervals for MAD across 50 replicates. Algorithm 1 is used to acquire 200 parameters.

rior density and $\hat{p}_t(\theta|\mathbf{y})$ is the estimated unnormalized posterior at stage t , which is obtained via Equation (9). For the examples in Figure 5, Θ_{ref} is generated in a two-dimensional grid of 50^2 points. For higher dimensional examples, since the size of a grid covering the parameter space becomes very large, 10^4 points are generated using LHS. The set Θ_{ref} is also used in Equation (26) to approximate the integral.

Figures 6-7 summarize the accuracy results and indicate that EIVAR has the lowest MAD between the estimated posterior and the ground truth posterior for all seven test functions. From our observations of acquired parameters in low dimensional functions (see Figure 15 in Appendix A.2), MAXVAR tends to ignore the regions where the cumulative uncertainty is larger

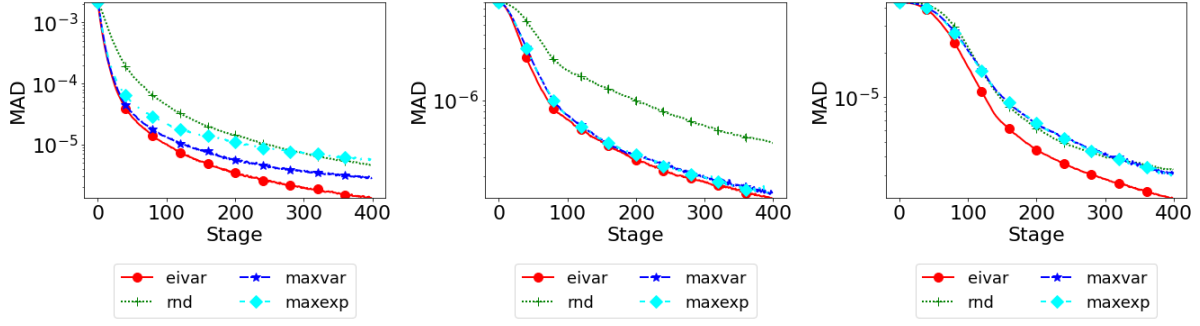


Figure 7: Comparison of four acquisition functions using three- (left), six- (center), and ten-dimensional (right) test functions. Error bars denote 95% confidence intervals for MAD across 50 replicates. Algorithm [1](#) is used to acquire 400 parameters.

and instead picks points at the boundary where the variance is larger. On the other hand, EIVAR recognizes that the acquired parameters impact the posterior estimation globally, and it acquires points both at the interior and boundary of the high-posterior region. The RND strategy performs better than MAXEXP for the functions in Figure [6](#), especially during later stages of the sequential procedure. However, for higher-dimensional spaces (see left and center panels) it takes more time to accurately estimate the posterior for RND strategy as compared to other methods. The prior range of parameters is kept the same in three- and six-dimensional examples; Figure [6](#) clearly shows that when the high posterior region becomes smaller relative to the support of the prior distribution, it becomes more challenging for RND to target the high posterior region. One disadvantage of EIVAR is the computational time needed to acquire a single parameter on higher-dimensional spaces since EIVAR requires an integral evaluation. This limitation can be mitigated as p grows, and in order to overcome the curse of dimensionality, one can consider sparse grids or other clever sampling or quadrature schemes for estimating higher dimensional integrals. Naturally, this computational expense is less of an issue when it is expensive to obtain a simulation output, as is the case described in the next subsection.

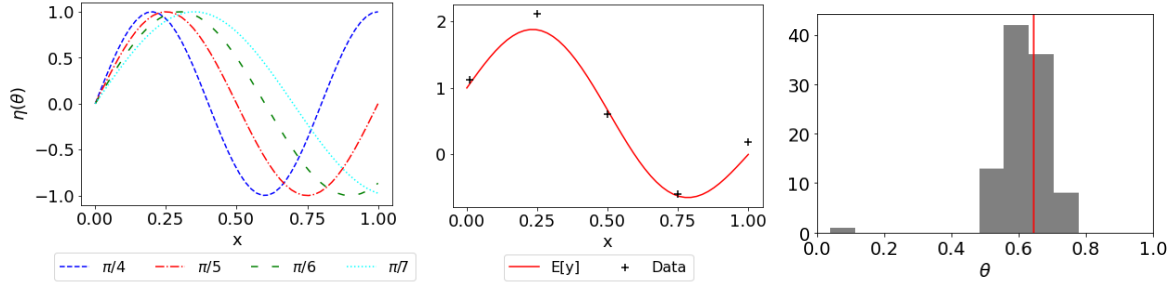


Figure 8: An illustrative example in the case of discrepancy term with $\mathbf{y} = \boldsymbol{\eta}(\theta = \frac{\pi}{5}) + \mathbf{b} + \boldsymbol{\varepsilon}$, $\boldsymbol{\eta}(\theta) = (\eta(x_1, \theta), \dots, \eta(x_5, \theta))$, and $\mathbf{b} = (b(x_1), \dots, b(x_5))$ where $\eta(x, \theta) = \sin(10x\theta)$ and $b(x) = 1 - \frac{1}{3}x - \frac{2}{3}x^2$ at design points $(x_1, x_2, x_3, x_4, x_5) = (0, 0.25, 0.50, 0.75, 1)$. In the left panel, lines correspond to the simulation outputs at $\theta \in \{\frac{\pi}{4}, \frac{\pi}{5}, \frac{\pi}{6}, \frac{\pi}{7}\}$. In the middle panel, the red line corresponds to $\mathbb{E}(\mathbf{y}) = \boldsymbol{\eta}(\theta = \frac{\pi}{5}) + \mathbf{b}$, and the black markers show the observed data with $\sigma_\varepsilon^2 = 0.2^2$. The right panel illustrates the distribution of 100 acquired parameters along with the red vertical line at the least-squares fit parameter value minimizing the sum of the squared errors between $\mathbf{y}$ and $\mathbb{E}(\mathbf{y})$.

8.2 Simulations in the Case of Discrepancy

The results in Section 8.1 assume known Σ with no discrepancy between the simulation model and the observed data. In this section, we observe the acquired parameters via the EIVAR criterion when the covariance matrix of errors between the observed data and the simulation model is unknown. As an illustration, consider the example in Figure 8 with a single parameter in $[0, 1]$; see also the illustrative example presented in Koermer et al. (2023) for details. The left panel of Figure 8 shows simulation outputs at four different parameter values. In the middle panel, the mean of the observed data $\mathbb{E}(\mathbf{y}) = \boldsymbol{\eta}(\theta) + \mathbf{b}$ and the observed data at five design points on an equally spaced grid are shown. We set the initial sample size $n_0 = 10$, and acquire 100 parameters using the procedure described in Section 6 to account for the observation noise and the uncertainty in the model in the form of discrepancy. We adopt a covariance form Σ^e defined by $\Sigma_{i,j}^e = \sigma_\varepsilon^2 \delta_{i,j} + \sigma_b^2 \exp(-\lambda|x_i - x_j|)$ for $i, j = 1, \dots, d$, where $\delta_{i,j}$ is a Kronecker delta func-

Label	Input	Range	Unit
θ_1	V : depth (real volume)	[40, 60]	MeV
θ_2	r : radii	[0.7, 1.2]	fm
θ_3	W_s : depth (imaginary surface)	[2.5, 4.5]	MeV

Table 1: Parameters and their ranges for the $^{48}\text{Ca}(n,n)^{48}\text{Ca}$ reaction model.

tion, and the likelihood in Equation (29) is optimized to find the plausible estimates of ancillary parameters $\theta^e = [\sigma_\epsilon^2, \sigma_b^2, \lambda]$ at each stage. The least-squares fit parameter value is obtained by minimizing the sum of the squared errors between $\mathbf{y}$ and $\mathbb{E}(\mathbf{y})$. In this example, the EIVAR criterion with unknown covariance matrix is able to target the high posterior region by acquiring almost all parameters concentrated around the least-squares fit parameter value.

8.3 Application to a Reaction Model

We now illustrate our sequential strategy with a nuclear physics model to predict differential cross section as a function of angle. Different parameterizations of the optical potential can result in different angular cross sections, and the reaction code `FRESCOX` (Thompson, 1988) takes the optical model parameters as input and generates the corresponding cross sections across angles from 0° to 180° (Lovell et al., 2017). This case study uses the elastic scattering data for the $^{48}\text{Ca}(n,n)^{48}\text{Ca}$ reaction presented in Lovell (2018) to obtain the best-fit parameterization for optical parameters (see Appendix B in Lovell (2018) for details). To simplify the example, some parameters from Lovell (2018) are fixed, since these parameters reproduce the data $\mathbf{y}$ well, while the remaining three parameters are acquired with the proposed approach. Table 1 summarizes the parameters used in this case study and their prior ranges.

The elastic scattering data $\mathbf{y}$ is available only at 15 different degree values, and the simulation outputs (i.e., cross sections) are obtained at the corresponding degree values. We assume a known

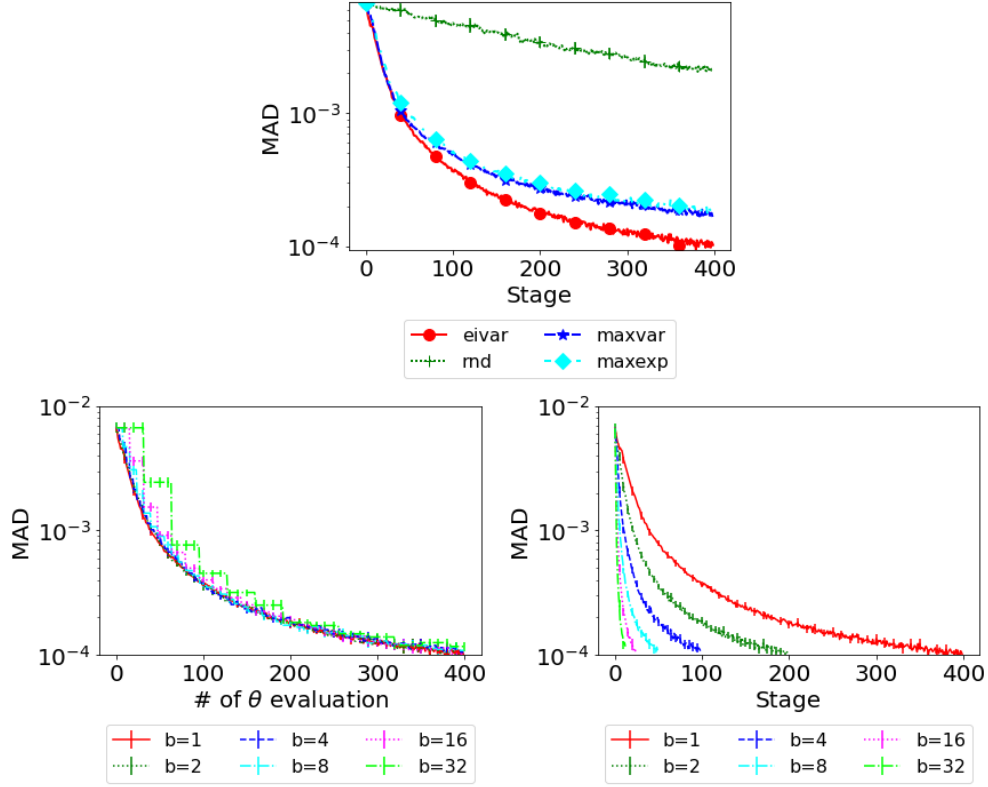


Figure 9: Results of physics reaction model. Error bars denote 95% confidence intervals for MAD across 50 replicates. Panel in the first row compares different acquisition functions; panels in the second row compare different batch sizes for the proposed EIVAR criterion with two perspectives.

diagonal covariance matrix such that $\Sigma_{i,i} = 0.1$ for $i = 1, \dots, 15$. We take 32 initial samples and obtain 400 parameters with each acquisition function. Each of the acquired parameters is chosen from the candidate list $\mathcal{L}_t$ of size 100, which is generated by sampling from the uniform prior at each stage t . First, we compare different acquisition strategies with Algorithm 1 (e.g., $b = 1$). Then, we focus on batch strategies with EIVAR and obtain the results with $b \in \{2, 4, 8, 16, 32\}$ via a batch synchronous procedure presented in Algorithm 2 (the number of workers equals to b to run simulations in parallel). Figure 9 summarizes the results with both sequential and batched approaches.

As illustrated in the first row of Figure 9, with the sequential procedure the posterior is well-

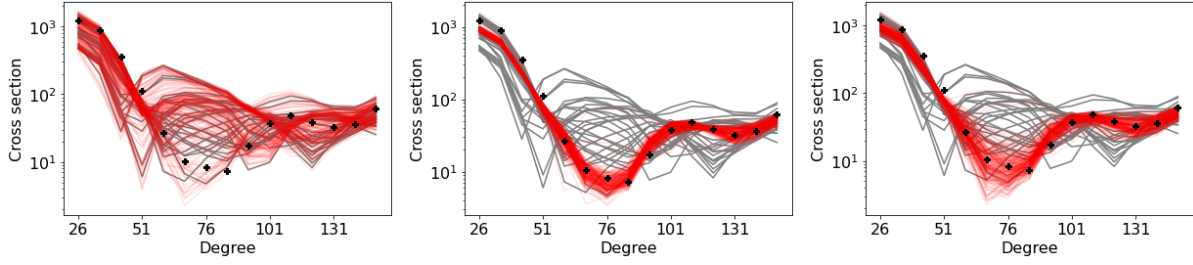


Figure 10: Cross sections (light-red lines) obtained with the sequential strategy via acquisition functions RND (left), MAXVAR (center), and EIVAR (right). Light-gray lines illustrate the cross sections for the initial sample. Black plus markers represent the observed data y .

estimated with EIVAR as compared to the other methods, whereas the prediction accuracy with RND does not even improve significantly over time. Figure 10 illustrates the cross sections of the acquired parameters for RND, MAXVAR, and EIVAR. Since RND does not consider the system data (e.g., black plus markers in Figure 10), most of the cross sections are placed onto the outside of the region of interest. While MAXVAR considers densely the regions around the data, EIVAR focuses on a wider region where the posterior is smaller but not exactly zero. Therefore, EIVAR covers the space well enough by not considering the zero-posterior region but at the same time distributing the acquired parameters around both high and low posterior regions to better learn the posterior. Finally, we observe the distribution of acquired parameters with EIVAR in Figure 11 (left panel) along with the values of the best-fit parameters. Since the best-fit parameters are included into region specified with the acquired parameters, EIVAR is also able to identify the highest posterior region.

The results of the batch synchronous approach with batch sizes $b \in \{2, 4, 8, 16, 32\}$ are provided in the second row of Figure 9. While the left panel summarizes MAD after obtaining each simulation output (x -axis represents each of 400 simulation evaluations), the right panel shows MAD after each stage (x -axis represents each stage). One can consider the figure in the right panel as a proxy to the total wall-clock time if the simulation was the dominant computational cost. Since 400 parameters are acquired with each of the batch sizes, sequential approaches

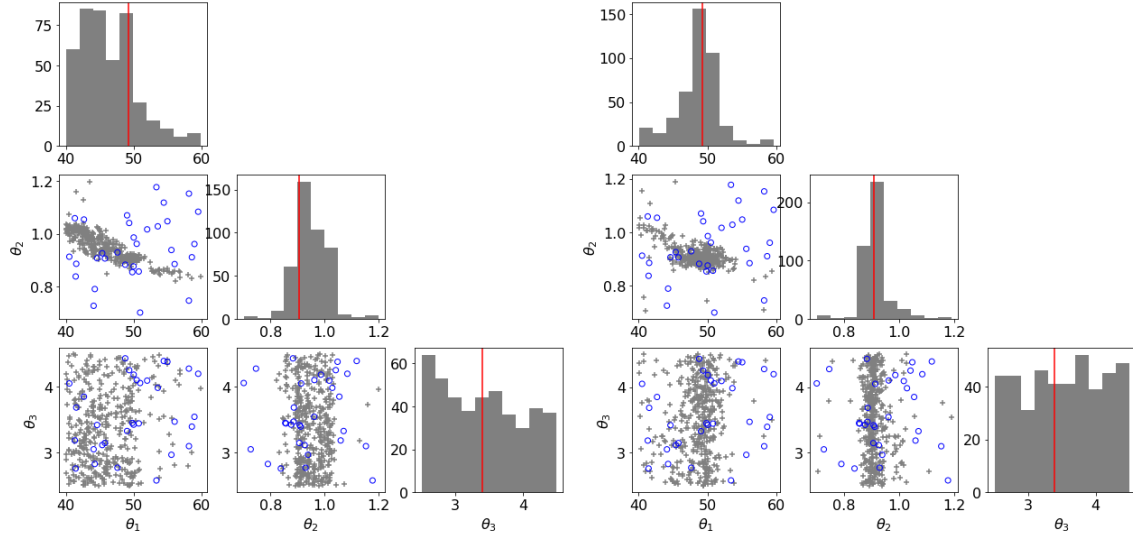


Figure 11: Pairwise plot for acquired parameters with EIVAR (gray plus markers) with known (left panel) and unknown (right panel) covariance matrix. Blue circles are the initial sample obtained with LHS. The red line corresponds to the parameter values obtained with the best-fit parameterization.

with larger batch sizes terminate in fewer batches. Although the final accuracy of different batch sizes does not differ significantly between each other, there is a significant difference between the sequential strategy with $b = 1$ and the batched strategy with $b = 32$ up until 200 simulation evaluations are completed. This indicates that the quality of emulator with the kriging-believer strategy starts degrading with an increasing batch size, especially during the early periods of data collection when the acquired parameters and their corresponding simulation outputs are not enough to explore the space. Thus, a user should decide on the best batch size considering the trade-off between the computational time and accuracy.

All the results so far use the known diagonal covariance matrix Σ with diagonal elements of 0.1 to acquire parameters. Although the assumption comes from the experts' best understanding of the experiment, one can also use the procedure presented in Section 6 to learn aspects of the covariance. To illustrate this, we use the reaction code FRESKOX with the same settings introduced above except that we assume an unknown covariance Σ^e such that $\Sigma_{i,j}^e =$

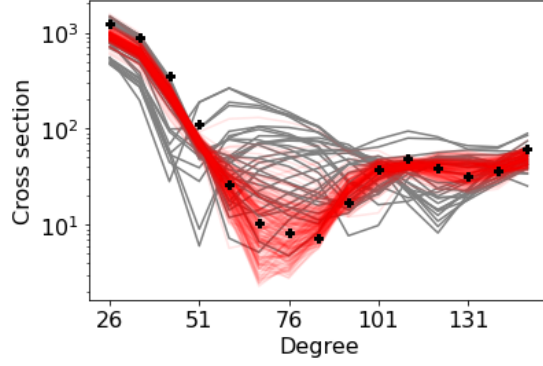


Figure 12: Cross sections (light-red lines) obtained with the sequential strategy via EIVAR with the discrepancy term. Light-gray lines illustrate the cross sections for the initial sample. Black plus markers represent the observed data y .

$\sigma_\varepsilon^2 \delta_{i,j} + \sigma_b^2 \exp(-\lambda |x_i - x_j|)$, for $i, j = 1, \dots, d$, to account for both the observation noise and the model uncertainty in the form of discrepancy where $\theta^e = (\sigma_\varepsilon^2, \sigma_b^2, \lambda)$. The observed data y is collected at 15 design points (i.e., at angles) $(x_1, \dots, x_{15}) = (26, 31, 41, 51, 61, 71, 76, 81, 91, 101, 111, 121, 131, 141, 151)$. The right panel of Figure 11 shows the distribution of acquired parameters with an unknown covariance matrix and demonstrates that EIVAR is able to identify the high posterior region in this situation as well. Figure 12 illustrates the cross sections of the acquired parameters; similar to the known covariance matrix situation, EIVAR focuses on a wider region as compared to the other approaches tested. Additionally, we have tested EIVAR with an unknown covariance matrix Σ^e to incorporate only the observation noise such that $\Sigma_{i,j}^e = \sigma_\varepsilon^2 \delta_{i,j}$, for $i, j = 1, \dots, d$. For this case, the cross sections and the pairwise plot of acquired parameters are similar to the ones obtained with the known covariance illustrated in Figure 10 (right panel) and Figure 11 (left panel), respectively. Although the resulting cross sections cover a similar region in all three of these cases, the parameters (especially θ_1 and θ_2) are tightly acquired around the best-fit parameters by the unknown covariance with the discrepancy term assumption (as in the right panel of Figure 11) indicating that the discrepancy term helps better constrain the range of acquired parameters.

9 Conclusion

This article proposes a novel sequential Bayesian experimental design procedure to estimate the posterior that is known up to a proportionality constant. Our approach selects the most informative parameters using previous simulation evaluations and places a majority of them in the high-posterior region, which facilitates the learning of the posterior region. The proposed approach can be advantageous especially when dealing with computationally expensive simulation models since it requires fewer number of simulation evaluations than space-filling approaches to learn the posterior at the same accuracy level. We extend a single acquisition case to batch acquisitions as well. Both sequential and batched settings require only the simulation model and the prior distribution for the parameters as input. Thus, practitioners can easily use the proposed approaches for a wide variety of calibration applications.

In this paper, a novel EIVAR criterion is derived for calibrating computationally expensive simulation models, and its closed-form expression makes it appealing for a variety of future work. Inspiring from the aggregate variance criterion when building an emulator with globally high prediction accuracy, one area is to extend the criterion for stochastic simulation models (Binois et al., 2019). In parallel to this, our proposed EIVAR criterion can be further developed to decide whether a new parameter is acquired for replication or exploration of the posterior surface when calibrating stochastic simulation models. Another area is to replace the aforementioned emulator with enhanced emulators such as deep GPs (Sauer et al., 2022), and the parameters can be selected via EIVAR sequentially for the purpose of calibration when dealing with complex response surfaces. Moreover, in the presence of non-identifiability, multiple parameter sets may result in high posterior density, and this issue may be diminished or controlled with the right active learning criteria. In such a setting, one can utilize both the proposed EIVAR and the new active learning criteria in an hybrid way to better estimate the posterior. In the case where a simulation model returns one-dimensional output $\eta(\mathbf{x}, \boldsymbol{\theta})$ at a design point $\mathbf{x}$ and calibration parameter

θ , one can consider acquiring a design point and a calibration parameter simultaneously. For this new setup, one can consider replacing the proposed emulation strategy with a global GP emulator. Since our results do not depend on the way the emulator is built, our derivations can still be useful for this setting with minor adjustments to the resulting statements and proofs. Acquiring design points and the calibration parameters simultaneously would be especially of interest when a physical experimental design is the goal and the sequential collection of data $\mathbf{y}$ is possible. Additionally, in this paper, our results assume that the residual error is sampled from a multivariate normal distribution; another line of future development would be to derive the EIVAR criterion for different distributions of the residual error.

Batch strategies allow computationally expensive simulation models to be run in parallel. In cases where the run times of each simulation evaluation are roughly equal, the batch synchronous sequential approach described here is effective to reduce the wall-clock time. However, if the run times in a batch vary, the batch synchronous approach may lead to inefficient utilization of the computational resources since workers are left idle while waiting for the slowest job in the batch to finish. In order to improve the utilization of parallel computing resources, we can run function evaluations asynchronously: as soon as b workers complete their jobs, b new parameters are chosen and sent to workers choose new tasks for them (with $b < \text{the number of workers}$). A batch asynchronous procedure has received little attention as compared to synchronous one. A more-detailed performance analysis for both the synchronous and asynchronous procedures and a guideline for the selection of the best batch size would be of great interest to practitioners as a future study.

Acknowledgement

The authors are grateful to Amy Lovell and Filomena M. Nunes for taking time to provide us physics reaction model.

A Appendix

A.1 Comparison of EIVAR with Common Acquisition Functions

The expected improvement (EI) and the integrated mean squared error (IMSE) are the two common acquisition functions in the literature. While the EI criterion acquires the set of parameters offering the greatest expected improvement over the current best parameter to maximize a function (e.g., the posterior), IMSE selects the set of parameters that minimizes the overall uncertainty in the estimate of the simulation model at each stage to build a globally accurate emulator. In this section, we demonstrate the differences between EIVAR, IMSE, and EI acquisition functions with a one-dimensional illustrative example.

To implement the EI criterion, a GP emulator of the posterior is obtained at each stage t (see Section 2.2 for an example), and then EI uses

$$\boldsymbol{\theta}^{\text{new}} \in \arg \max_{\boldsymbol{\theta} \in \mathcal{L}_t} \left((\check{m}_t(\boldsymbol{\theta}) - p_{\max}) \Phi \left(\frac{\check{m}_t(\boldsymbol{\theta}) - p_{\max}}{\check{\sigma}_t(\boldsymbol{\theta})} \right) + \check{\sigma}_t(\boldsymbol{\theta}) f_{\mathcal{N}} \left(\frac{\check{m}_t(\boldsymbol{\theta}) - p_{\max}}{\check{\sigma}_t(\boldsymbol{\theta})}; 0, 1 \right) \right), \quad (30)$$

where $p_{\max} = \max_{i=1, \dots, n_t} \tilde{p}(\boldsymbol{\theta}_i | \mathbf{y})$, $\Phi(\cdot)$ is the cumulative density function of the standard normal distribution, and $\check{m}_t(\boldsymbol{\theta})$ and $\check{\sigma}_t^2(\boldsymbol{\theta})$ are the prediction mean and variance resulting from a GP model for the posterior. To compute the IMSE at any parameter, an emulator of the simulation model is obtained at each stage as described in the paper and the parameter that minimizes the total uncertainty of the emulator is acquired via

$$\boldsymbol{\theta}^{\text{new}} \in \arg \min_{\boldsymbol{\theta}^* \in \mathcal{L}_t} \sum_{\boldsymbol{\theta} \in \Theta_{\text{ref}}} \varsigma_{t+1}^2(\boldsymbol{\theta}), \quad (31)$$

where $\varsigma_{t+1}^2(\boldsymbol{\theta})$ depends on the chosen parameter $\boldsymbol{\theta}^*$ and is obtained as in Equation (19) except that the index j is dropped as we only explore the IMSE criterion for $d = 1$.

Figure 13 summarizes the accuracy results across 50 replicates to acquire 10 parameters and Figure 14 illustrates 10 acquired points for a single replication using the one-dimensional example with EI (first column), IMSE (second column), and EIVAR (third column) acquisition

functions to gain insights about how the acquisition functions perform. The EI criterion acquires hyper-localized parameters and results in a myopic focus on the highest posterior region. Moreover, EI is not able to explore the medium-level posterior region, and thus the posterior is not correctly estimated at the corresponding parameter values. On the other hand, IMSE acquires parameters to build an emulator with a high prediction accuracy on the entire support of the parameters. Consequently, IMSE acquires many parameters at near-zero posterior regions, and it does not bring any additional value for estimating the posterior. This can be problematic especially if the available computational budget is limited and the simulation model is expensive to evaluate since the computational resources would be wasted for evaluations outside of the region of interest. In contrast, EIVAR faithfully acquires parameters from the nonzero posterior region to better estimate the posterior. Moreover, since IMSE does not specifically target the high posterior region, the uncertainty is higher and the predictive accuracy is lower on the high posterior region as compared to EIVAR.

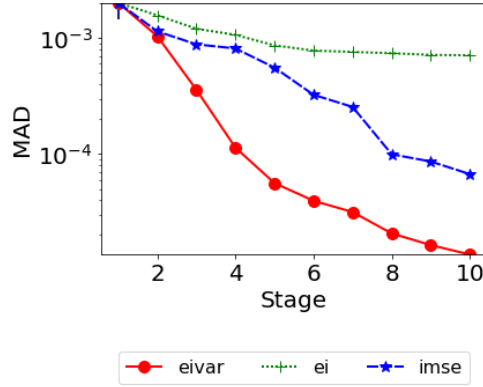


Figure 13: Comparison of three acquisition functions using a test function $\eta(\theta) = \sin(\theta) + 0.5\theta$, observation $y = -5$, and observation variance one. Algorithm 1 is used to acquire 10 parameters.

A.2 Additional Details and Experiments

In this section, we provide the details of the test functions used in our numerical experiments.

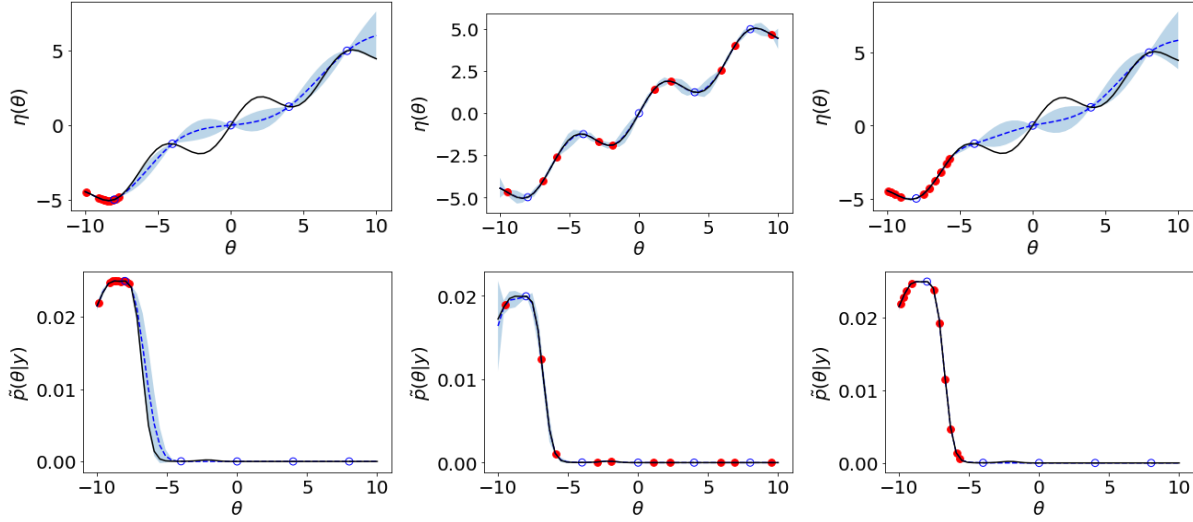


Figure 14: 10 samples (red dots) with EI (first column), IMSE (second column), and EIVAR (third column) acquisition functions for a test function $\eta(\theta) = \sin(\theta) + 0.5\theta$, observation $y = -5$, and observation variance one for $\theta \in [-10, 10]$. Red dots indicate the acquired parameters and the black line illustrates the simulation model (first row) and the posterior (second row). The blue dashed-line shows the prediction mean and the shaded area illustrates two predictive standard deviation from the mean. Blue circles illustrate the 5 samples used to initiate the proposed procedure.

- For the banana function (first plot in Figure 5), we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|\mathbf{y}) = (2\pi)^{-1}|\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} \frac{(\theta_1 - y_1)^2}{100} - \frac{1}{2} (\theta_2 + 0.03\theta_1^2 - y_2)^2 \right\} p(\boldsymbol{\theta})$$

where $\theta_1 \in [-20, 20]$ and $\theta_2 \in [-10, 5]$. In the calibration context, we take $\boldsymbol{\theta} = (\theta_1, \theta_2)$ and $\boldsymbol{\eta}(\boldsymbol{\theta}) = (\theta_1, \theta_2 + 0.03\theta_1^2)$ with $\mathbf{y} = (y_1, y_2) = (0, 3)$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements $(10^2, 1)$.

- For the bimodal function (second plot in Figure 5), we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|\mathbf{y}) = (2\pi)^{-1}|\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{\sqrt{0.2}}{2} (\theta_2 - \theta_1^2 - y_1)^2 - \frac{\sqrt{0.75}}{2} (\theta_2 - \theta_1 - y_2)^2 \right\} p(\boldsymbol{\theta})$$

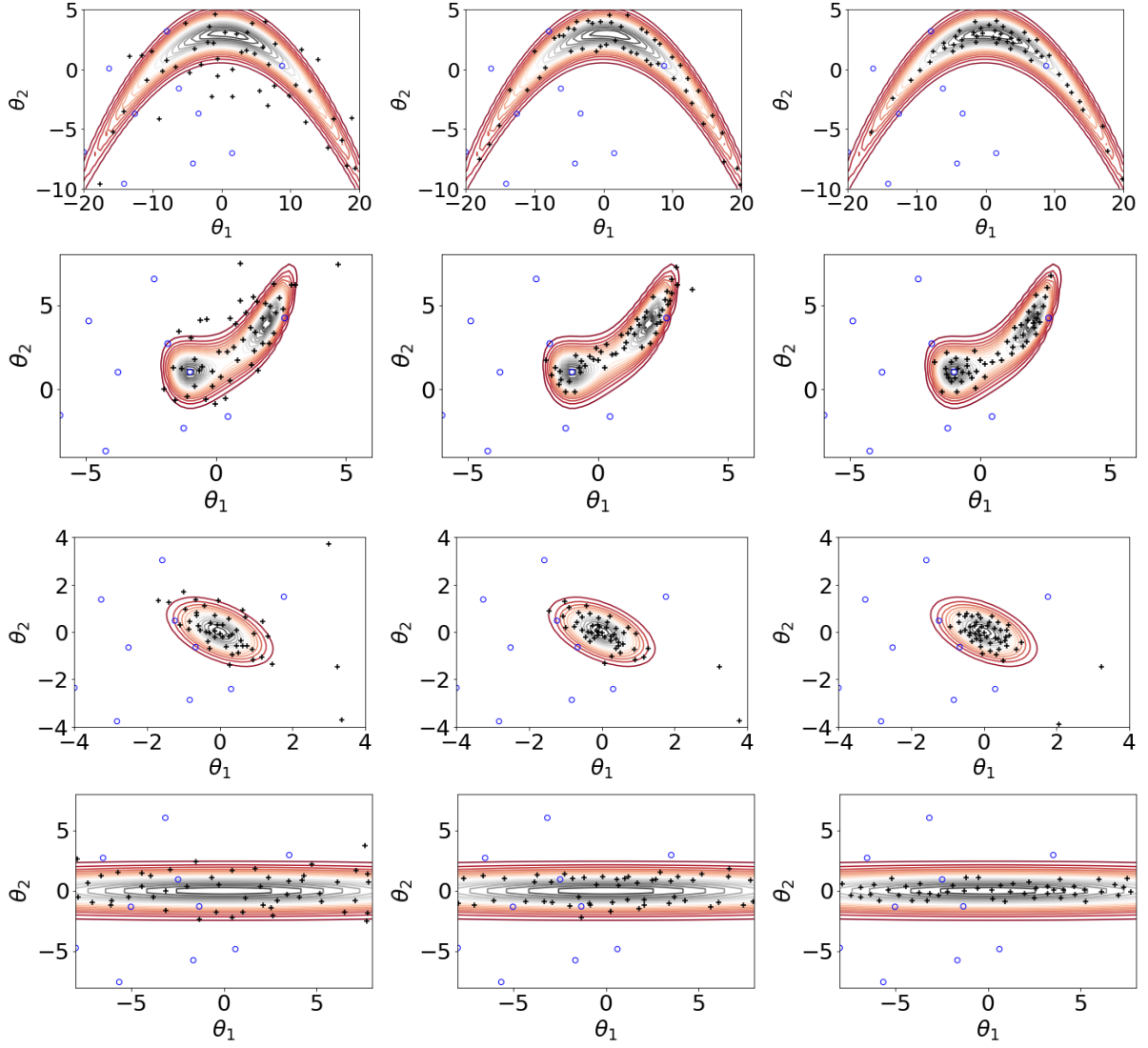


Figure 15: 50 samples (plus markers) with different acquisition functions: EIVAR (first column), MAXVAR (second column), MAXEXP (last column). Test functions illustrated in Figure 5: banana (first row), bimodal (second row), unimodal (third row), and unidentifiable (last row). Blue circles illustrate 10 samples used to initiate the proposed procedure.

where $\theta_1 \in [-6, 6]$ and $\theta_2 \in [-4, 8]$. In the calibration context, we take $\boldsymbol{\theta} = (\theta_1, \theta_2)$ and $\boldsymbol{\eta}(\boldsymbol{\theta}) = (\theta_2 - \theta_1^2, \theta_2 - \theta_1)$ with $\mathbf{y} = (y_1, y_2) = (0, 2)$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements $(1/\sqrt{0.2}, 1/\sqrt{0.75})$.

- For the unimodal function (third plot in Figure 5), we consider the unnormalized posterior

given by

$$\tilde{p}(\boldsymbol{\theta}|y) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left\{ -\frac{1}{2} \left(\frac{\theta_1^2 + \theta_1\theta_2 + \theta_2^2 - y}{2} \right)^2 \right\} p(\boldsymbol{\theta})$$

where $\theta_1 \in [-4, 4]$ and $\theta_2 \in [-4, 4]$. In the calibration context, we take $\boldsymbol{\theta} = (\theta_1, \theta_2)$ and

$\boldsymbol{\eta}(\boldsymbol{\theta}) = \theta_1^2 + \theta_1\theta_2 + \theta_2^2$ with $y = -6$ and $\sigma^2 = 4$.

- For the unidentifiable function (last plot in Figure 5), we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|\mathbf{y}) = (2\pi)^{-1} |\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} \frac{(\theta_1 - y_1)^2}{100} - \frac{1}{2} (\theta_2 - y_2)^2 \right\} p(\boldsymbol{\theta})$$

where $\theta_1 \in [-8, 8]$ and $\theta_2 \in [-8, 8]$. In the calibration context, we take $\boldsymbol{\theta} = (\theta_1, \theta_2)$

and $\boldsymbol{\eta}(\boldsymbol{\theta}) = (\theta_1, \theta_2)$ with $\mathbf{y} = (y_1, y_2) = (0, 0)$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements $(10^2, 1)$.

- For 3d function, we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|y) = (2\pi)^{-3/2} |\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y})^\top \boldsymbol{\Sigma}^{-1} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y}) \right\} p(\boldsymbol{\theta})$$

where $\boldsymbol{\theta} = (\theta_1, \theta_2, \theta_3) \in [-4, 4]^3$, $\boldsymbol{\eta}(\boldsymbol{\theta}) = (\theta_1^2 + 0.5\theta_1\theta_2 + 0.5\theta_1\theta_3, \theta_2^2 + 0.5\theta_2\theta_1 + 0.5\theta_2\theta_3, \theta_3^2 + 0.5\theta_3\theta_1 + 0.5\theta_3\theta_2)$ with $\mathbf{y} = \mathbf{0}_3$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements each of which is 0.5.

- For 6d function, we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|y) = (2\pi)^{-3} |\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y})^\top \boldsymbol{\Sigma}^{-1} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y}) \right\} p(\boldsymbol{\theta})$$

where $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_6) \in [-4, 4]^6$, $\boldsymbol{\eta}(\boldsymbol{\theta}) = (0.5\theta_1^2 + 0.5\theta_1\theta_1 + \dots + 0.5\theta_1\theta_6, 0.5\theta_2^2 + 0.5\theta_2\theta_1 + \dots + 0.5\theta_2\theta_6, \dots, 0.5\theta_6^2 + 0.5\theta_6\theta_1 + \dots + 0.5\theta_6\theta_6)$ with $\mathbf{y} = \mathbf{0}_6$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements each of which is 0.5.

- For 10d function, we consider the unnormalized posterior given by

$$\tilde{p}(\boldsymbol{\theta}|y) = (2\pi)^{-5} |\boldsymbol{\Sigma}|^{-1/2} \exp \left\{ -\frac{1}{2} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y})^\top \boldsymbol{\Sigma}^{-1} (\boldsymbol{\eta}(\boldsymbol{\theta}) - \mathbf{y}) \right\} p(\boldsymbol{\theta})$$

where $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_{10}) \in [-2, 2]^{10}$, $\boldsymbol{\eta}(\boldsymbol{\theta}) = (0.5\theta_1^2 + 0.5\theta_1\theta_1 + \dots + 0.5\theta_1\theta_6, 0.5\theta_2^2 + 0.5\theta_2\theta_1 + \dots + 0.5\theta_2\theta_6, \dots, 0.5\theta_{10}^2 + 0.5\theta_{10}\theta_1 + \dots + 0.5\theta_{10}\theta_{10})$ with $\mathbf{y} = \mathbf{0}_{10}$ and $\boldsymbol{\Sigma}$ is a diagonal matrix with diagonal elements each of which is 0.25.

Figure 15 illustrates 50 acquired points for a single replication using two-dimensional synthetic functions to gain insights about how the acquisition functions perform.

A.3 Additional Results with Minimum Energy Design

We also include the minimum energy design (a.k.a. MINED) criterion (Joseph et al., 2015, 2019) into the benchmark. To do that, we use the R package MINED (Wang and Joseph, 2022) to acquire 200 parameters, and replicate the results 50 times. We set the iteration number $K_{\text{iter}} = 8$ as illustrated in the two-dimensional example provided in the vignette. Since the acquired parameters are not ordered, we did not include MINED in the results presented in Figure 6, and instead we compare the results with a total of 200 parameters. Figure 16 illustrates the results for four toy problems presented in the paper.

A.4 Impact of Prior on EIVAR

Similar to the other Bayesian calibration procedures, the prior probability density $p(\boldsymbol{\theta})$ is critical in the proposed approach. Figure 17 illustrates an example with two parameters $\boldsymbol{\theta} = (\theta_1, \theta_2)$ in $[-2, 5] \times [-2, 5]$ to test the impact of prior on the proposed approach. In Figure 17, for all three cases, the likelihood is the same and the truncated Gaussian prior is used for each parameter with the same mean. We vary the variance of the prior to test the impact of prior on the acquired parameters. In all cases, EIVAR has places points to learn the posterior and majority of the acquired parameters are located onto the high posterior density regions. This example shows that our approach selects the most informative parameters consistently for different strength of prior density.

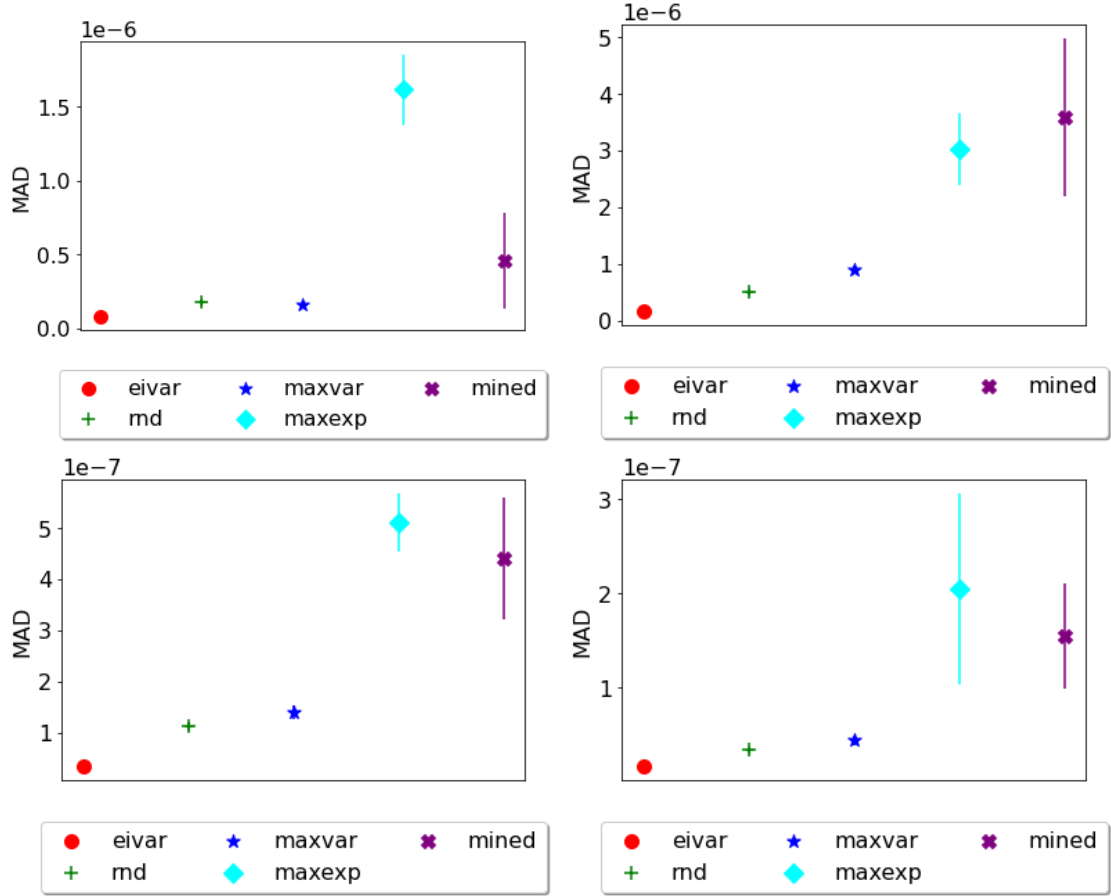


Figure 16: Comparison of different acquisition functions using the test functions illustrated in Figure 5: banana (upper-left), bimodal (upper-right), unimodal (lower-left), and unidentifiable (lower-right).

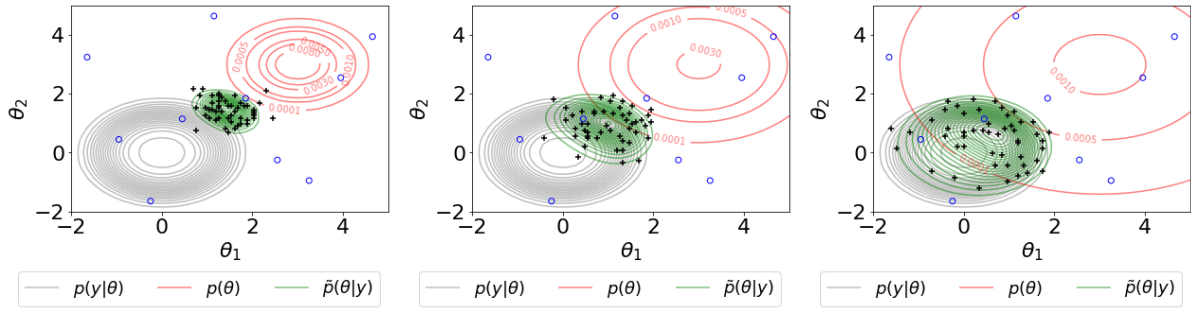


Figure 17: 50 samples (plus markers) with EIVAR acquisition function for a test function $\eta(\theta_1, \theta_2) = \theta_1^2 + \theta_2^2$, observation $y = 0$ and observation variance 2. Shown are the likelihood $p(y|\boldsymbol{\theta})$ (gray contour plot), the prior $p(\boldsymbol{\theta})$ (red contour plot), and the posterior $\tilde{p}(\boldsymbol{\theta}|y)$ (green contour plot) with varying values of variance for a truncated Gaussian prior: 0.5^2 (left panel), 1 (middle panel), and 2^2 (right panel). Blue circles illustrate the 10 samples used to initiate the proposed procedure.

References

- Azimi, J., Fern, A., and Fern, X. “Batch Bayesian Optimization via Simulation Matching.” In *Advances in Neural Information Processing Systems*. Red Hook, NY, USA: Curran Associates Inc. (2010).
- Bayarri, M. J., Berger, J. O., and Liu, F. “Modularization in Bayesian analysis, with emphasis on analysis of computer models.” *Bayesian Analysis*, 4(1):119–150 (2009).
- Bayarri, M. J., Berger, J. O., Paulo, R., Sacks, J., Cafeo, J. A., Cavendish, J., Lin, C.-H., and Tu, J. “A Framework for Validation of Computer Models.” *Technometrics*, 49(2):138–154 (2007).
- Binois, M., Huang, J., Gramacy, R. B., and Ludkovski, M. “Replication or Exploration? Sequential Design for Stochastic Simulation Experiments.” *Technometrics*, 61(1):7–23 (2019).
- Chipman, H. A., George, E. I., and McCulloch, R. E. “BART: Bayesian Additive Regression Trees.” *The Annals of Applied Statistics*, 4(1):266–298 (2010).
- Damblin, G., Barbillon, P., Keller, M., Pasanisi, A., and Parent, E. “Adaptive Numerical Designs for the Calibration of Computer Codes.” *SIAM/ASA Journal on Uncertainty Quantification*, 6(1):151–179 (2018).
- Fernández, F. L., Martino, L., Elvira, V., Delgado, D., and López-Santiago, J. “Adaptive Quadrature Schemes for Bayesian Inference via Active Learning.” *IEEE Access*, 8:208462–208483 (2020).
- Frazier, P. I. “A Tutorial on Bayesian Optimization.” (2018).
URL <https://arxiv.org/abs/1807.02811>
- Gelman, A., Carlin, J. B., Stern, H. S., and Rubin, D. B. *Bayesian Data Analysis*. Chapman and Hall/CRC, second edition (2004).

- Ginsbourger, D., Le Riche, R., and Carraro, L. “Kriging is Well-Suited to Parallelize Optimization.” *Computational Intelligence in Expensive Optimization Problems*, 131–162 (2010).
- Gramacy, R. B. *Surrogates: Gaussian Process Modeling, Design, and Optimization for the Applied Sciences*. New York, NY: CRC Press; Taylor & Francis Group (2020).
- Gu, M. and Wang, L. “Scaled Gaussian Stochastic Process for Computer Model Calibration and Prediction.” *SIAM/ASA Journal on Uncertainty Quantification*, 6(4):1555–1583 (2018).
- Higdon, D., Gattiker, J., Williams, B., and Rightley, M. “Computer Model Calibration Using High-Dimensional Output.” *Journal of the American Statistical Association*, 103(482):570–583 (2008).
- URL <http://www.jstor.org/stable/27640080>
- Higdon, D., Kennedy, M., Cavendish, J. C., Cafoe, J. A., and Ryne, R. D. “Combining Field Data and Computer Simulations for Calibration and Prediction.” *SIAM Journal on Scientific Computing*, 26(2):448–466 (2004).
- Huang, J., Gramacy, R. B., Binois, M., and Libraschi, M. “On-site surrogates for large-scale calibration.” *Applied Stochastic Models in Business and Industry*, 36(2):283–304 (2020).
- Järvenpää, M., Gutmann, M. U., Pleska, A., Vehtari, A., and Marttinen, P. “Efficient Acquisition Rules for Model-Based Approximate Bayesian Computation.” *Bayesian Analysis*, 14(2):595–622 (2019).
- Jones, D. R., Schonlau, M., and Welch, W. J. “Efficient Global Optimization of Expensive Black-Box Functions.” *Journal of Global Optimization*, 13(4):455–492 (1998).
- Joseph, V. R., Dasgupta, T., Tuo, R., and Wu, C. F. J. “Sequential Exploration of Complex Surfaces Using Minimum Energy Designs.” *Technometrics*, 57(1):64–74 (2015).

- Joseph, V. R., Wang, D., Gu, L., Lyu, S., and Tuo, R. “Deterministic Sampling of Expensive Posteriors Using Minimum Energy Designs.” *Technometrics*, 61(3):297–308 (2019).
- Kandasamy, K., Schneider, J., and Póczos, B. “Bayesian Active Learning for Posterior Estimation.” In *Proceedings of the 24th International Conference on Artificial Intelligence, IJCAI’15*, 3605–3611. AAAI Press (2015).
- Kennedy, M. C. and O’Hagan, A. “Bayesian Calibration of Computer Models.” *Journal of the Royal Statistical Society. Series B*, 63(3):425–464 (2001).
- Koermer, S., Loda, J., Noble, A., and Gramacy, R. B. “Active Learning for Simulator Calibration.” (2023).
URL <https://arxiv.org/abs/2301.10228>
- Lovell, A. E. “A Deeper Understanding of the Inputs for Reaction Theory Through Uncertainty Quantification.” Ph.D. thesis, Michigan State University (2018).
- Lovell, A. E., Nunes, F. M., Sarich, J., and Wild, S. M. “Uncertainty Quantification for Optical Model Parameters.” *Physical Review C*, 95(2) (2017).
- Marin, J.-M., Pudlo, P., Robert, C. P., and Ryder, R. J. “Approximate Bayesian Computational Methods.” *Statistics and Computing*, 22(6):1167–1180 (2012).
- O’Hagan, A. “Bayes–Hermite Quadrature.” *Journal of Statistical Planning and Inference*, 29(3):245–260 (1991).
- Plumlee, M. “Bayesian Calibration of Inexact Computer Models.” *Journal of the American Statistical Association*, 112(519):1274–1285 (2017).
- Rasmussen, C. E. and Williams, C. K. I. *Gaussian Processes for Machine Learning (Adaptive Computation and Machine Learning)*. The MIT Press (2005).

- Sacks, J., Welch, W. J., Mitchell, T. J., and Wynn, H. P. “Statistical Science.” *Design and Analysis of Computer Experiments*, 4(4):409–423 (1989).
- Santner, T. J., Williams, B. J., and Notz, W. I. *The Design and Analysis of Computer Experiments*. Springer Series in Statistics. New York: Springer, second edition (2018).
- Sauer, A., Gramacy, R. B., and Higdon, D. “Active Learning for Deep Gaussian Process Surrogates.” *Technometrics*, 1–15 (2022). To appear.
- Seo, S., Wallat, M., Graepel, T., and Obermayer, K. “Gaussian Process Regression: Active Data Selection and Test Point Rejection.” *Proceedings of the IEEE-INNS-ENNS International Joint Conference on Neural Networks. IJCNN 2000. Neural Computing: New Challenges and Perspectives for the New Millennium*, 3:241–246 (2000).
- Sürer, O., Nunes, F. M., Plumlee, M., and Wild, S. M. “Uncertainty Quantification in Breakup Reactions.” *Phys. Rev. C*, 106:024607 (2022a).
- Sürer, O., Plumlee, M., and Wild, S. M. “Parallel Implementation of Novel Uncertainty Quantification Methods (PUQ) Python Package.” <https://github.com/parallelUQ/PUQ> (2022b). Accessed: 2022-09-18.
- Thompson, I. J. “Coupled Reaction Channels Calculations in Nuclear Physics.” *Computer Physics Reports*, 7:167–212 (1988).
- Wang, D. and Joseph, V. R. *mined: Minimum Energy Designs* (2022). R package version 1.0-3. URL <https://CRAN.R-project.org/package=mined>
- Yang, H., Sürer, O., Duque, D., Morton, D. P., Singh, B., Fox, S. J., Pasco, R., Pierce, K., Rathouz, P., Valencia, V., Du, Z., Pignone, M., Escott, M. E., Adler, S. I., Johnston, S. C., and Meyers, L. A. “Design of COVID-19 Staged Alert Systems to Ensure Healthcare Capacity with Minimal Closures.” *Nature Communications*, 12(1):3767 (2021).