

Variational Bayesian surrogate modelling with application to robust design optimisation

Thomas A. Archbold, Ieva Kazlauskaite, Fehmi Cirak*

Department of Engineering, University of Cambridge, Cambridge, CB2 1PZ, UK

Abstract

Surrogate models provide a quick-to-evaluate approximation to complex computational models and are essential for multi-query problems like design optimisation. The inputs of current computational models are usually high-dimensional and uncertain. We consider Bayesian inference for constructing statistical surrogates with input uncertainties and intrinsic dimensionality reduction. The surrogates are trained by fitting to data from prevalent deterministic computational models. The assumed prior probability density of the surrogate is a Gaussian process. We determine the respective posterior probability density and parameters of the posited statistical model using variational Bayes. The non-Gaussian posterior is approximated by a simpler trial density with free variational parameters and the discrepancy between them is measured using the Kullback-Leibler (KL) divergence. We employ the stochastic gradient method to compute the variational parameters and other statistical model parameters by minimising the KL divergence. We demonstrate the accuracy and versatility of the proposed reduced dimension variational Gaussian process (RDVGP) surrogate on illustrative and robust structural optimisation problems with cost functions depending on a weighted sum of the mean and standard deviation of model outputs.

Keywords: Surrogate modelling; Multi-query problems; Bayesian inference; Variational Bayes; Gaussian processes; Robust optimisation

1. Introduction

The industrial design of engineering products relies on complex computational models typically consisting of several submodels responsible for different aspects of design and analysis, like geometric modelling, finite element analysis, etc. Prevailing computational models are costly to evaluate and quickly become a bottleneck in multi-query problems such as design optimisation, uncertainty quantification, inverse problems and model calibration. For surrogate modelling, the computational model is interpreted as a black-box function that can be interrogated by recording its outputs for given inputs. The quick-to-evaluate surrogate model is learned by fitting to data collected from the black-box function. Common surrogate modelling techniques include polynomial chaos expansions [1–4], neural networks [5–8] and Gaussian process (GP) regression [9–13]. In our primary application, robust design optimisation (RDO) [14–18], the model inputs consist of design variables like geometric dimensions and fixed (immutable) variables like forcing or constitutive parameters. The model outputs are quantities of interest like the maximum stress or compliance.

*Corresponding author

Email address: f.cirak@eng.cam.ac.uk (Fehmi Cirak)

Although most prevailing computational models are deterministic, it is critical to consider the input uncertainties (see Figure 1) to achieve more efficient and robust designs [19, 20]. The uncertainties typically arise from variations in manufacturing processes and insufficiently known operational conditions. Such intrinsic aleatoric uncertainties are critical for RDO, necessitating their inclusion in surrogate modelling, see e.g. [21, 22] for a discussion on aleatoric and epistemic uncertainties. Robust optimisation yields a design which is less sensitive to the variations in the inputs by considering a cost function consisting of the weighted sum of the expectation and standard deviation of typical objective and constraint functions used in deterministic optimisation. Surrogate-based techniques are widely used in RDO, see e.g. [23–29].

We construct a statistical surrogate model by fitting to data from the costly-to-evaluate deterministic computational model. As common in probabilistic approaches, most prominently GP regression, we pose surrogate model construction as a Bayesian inference problem [30–34]. That is, the posterior probability density representing a surrogate $f(s)$ is determined by updating a prior probability density in light of a training data set $\mathcal{D}$ using the Bayes’ rule. The prior probability density for $f(s)$ is a GP with a prescribed mean and covariance with several hyperparameters. The training data set $\mathcal{D} = \{(s_i, y_i)\}_{i=1}^n$ collects the n evaluations of the computational model with each pair consisting of an input $s_i \in \mathbb{R}^{d_s}$ and its output $y_i \in \mathbb{R}$. The assumed prior probability density is a GP defining a distribution over functions with a certain mean and covariance (smoothness). The required likelihood density is obtained from the posited statistical observation model $y = f(s) + \epsilon_y$ with an additive noise or error term ϵ_y . A crucial difference of the sketched approach from standard GP regression is that the input vector s has intrinsic aleatoric uncertainties, e.g. due to manufacturing. Consequently, the posterior probability density representing the surrogate $f(s)$ encodes in addition to epistemic also intrinsic aleatoric uncertainties that do not vanish in the limit of infinite data. In standard GP regression with a deterministic input s , all uncertainty is epistemic and tends toward zero when more data is taken into account [9, 21]. The key challenge in Bayesian inference with an uncertain input s is that the posterior density is not a GP with a closed-form solution (even when both the prior probability density and the likelihood density are Gaussians). Hence, it is necessary to use approximation techniques, like Markov chain Monte Carlo (MCMC), Laplace approximation, or variational Bayes (VB) [32].

In design optimisation, as in most other multi-query applications, the dimension of the input variable s must be restricted for computational tractability. The dimension of the input variable is typically large from the outset, so a suitable dimensionality reduction technique must be employed. Dimensionality reduction is an extensively studied topic, and many techniques are available [35–39]. Two-step approaches for sequentially combining dimensionality reduction with standard GP regression have also been proposed [40–42]. Departing from earlier approaches and adopting a fully Bayesian viewpoint, we implement dimensionality reduction by positing the statistical observation model $y = f(z) + \epsilon_y$ with the low-dimensional (latent) input vector $z = W^T s$ and the non-square (tall) orthogonal matrix W , i.e. $W^T W = I$. A reduction in input dimension is possible because in most applications the intrinsic dimensionality of the input-output manifold is usually low. We assume here that the solution lies on a hyperplane in the high-dimensional input-output space. Furthermore, as noted before, the input vector s and consequently, the latent vector z have intrinsic aleatoric uncertainties whereas the matrix W is deterministic. The posterior probability density of $f(s)$ and all statistical model parameters including the dimension of z and entries of W , can be obtained by consequent application of the Bayes’ rule. Expectedly, the posterior probability densities and the marginal likelihood density for determining the model parameters have no analytically tractable solution and must be approximated.

VB provides an optimisation-based formulation for approximating the posterior density and parameters of statistical models [43–46]. In VB the true posterior density is approximated by a trial density and the distance between the two is measured using the Kullback-Leibler (KL) divergence. The chosen trial density

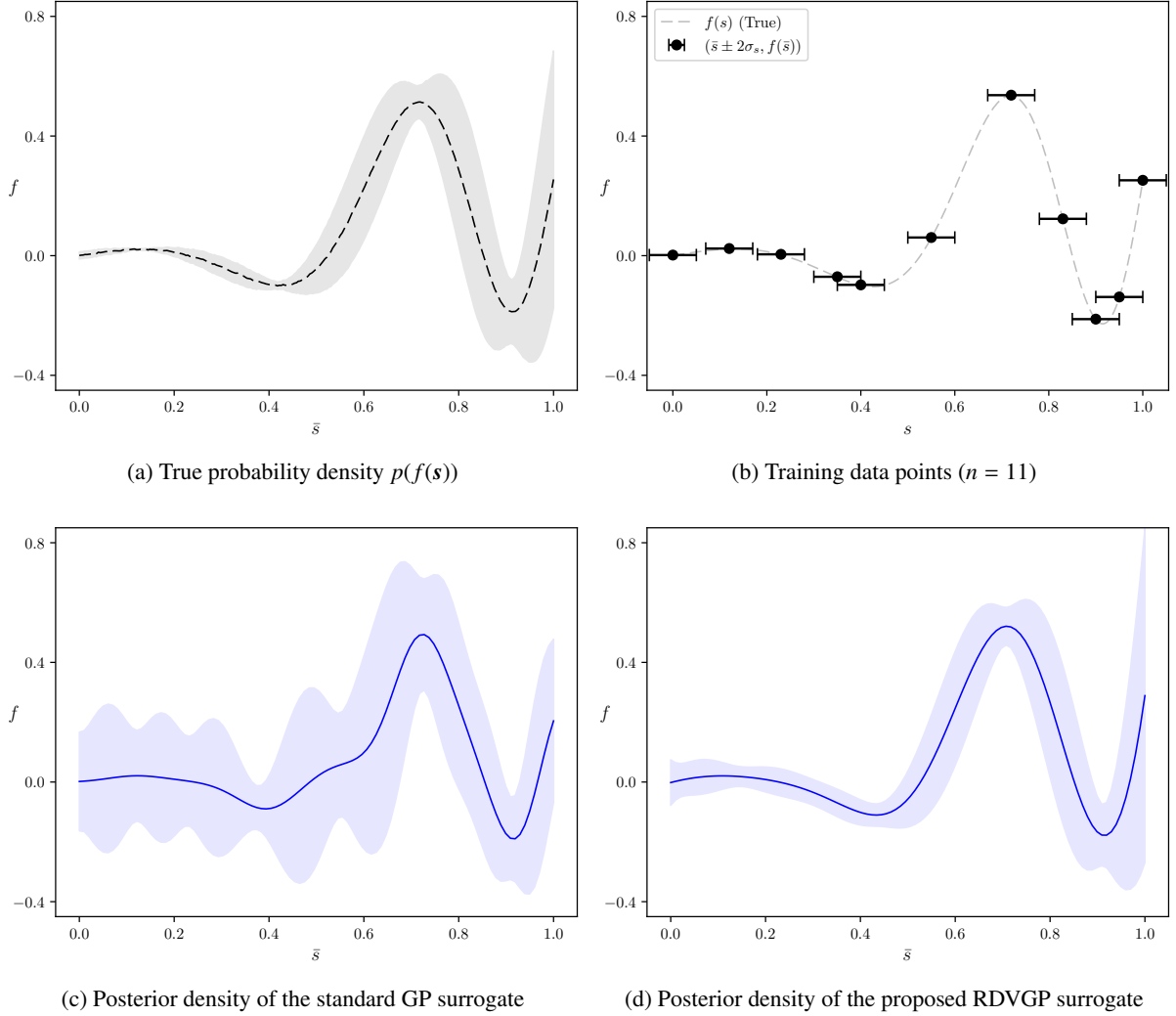


Figure 1: Comparison of standard GP and the proposed RDVGP surrogates for a function $f(s) = -0.5s \sin(3\pi s^2) + 0.25s$ and the normally distributed input variable with $p(s) = \mathcal{N}(\bar{s}, \sigma_s^2)$; see also Section 4.1. In the plots the lines indicate the mean and the shaded areas the 95% confidence intervals. Shown is the (a) true probability density $p(f(\bar{s})) = \int p(f|s)p(s) ds$ obtained by MC sampling, (b) the $n = 11$ training data points comprising the training set $\mathcal{D}$, (c) the inferred GP posterior probability density $p(f(\bar{s})|\mathcal{D})$, and (d) the inferred RDVGP posterior probability density $p(f(\bar{s})|\mathcal{D})$. Comparing (a), (c), and (d) it is evident that the RDVGP posterior is much closer to the true probability density than the standard GP posterior.

is a parameterised probability density with free variational parameters. It can be shown that minimising the KL divergence is equivalent to maximising the evidence lower bound (ELBO). As the name suggests, the ELBO is a lower bound for the log-evidence also called the log-marginal likelihood, which quantifies how well the posited statistical observation model explains the training data. In our statistical surrogate model described above, the approximated density is the posterior density of the surrogate $f(s)$ and the model hyperparameters consist of the entries of the projection matrix $\mathbf{W}$ and the hyperparameters of the prior. We select as the trial density a Gaussian with a diagonal covariance structure and choose its mean and variance as the variational parameters. All the model hyperparameters and the variational parameters are determined by maximising the ELBO. In the proposed implementation, we maximise the ELBO using a

stochastic gradient method [47], specifically the Cayley ADAM optimiser [48], which automatically ensures the orthogonality of $\mathbf{W}$. Unfortunately, the evaluation of the ELBO has the computational complexity $O(n^3)$, where n is the size of the training data set $\mathcal{D}$. Therefore, we introduce a sparse approximation of $f(\mathbf{z})$ using m pseudo training data points reducing complexity to $O(nm^2)$ [49, 50]. Although the number m is user prescribed, the coordinates of the pseudo input points are determined as part of VB. The overall solution approach implemented is akin to the latent variable GP models widely used in machine learning [51–53]. For brevity, we refer to this as the reduced dimension variational Gaussian process (RDVGP) surrogate.

This paper is structured as follows. In Section 2, we define RDO problems in terms of objective and constraint functions describing the behavior of complex engineering systems, and review GP regression for emulating these functions. In Section 3, we present the statistical surrogate model by postulating a graphical model with intrinsic statistical assumptions and deriving the ELBO from the joint probability density using VB, later extending to include sparsity. We detail the algorithm used to train the RDVGP surrogate, and provide the metrics used to verify its accuracy when compared to the true solution estimated using Monte Carlo (MC) sampling. Four examples are introduced in Section 4 to demonstrate the accuracy and versatility of the predicted marginal posterior probability densities for solving RDO problems. Finally, Section 5 concludes the paper and provides promising directions for further research.

2. Robust design optimisation and Gaussian processes

In this section, we review RDO and GP regression. We introduce the problem in terms of inputs with uncertainties as set out by prescribed probability densities and we show how the problem is formulated in terms of objective and constraint functions from the computational model.

2.1. Review of robust design optimisation

The input variables $\mathbf{s} \in \mathbb{R}^{d_s}$ of deterministic computational models may be composed of immutable variables $\mathbf{s}_f \in \mathbb{R}^{d_f}$ that cannot be varied (such as the partial differential equation (PDE) source terms), and design variables $\mathbf{s}_d \in \mathbb{R}^{d_d}$ (such as CAD geometry parameters) such that $\mathbf{s} = (\mathbf{s}_d^\top \mathbf{s}_f^\top)^\top$.

In many applications, the computational model output is not directly the quantity of interest, but an argument provided to an objective or constraint function representing model behavior, and may be computed from a finite element (FE) model. For example, in structural mechanics, compliance objectives and stress constraints are both functions of the PDE solution. For brevity, the objective $J : \mathbb{R}^{d_s} \rightarrow \mathbb{R}$ and constraint $H^{(j)} : \mathbb{R}^{d_s} \rightarrow \mathbb{R}$ may be denoted as functions of the input variables $\mathbf{s}$. A total of d_y functions, consisting (without loss of generality) of a single objective and $(d_y - 1)$ constraint functions are considered.

This paper assumes the input variable probability density is prescribed a priori, which could be Gaussian $p(\mathbf{s}) = \mathcal{N}(\bar{\mathbf{s}}, \Sigma_s)$, where $\bar{\mathbf{s}} \in \mathbb{R}^{d_s}$ is the mean and $\Sigma_s \in \mathbb{R}^{d_s \times d_s}$ is the covariance matrix. For example, a CAD embodiment of engineering products may consist of components with geometric variation represented by a Gaussian probability density, where an engineer models the product in CAD using the mean design variables $\bar{\mathbf{s}}$. Since the input variables are random, the objective $J(\mathbf{s})$ and constraint $H^{(j)}(\mathbf{s})$ functions are also random and the RDO problem may be solved using the weighted sum cost function

$$\begin{aligned} \min_{\bar{\mathbf{s}} \in D_{\bar{\mathbf{s}}}} \quad & \frac{(1 - \alpha)}{\bar{\mu}} \mathbb{E}(J(\mathbf{s})) + \frac{\alpha}{\bar{\sigma}} \sqrt{\text{var}(J(\mathbf{s}))}, \\ \text{s.t.} \quad & \mathbb{E}(H^{(j)}(\mathbf{s})) + \beta^{(j)} \sqrt{\text{var}(H^{(j)}(\mathbf{s}))} \leq 0, \quad j \in \{1, 2, \dots, d_y - 1\}, \\ & \sqrt{\text{var}(H^{(j)}(\mathbf{s}))} \leq \sigma^{(j)}, \\ & D_{\bar{\mathbf{s}}} = \{\bar{\mathbf{s}} \in \mathbb{R}^{d_s} \mid \bar{\mathbf{s}}^{(l)} \leq \bar{\mathbf{s}} \leq \bar{\mathbf{s}}^{(u)}\}, \end{aligned} \tag{1}$$

where $\alpha \in \mathbb{R}^+$ is a prescribed weighting factor ($0 \leq \alpha \leq 1$), $\bar{\mu} \in \mathbb{R}$ and $\bar{\sigma} \in \mathbb{R}$ are normalisation constants, $\beta^{(j)} \in \mathbb{R}^+$ is a prescribed feasibility index, and $\sigma^{(j)} \in \mathbb{R}^+$ is an upper limit on the standard deviation of the j^{th} constraint. The mean input variables $\bar{s}$ are bounded between lower and upper limits $\bar{s}^{(l)} = ((\bar{s}_d^{(l)})^\top \bar{s}_f^{(l)})^\top$ and $\bar{s}^{(u)} = ((\bar{s}_d^{(u)})^\top \bar{s}_f^{(u)})^\top$ respectively. The objective of RDO is to optimise the mean design variables $\bar{s}_d$.

2.2. Review of Gaussian processes

In RDO, the objective $J(s)$ and constraint $H^{(j)}(s)$ functions are emulated by a GP surrogate denoted $f(s)$. A single observation from each function denoted $y \in \mathbb{R}$, follows a statistical observation model, see the graphical model in Figure 2,

$$y = f(s) + \epsilon_y, \quad \epsilon_y \sim \mathcal{N}(0, \sigma_y^2), \quad (2)$$

where $\sigma_y \in \mathbb{R}^+$ denotes the noise or error standard deviation (such as FE discretisation error), and the GP $f : \mathbb{R}^{d_s} \rightarrow \mathbb{R}$ maps between the input variables s and observation y , and has the prior

$$f(s) \sim \mathcal{GP}(\bar{f}(s), c(s, s')). \quad (3)$$

Without loss of generality, we choose a zero-mean prior $\bar{f}(s) = 0$ and squared exponential covariance function $c : \mathbb{R}^{d_s} \times \mathbb{R}^{d_s} \rightarrow \mathbb{R}$ (other types of kernels may be used when the objective is less smooth),

$$c(s, s') = \sigma_f^2 \exp\left(-\sum_{i=1}^{d_s} \frac{(s_i - s'_i)^2}{2\ell_i}\right), \quad (4)$$

where $\sigma_f \in \mathbb{R}^+$ is a scaling parameter, $\ell_i \in \mathbb{R}^+$ is a length scale parameter, and s_i denotes the i^{th} component of input vector s . The hyperparameters of the GP are summarised as $\theta = \{\sigma_f, \ell_1, \ell_2, \dots, \ell_{d_s}\}$.

Let $\mathcal{D} = \{(s_i, y_i) | i \in \mathbb{N}, i \leq n\}$ denote the training data set, which is collected into the input variable matrix $\mathbf{S} = (s_1 \ s_2 \ \dots \ s_n)^\top \in \mathbb{R}^{n \times d_s}$ and observation vector $\mathbf{y} = (y_1 \ y_2 \ \dots \ y_n)^\top \in \mathbb{R}^n$. The joint probability density of observations $\mathbf{y}$ and target output variables $\mathbf{f} = (f_1 \ f_2 \ \dots \ f_n)^\top \in \mathbb{R}^n$ can be factorised as

$$p_\theta(\mathbf{y}, \mathbf{f}) = p_{\sigma_y}(\mathbf{y}|\mathbf{f})p_\theta(\mathbf{f}), \quad (5)$$

which is expanded using the postulated conditional independence with the likelihood (observations are assumed to be independent and identically distributed (iid))

$$p_{\sigma_y}(\mathbf{y}|\mathbf{f}) = \prod_{i=1}^n p_{\sigma_y}(y_i|f_i) = \mathcal{N}(\mathbf{f}, \sigma_y^2 \mathbf{I}), \quad (6)$$

where y_i and f_i denote the observation and target output variable for the i^{th} training data sample, and the GP prior probability density

$$p_\theta(\mathbf{f}) = \mathcal{N}(\mathbf{0}, \mathbf{C}_{SS}). \quad (7)$$

The entries of the covariance $\mathbf{C}_{SS} \in \mathbb{R}^{n \times n}$ consist of the covariance function $c(s, s')$ where s and s' are rows of $\mathbf{S}$. The marginal likelihood

$$p_\theta(\mathbf{y}) = \int p_{\sigma_y}(\mathbf{y}|\mathbf{f})p_\theta(\mathbf{f}) d\mathbf{f} = \mathcal{N}(\mathbf{0}, \mathbf{C}_{SS} + \sigma_y^2 \mathbf{I}), \quad (8)$$

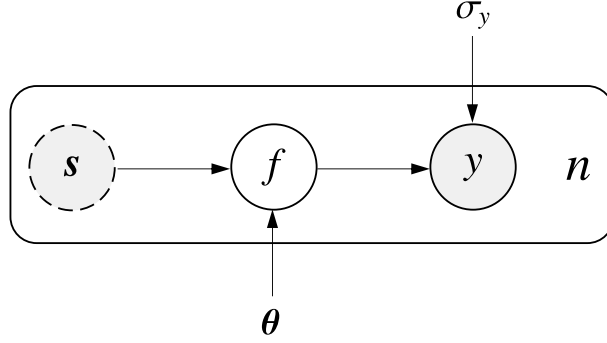


Figure 2: Graphical model of a standard GP where n is the size of the training data set $\mathcal{D}$, s is the vector of (deterministic) input variables, f is the (random) target output variable, y is the (random) noisy observation, and θ and σ_y are model hyperparameters.

is derived by marginalising over the target output variables f . In empirical Bayes' methods, the optimum hyperparameters $\Theta = \theta \cup \{\sigma_y\}$ are computed by maximising the log marginal likelihood

$$\Theta^* = \arg \max_{\Theta} \ln p_{\Theta}(\mathbf{y}), \quad (9)$$

where

$$\ln p_{\Theta}(\mathbf{y}) = -\frac{n}{2} \ln(2\pi) - \frac{1}{2} \ln |\mathbf{C}_{SS} + \sigma_y^2 \mathbf{I}| - \frac{1}{2} \mathbf{y}^T (\mathbf{C}_{SS} + \sigma_y^2 \mathbf{I})^{-1} \mathbf{y}, \quad (10)$$

and $|\cdot|$ denotes the determinant. For any new test input variables $\mathbf{S}_* \in \mathbb{R}^{n_* \times d_s}$, the posterior probability density of the test output variables $\mathbf{f}_* \in \mathbb{R}^{n_*}$ is given by

$$p_{\Theta}(\mathbf{f}_*|\mathbf{y}) = \int p_{\Theta}(\mathbf{f}_*|\mathbf{f}) p_{\Theta}(\mathbf{f}|\mathbf{y}) d\mathbf{f}, \quad (11)$$

which assumes conditional independence between the test output variables $\mathbf{f}_*$ and target output variables $\mathbf{f}$ when observations $\mathbf{y}$ are given. The posterior probability densities $p_{\Theta}(\mathbf{f}_*|\mathbf{f})$ and $p_{\Theta}(\mathbf{f}|\mathbf{y})$ can be determined using Bayes rule

$$p_{\Theta}(\mathbf{f}_*|\mathbf{f}) = \frac{p_{\Theta}(\mathbf{f}|\mathbf{f}_*) p_{\Theta}(\mathbf{f}_*)}{p_{\Theta}(\mathbf{f})}, \quad p_{\Theta}(\mathbf{f}|\mathbf{y}) = \frac{p_{\sigma_y}(\mathbf{y}|\mathbf{f}) p_{\Theta}(\mathbf{f})}{p_{\Theta}(\mathbf{y})}. \quad (12)$$

Alternatively because all involved probability densities are Gaussian, the posterior probability density $p_{\Theta}(\mathbf{f}_*|\mathbf{y})$ is derived using the joint probability density

$$p_{\Theta}(\mathbf{f}_*, \mathbf{y}) = \mathcal{N} \left(\begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{C}_{S_* S_*} & \mathbf{C}_{S_* S} \\ \mathbf{C}_{S S_*} & \mathbf{C}_{SS} + \sigma_y^2 \mathbf{I} \end{pmatrix} \right). \quad (13)$$

The posterior probability density of test output variables $\mathbf{f}_*$ is given by

$$p_{\Theta}(\mathbf{f}_*|\mathbf{y}) = \mathcal{N} \left(\mathbf{C}_{S_* S} (\mathbf{C}_{SS} + \sigma_y^2 \mathbf{I})^{-1} \mathbf{y}, \mathbf{C}_{S_* S_*} - \mathbf{C}_{S_* S} (\mathbf{C}_{SS} + \sigma_y^2 \mathbf{I})^{-1} \mathbf{C}_{S S_*} \right), \quad (14)$$

which is derived by conditioning the joint probability density $p_{\Theta}(\mathbf{f}_*, \mathbf{y})$ on observations $\mathbf{y}$. The evaluation of the posterior probability density becomes computationally expensive for large training data sets due to

the inversion of the covariance matrix $\mathbf{C}_{SS}$, which involves $O(n^3)$ operations.

So far the input variables $\mathbf{S}$ have been assumed deterministic; however, if they are random and sampled from a prescribed probability density $p(\mathbf{S})$, the posterior probability density of the target output variables

$$p_{\Theta}(\mathbf{f}|\mathbf{y}) = \frac{p_{\sigma_y}(\mathbf{y}|\mathbf{f}) \int p_{\theta}(\mathbf{f}|\mathbf{S})p(\mathbf{S}) d\mathbf{S}}{p_{\Theta}(\mathbf{y})}, \quad (15)$$

becomes analytically intractable because the covariance $\mathbf{C}_{SS}$ of the GP prior $p_{\theta}(\mathbf{f}|\mathbf{S})$ depends non-linearly according to (4), on the input variables $\mathbf{S}$.

3. Statistical surrogate model

In this section, we generalise standard GP regression to random input variables with intrinsic dimensionality reduction using a latent variable formulation. This is later extended to a sparse formulation of the latent variables, which reduces the computational expense of performing inference with input variable uncertainty. Finally, an algorithm for training the model and metrics for measuring accuracy are provided.

3.1. Reduced dimension variational Gaussian process

We introduce a low-dimensional latent variable representation for each input variable $\mathbf{s}$, which is computed using the statistical observation model given by

$$\mathbf{z} = \mathbf{W}^T \mathbf{s}, \quad \mathbf{s} \sim \mathcal{N}(\bar{\mathbf{s}}, \mathbf{\Sigma}_s), \quad (16a)$$

$$y = f(\mathbf{z}) + \epsilon_y, \quad \epsilon_y \sim \mathcal{N}(0, \sigma_y^2). \quad (16b)$$

The low-dimensional latent variables $\mathbf{z} \in \mathbb{R}^{d_z}$, ($d_z < d_s$) are mapped by a GP surrogate $f(\mathbf{z})$ to an observation y following the posited statistical observation model. The linear projection matrix $\mathbf{W} \in \mathbb{R}^{d_s \times d_z}$ is composed of basis vectors of the low-dimensional subspace, and the mean $\bar{\mathbf{s}}$ and covariance $\mathbf{\Sigma}_s$ of the input variable probability density $p(\mathbf{s})$ are prescribed a priori. For a training data set $\mathcal{D}$ of size n , the joint probability density of the observations $\mathbf{y}$, target output variables $\mathbf{f}$, latent variables $\mathbf{Z} = (\mathbf{z}_1 \ \mathbf{z}_2 \ \dots \ \mathbf{z}_n)^T \in \mathbb{R}^{n \times d_z}$, and input variables $\mathbf{S}$, is factorised using the postulated conditional independence, see Figure 3, such that

$$p_{\Theta}(\mathbf{y}, \mathbf{f}, \mathbf{Z}, \mathbf{S}) = p_{\sigma_y}(\mathbf{y}|\mathbf{f})p_{\theta}(\mathbf{f}|\mathbf{Z})p_W(\mathbf{Z}|\mathbf{S})p(\mathbf{S}). \quad (17)$$

The likelihood $p_{\sigma_y}(\mathbf{y}|\mathbf{f})$ is the same as in GP regression (6). The zero-mean GP prior probability density

$$p_{\theta}(\mathbf{f}|\mathbf{Z}) = \mathcal{N}(\mathbf{0}, \mathbf{C}_{ZZ}), \quad (18)$$

has the covariance matrix $\mathbf{C}_{ZZ} \in \mathbb{R}^{n \times n}$ with the entries $c(\mathbf{z}, \mathbf{z}')$ as defined in (4), where $\mathbf{z}$ and $\mathbf{z}'$ are rows of $\mathbf{Z}$. The conditional probability density

$$p_W(\mathbf{Z}|\mathbf{S}) = \prod_{i=1}^n \delta(\mathbf{z}_i - \mathbf{W}^T \mathbf{s}_i), \quad (19)$$

readily follows from (16a), where $\delta(\cdot)$ is the Dirac delta function. The input variable probability density

$$p(\mathbf{S}) = \prod_{i=1}^n \mathcal{N}(\bar{\mathbf{s}}_i, \mathbf{\Sigma}_s), \quad (20)$$

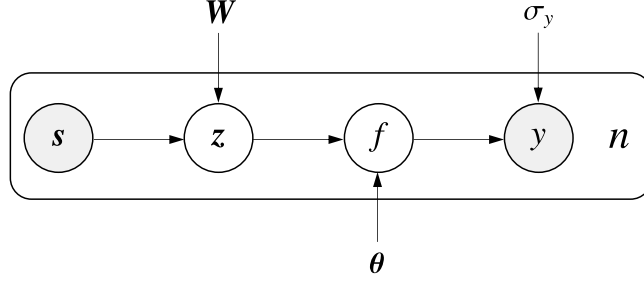


Figure 3: Graphical model of the RDVGP surrogate where n is the size of the training data set $\mathcal{D}$, s is the random input variable vector, z is the low-dimensional unobserved latent variable vector, f is the target output variable, y is the noisy observations, and W , θ and σ_y are model hyperparameters.

is prescribed a priori. We assume the input variables are uncorrelated and the covariance Σ_s does not depend on the sample index i . The input variables S are marginalised out of the joint probability density, yielding

$$p_{\Theta}(y, f, Z) = p_{\sigma_y}(y|f)p_{\theta}(f|Z)p_W(Z), \quad (21)$$

where

$$p_W(Z) = \int p_W(Z|S)p(S) dS = \prod_{i=1}^n \mathcal{N}(W^T \bar{s}_i, W^T \Sigma_s W). \quad (22)$$

The model hyperparameters $\Theta = \theta \cup \{W, \sigma_y\}$ include in addition to the hyperparameters of the GP prior θ , the observation noise or error standard deviation σ_y and entries of the linear projection matrix W . Although the posterior probability density of all unobserved variables

$$p_{\Theta}(f, Z|y) = \frac{p_{\Theta}(y, f, Z)}{p_{\Theta}(y)}, \quad (23)$$

follows from Bayes' rule, it is not available in closed-form without sampling.

In VB, the posterior is approximated with a trial probability density $q_{\theta, \psi}(f, Z)$ where the KL divergence between the two probability densities is given by

$$\begin{aligned} D_{KL}(q_{\theta, \psi}(f, Z) \| p_{\Theta}(f, Z|y)) &= \int q_{\theta, \psi}(f, Z) \ln \left(\frac{q_{\theta, \psi}(f, Z)}{p_{\Theta}(f, Z|y)} \right) dZ df \\ &= \int q_{\theta, \psi}(f, Z) \ln \left(\frac{q_{\theta, \psi}(f, Z)p_{\Theta}(y)}{p_{\Theta}(y, f, Z)} \right) dZ df \\ &= \int q_{\theta, \psi}(f, Z) \ln \left(\frac{q_{\theta, \psi}(f, Z)}{p_{\Theta}(y, f, Z)} \right) dZ df + \ln p_{\Theta}(y) \\ &= -\mathcal{F}(y) + \ln p_{\Theta}(y). \end{aligned} \quad (24)$$

Since KL divergence is non-negative $D_{KL}(\cdot \| \cdot) \geq 0$,

$$\ln p_{\Theta}(y) \geq \int q_{\theta, \psi}(f, Z) \ln \left(\frac{p_{\Theta}(y, f, Z)}{q_{\theta, \psi}(f, Z)} \right) dZ df. \quad (25)$$

The KL divergence is minimised by maximising the ELBO $\mathcal{F}(\mathbf{y})$. We assume a trial probability density

$$q_{\theta,\psi}(\mathbf{f}, \mathbf{Z}) = p_{\theta}(\mathbf{f}|\mathbf{Z})q_{\psi}(\mathbf{Z}), \quad (26)$$

with the Gaussian prior probability density $p_{\theta}(\mathbf{f}|\mathbf{Z})$ given by (18), and the trial probability density

$$q_{\psi}(\mathbf{Z}) := \prod_{i=1}^n q_{\psi}(z_i) := \prod_{i=1}^n \mathcal{N}(\tilde{\boldsymbol{\mu}}_{z,i}, \tilde{\boldsymbol{\Sigma}}_{z,i}). \quad (27)$$

Variational parameters $\boldsymbol{\psi} = \{(\tilde{\boldsymbol{\mu}}_{z,i}, \tilde{\boldsymbol{\Sigma}}_{z,i}) \mid i \in \mathbb{N}, i \leq n\}$ consist of the components of the mean vector $\tilde{\boldsymbol{\mu}}_{z,i} \in \mathbb{R}^{d_z}$ and the entries of the diagonal covariance matrix $\tilde{\boldsymbol{\Sigma}}_{z,i} \in \mathbb{R}^{d_z \times d_z}$. The ELBO is expanded by introducing the factorised joint probability density $p_{\Theta}(\mathbf{y}, \mathbf{f}, \mathbf{Z})$ that follows from (21). The trial probability density $q_{\theta,\psi}(\mathbf{f}, \mathbf{Z})$ such that

$$\begin{aligned} \mathcal{F}(\mathbf{y}) &= \int p_{\theta}(\mathbf{f}|\mathbf{Z})q_{\psi}(\mathbf{Z}) \ln \left(\frac{p_{\sigma_y}(\mathbf{y}|\mathbf{f})p_{\theta}(\mathbf{f}|\mathbf{Z})p_W(\mathbf{Z})}{p_{\theta}(\mathbf{f}|\mathbf{Z})q_{\psi}(\mathbf{Z})} \right) d\mathbf{Z}d\mathbf{f} \\ &= \int p_{\theta}(\mathbf{f}|\mathbf{Z})q_{\psi}(\mathbf{Z}) \left(\ln p_{\sigma_y}(\mathbf{y}|\mathbf{f}) + \ln \left(\frac{p_W(\mathbf{Z})}{q_{\psi}(\mathbf{Z})} \right) \right) d\mathbf{Z}d\mathbf{f} \\ &= \mathbb{E}_{q_{\psi}(\mathbf{Z})} \left(\mathbb{E}_{p_{\theta}(\mathbf{f}|\mathbf{Z})} \left(\ln p_{\sigma_y}(\mathbf{y}|\mathbf{f}) \right) \right) - D_{KL} \left(q_{\psi}(\mathbf{Z}) \parallel p_W(\mathbf{Z}) \right). \end{aligned} \quad (28)$$

The KL divergence term in the ELBO is analytically tractable given that $q_{\psi}(\mathbf{Z})$ and $p_W(\mathbf{Z})$ belong to the exponential family of probability densities (Appendix A). The expectations can be evaluated by MC sampling with $\mathbf{Z} \sim q_{\psi}(\mathbf{Z})$ and $\mathbf{f} \sim p_{\theta}(\mathbf{f}|\mathbf{Z})$. This is computationally expensive for large sample sizes (when n is large) since sampling of $\mathbf{f}$ involves $O(n^3)$ operations. The ELBO $\mathcal{F}(\mathbf{y})$ can be extended to multiple observation vectors using the mean field approximation (Appendix B).

3.2. Sparse formulation

To reduce the computational expense while retaining the properties of the RDVGP surrogate, the training data set $\mathcal{D}$ is augmented with m (where $m \ll n$) pseudo output variables $\tilde{\mathbf{f}} = (\tilde{f}_1 \ \tilde{f}_2 \ \dots \ \tilde{f}_m)^{\top} \in \mathbb{R}^m$. The pseudo output variables are trainable model parameters that when sampled at the pseudo latent variables $\tilde{\mathbf{Z}} = (\tilde{z}_1 \ \tilde{z}_2 \ \dots \ \tilde{z}_m)^{\top} \in \mathbb{R}^{m \times d_z}$, accurately explain the observations, see e.g. [50]. The pseudo latent variables also called inducing points, are also considered as hyperparameters. The joint probability density $p_{\Theta}(\mathbf{y}, \mathbf{f}, \mathbf{Z})$ given by (21) is augmented with the pseudo output variables $\tilde{\mathbf{f}}$ such that

$$p_{\Theta}(\mathbf{y}, \mathbf{f}, \tilde{\mathbf{f}}, \mathbf{Z}) = p_{\sigma_y}(\mathbf{y}|\mathbf{f})p_{\theta}(\mathbf{f}|\tilde{\mathbf{f}}, \mathbf{Z})p_{\theta}(\tilde{\mathbf{f}})p_W(\mathbf{Z}), \quad (29)$$

using the assumed conditional independence structure (Figure 4). The new GP prior probability density

$$p_{\theta}(\mathbf{f}|\tilde{\mathbf{f}}, \mathbf{Z}) = \mathcal{N}(\mathbf{C}_{Z\tilde{Z}}\mathbf{C}_{\tilde{Z}\tilde{Z}}^{-1}\tilde{\mathbf{f}}, \mathbf{C}_{ZZ} - \mathbf{C}_{Z\tilde{Z}}\mathbf{C}_{\tilde{Z}\tilde{Z}}^{-1}\mathbf{C}_{\tilde{Z}Z}), \quad (30)$$

is derived by conditioning the joint GP prior probability density $p_{\theta}(\mathbf{f}, \tilde{\mathbf{f}}|\mathbf{Z}) = p_{\theta}(\mathbf{f}|\mathbf{Z})p_{\theta}(\tilde{\mathbf{f}})$ on the pseudo output variables $\tilde{\mathbf{f}}$, where

$$p_{\theta}(\mathbf{f}, \tilde{\mathbf{f}}|\mathbf{Z}) = \mathcal{N} \left(\begin{pmatrix} \mathbf{0} \\ \mathbf{0} \end{pmatrix}, \begin{pmatrix} \mathbf{C}_{ZZ} & \mathbf{C}_{Z\tilde{Z}} \\ \mathbf{C}_{\tilde{Z}Z} & \mathbf{C}_{\tilde{Z}\tilde{Z}} \end{pmatrix} \right). \quad (31)$$

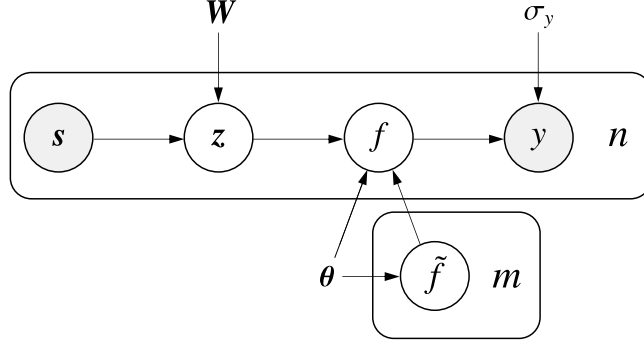


Figure 4: Graphical model of the sparse RDVGP surrogate where n is the size of the training data set $\mathcal{D}$, s is the random input variable vector, z is the low-dimensional unobserved latent variable vector, f is the target output variable, $\tilde{f}$ is the pseudo output variable with m realisations, y is the noisy observation, and W , θ (which includes pseudo latent variables $\tilde{Z}$) and σ_y are model hyperparameters.

This is similar to the derivation of standard GP regression, where $\theta = \{\tilde{Z}, \sigma_f, \ell_1, \ell_2, \dots, \ell_{d_z}\}$ are hyperparameters of $p_\theta(f, \tilde{f} | Z)$. The covariance $C_{\tilde{Z}\tilde{Z}} \in \mathbb{R}^{m \times m}$ consists of the entries $c(\tilde{z}, \tilde{z}')$ (according to (4)) where $\tilde{z}$ and $\tilde{z}'$ are rows of $\tilde{Z}$. Similarly, the covariance $C_{ZZ} \in \mathbb{R}^{n \times n}$ consists of entries $c(z, \tilde{z})$ where z is a row of Z . The posterior probability density of all unobserved variables

$$p_\theta(f, \tilde{f}, Z | y) = \frac{p_\theta(y, f, \tilde{f}, Z)}{p_\theta(y)} = p_\theta(f | \tilde{f}, Z) p_\theta(\tilde{f} | y) p_W(Z), \quad (32)$$

is analytically intractable and is therefore approximated by a Gaussian trial probability density

$$q_{\theta, \Psi}(f, \tilde{f}, Z) = p_\theta(f | \tilde{f}, Z) q_\omega(\tilde{f}) q_\Psi(Z), \quad (33)$$

where $\Psi = \psi \cup \omega$ are variational parameters. We approximate the analytically intractable posterior probability density $p_\theta(\tilde{f} | y)$ with the Gaussian trial density

$$q_\omega(\tilde{f}) := \mathcal{N}(\tilde{\mu}_{\tilde{f}}, \tilde{\Sigma}_{\tilde{f}}), \quad (34)$$

where $\omega = \{\tilde{\mu}_{\tilde{f}}, \tilde{\Sigma}_{\tilde{f}}\}$ are the free variational parameters. The mean vector $\tilde{\mu}_{\tilde{f}} \in \mathbb{R}^m$ and diagonal covariance matrix $\tilde{\Sigma}_{\tilde{f}} \in \mathbb{R}^{m \times m}$ are trainable parameters.

For determining the variational parameters, the KL divergence between the posterior probability density and Gaussian trial probability density

$$D_{KL}(q_{\theta, \Psi}(f, \tilde{f}, Z) \| p_\theta(f, \tilde{f}, Z | y)) = \int q_{\theta, \Psi}(f, \tilde{f}, Z) \ln \left(\frac{q_{\theta, \Psi}(f, \tilde{f}, Z)}{p_\theta(y, f, \tilde{f}, Z)} \right) dZ d\tilde{f} df + \ln p_\theta(y), \quad (35)$$

is minimised. By factorising the KL divergence using the joint probability density $p_\theta(y, f, \tilde{f}, Z)$ and trial probability density $q_{\theta, \Psi}(f, \tilde{f}, Z)$ as in (29) and (33) respectively, we obtain the ELBO

$$\mathcal{F}(y) = \mathbb{E}_{q_\psi(Z)}(\hat{\mathcal{F}}(y, Z)) - D_{KL}(q_\omega(\tilde{f}) \| p_\theta(\tilde{f})) - D_{KL}(q_\psi(Z) \| p_W(Z)), \quad (36)$$

where $\hat{\mathcal{F}}(\mathbf{y})$ according to [54], is given by

$$\begin{aligned}\hat{\mathcal{F}}(\mathbf{y}, \mathbf{Z}) &= \mathbb{E}_{q_{\omega}(\tilde{\mathbf{f}})} \left(\mathbb{E}_{p_{\theta}(\mathbf{f}|\tilde{\mathbf{f}}, \mathbf{Z})} \left(\ln p_{\sigma_y}(\mathbf{y}|\mathbf{f}) \right) \right) \\ &= \ln \mathcal{N}(\mathbf{y} | \mathbf{C}_{\mathbf{Z}\mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \tilde{\boldsymbol{\mu}}_{\tilde{\mathbf{f}}}, \sigma_y^2 \mathbf{I}) - \frac{1}{2\sigma_y^2} \left(\text{Tr}(\mathbf{C}_{\mathbf{Z}\mathbf{Z}} - \mathbf{C}_{\mathbf{Z}\mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}) + \text{Tr}(\tilde{\boldsymbol{\Sigma}}_{\tilde{\mathbf{f}}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \mathbf{C}_{\mathbf{Z}\mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1}) \right).\end{aligned}\quad (37)$$

MC sampling of $\mathbf{Z} \sim q_{\psi}(\mathbf{Z})$ is used to estimate the expectation in the ELBO $\mathcal{F}(\mathbf{y})$. The KL divergence terms are derived analytically given that all probability densities belong to the exponential family (Appendix A).

The approximate posterior probability density of the test output variables $\mathbf{f}_*$ evaluated at test latent variables $\mathbf{Z}_* \in \mathbb{R}^{n_* \times d_z}$ given by

$$\begin{aligned}q_{\theta, \omega}(\mathbf{f}_* | \mathbf{Z}_*) &= \int p_{\theta}(\mathbf{f}_* | \tilde{\mathbf{f}}, \mathbf{Z}_*) q_{\omega}(\tilde{\mathbf{f}}) d\tilde{\mathbf{f}} \\ &= \mathcal{N}(\mathbf{C}_{\mathbf{Z}_* \mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \tilde{\boldsymbol{\mu}}_{\tilde{\mathbf{f}}}, \mathbf{C}_{\mathbf{Z}_* \mathbf{Z}_*} - \mathbf{C}_{\mathbf{Z}_* \mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \mathbf{C}_{\mathbf{Z}\mathbf{Z}_*} + \mathbf{C}_{\mathbf{Z}_* \mathbf{Z}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \tilde{\boldsymbol{\Sigma}}_{\tilde{\mathbf{f}}} \mathbf{C}_{\mathbf{Z}\mathbf{Z}}^{-1} \mathbf{C}_{\mathbf{Z}\mathbf{Z}_*}),\end{aligned}\quad (38)$$

is derived by marginalising the conditional probability density $p_{\theta}(\mathbf{f}_* | \tilde{\mathbf{f}}, \mathbf{Z}_*)$ (of similar form to (30)) over the pseudo output variables $\tilde{\mathbf{f}}$ sampled from the trial probability density $q_{\omega}(\tilde{\mathbf{f}})$. By augmenting the training data set $\mathcal{D}$ with pseudo output variables $\tilde{\mathbf{f}}$ the computational complexity is reduced from $O(n^3)$ to $O(nm^2)$ (see also [49–51]), since the GP prior $p_{\theta}(\mathbf{f}|\mathbf{Z}) = \int p_{\theta}(\mathbf{f}|\tilde{\mathbf{f}}, \mathbf{Z}) p_{\theta}(\tilde{\mathbf{f}}) d\tilde{\mathbf{f}}$ now depends on the covariance matrix $\mathbf{C}_{\mathbf{Z}\mathbf{Z}}$ over the m pseudo latent variables $\tilde{\mathbf{Z}}$.

3.3. Inference with input uncertainty

It is computationally more efficient to optimise in RDO the RDVGP surrogate over the low-dimensional test latent variables $\mathbf{Z}_*$ instead of the high-dimensional test input variables $\mathbf{S}_*$. For any test latent variables $\mathbf{Z}_*$ and corresponding test output variable $\mathbf{f}_*$, the approximate marginal posterior probability density given by

$$q_{\theta, \omega, W}(\mathbf{f}_*) = \int q_{\theta, \omega}(\mathbf{f}_* | \mathbf{Z}_*) p_W(\mathbf{Z}_*) d\mathbf{Z}_*, \quad (39)$$

is obtained from the derived approximate posterior probability density $q_{\theta, \omega}(\mathbf{f}_* | \mathbf{Z}_*)$ in (38) and the marginal probability density

$$p_W(\mathbf{Z}_*) = \prod_{i_*=1}^{n_*} \mathcal{N}(\mathbf{W}^T \bar{\mathbf{s}}_{i_*}, \mathbf{W}^T \boldsymbol{\Sigma}_s \mathbf{W}), \quad (40)$$

where the objective of RDO is to compute $\bar{\mathbf{z}}^* = \mathbf{W}^T \bar{\mathbf{s}}^*$, which is mapped to the input domain using the Moore-Penrose inverse $\mathbf{W}^\dagger$. For any test latent variable $\mathbf{z}_*$ (that is any row of $\mathbf{Z}_*$), the mean

$$\mathbb{E}(\mathbf{f}_*) = \mathbb{E}_{p_W(\mathbf{z}_*)} \left(\mathbb{E}_{q_{\theta, \omega}(\mathbf{f}_* | \mathbf{z}_*)}(\mathbf{f}_*) \right), \quad (41)$$

and variance

$$\text{var}(\mathbf{f}_*) = \text{var}_{p_W(\mathbf{z}_*)} \left(\mathbb{E}_{q_{\theta, \omega}(\mathbf{f}_* | \mathbf{z}_*)}(\mathbf{f}_*) \right) + \mathbb{E}_{p_W(\mathbf{z}_*)} \left(\text{var}_{q_{\theta, \omega}(\mathbf{f}_* | \mathbf{z}_*)}(\mathbf{f}_*) \right), \quad (42)$$

of the approximate marginal posterior probability density $q_{\theta, \omega, W}(\mathbf{f}_*)$ given in (39), can be approximated using the laws of total expectation and total variance, with MC sampling.

3.4. Surrogate training

The RDVGP surrogate is trained by learning the set of hyperparameters $\Theta = \theta \cup \{W, \sigma_y\}$ and variational parameters $\Psi = \psi \cup \omega$, that maximise the ELBO $\mathcal{F}(y)$ given by (36), such that

$$\Theta^*, \Psi^* = \arg \max_{\Theta, \Psi} \mathcal{F}(y). \quad (43)$$

We use the stochastic gradient method with the adaptive moment estimation algorithm (ADAM) [55], which evaluates the gradient of the expectation taken with respect to variational parameters of the latent variable trial probability density $q_\psi(\mathbf{Z})$ given by (27). Each ADAM iteration can be repeated n_t times until convergence. We use the default ADAM step size and decay rates [55]. As discussed in [56], to compute the gradient, the latent variable must use the reparameterisation given by

$$\mathbf{z}_i = \tilde{\boldsymbol{\mu}}_{z,i} + \tilde{\mathbf{L}}_{z,i} \boldsymbol{\epsilon}_i, \quad (44)$$

where $\boldsymbol{\epsilon}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and $\tilde{\mathbf{L}}_{z,i}$ is the lower-triangular matrix computed using the Cholesky decomposition of $\tilde{\boldsymbol{\Sigma}}_{z,i}$ (this is known as the reparameterisation trick). The expectation in the ELBO $\mathcal{F}(y)$ is estimated using MC sampling with n_l samples. We compute the ELBO gradient using a single MC sample estimate, which is sufficient for many practical applications [57].

Since the linear projection matrix $\mathbf{W}$ from the statistical observation model (16a) diverges from the Stiefel manifold (is no-longer orthogonal) if using ADAM, the Cayley ADAM algorithm [48] is used (we use the default step size and decay rates). An automatic relevance determination (ARD) procedure [58] is used to determine the dimension d_z of the latent variables $\mathbf{z}$.

Since the ELBO is non-convex, optimisation on the Stiefel manifold using gradient methods may converge to local minima [59], n_r restarts with different initial hyperparameters are used (Algorithm 1). Exploration of highly redundant random input variables $\mathbf{s}$ can be prevented using a sparse orthogonal initialisation of the projection matrix $\mathbf{W}$; however, a random initialisation (using QR decomposition [60]) is more appropriate for input variables $\mathbf{s}$ with low redundancy (we use a combination of both). The covariance matrices of the trial probability densities $\tilde{\boldsymbol{\Sigma}}_{\tilde{f}}$ and $\tilde{\boldsymbol{\Sigma}}_{z,i}$ are sampled from the space of positive definite diagonal matrices.

Training data set $\mathcal{D}$ is proposed using Latin hypercube sampling (LHS) [61]. Input variables $\mathbf{s} = (\mathbf{s}_d^\top \mathbf{s}_f^\top)^\top$ consist of immutable variables sampled from the domain $\mathbf{s}_f \in \{\mathbf{s}_f \in \mathbb{R}^{d_f} \mid \bar{\mathbf{s}}_f - 2\boldsymbol{\sigma}_{s_f} \leq \mathbf{s}_f \leq \bar{\mathbf{s}}_f + 2\boldsymbol{\sigma}_{s_f}\}$, $\bar{\mathbf{s}}_f$ and $\boldsymbol{\sigma}_{s_f}$ are vectors of the mean and standard deviations corresponding to each dimension. Similarly, the design variables are sampled from the domain $\mathbf{s}_d \in \{\mathbf{s}_d \in \mathbb{R}^{d_d} \mid \bar{\mathbf{s}}_d^{(l)} - 2\boldsymbol{\sigma}_{s_d} \leq \mathbf{s}_d \leq \bar{\mathbf{s}}_d^{(u)} + 2\boldsymbol{\sigma}_{s_d}\}$ where $\bar{\mathbf{s}}_d^{(l)}, \bar{\mathbf{s}}_d^{(u)}$

Algorithm 1 Pseudocode for training the RDVGP surrogate

Data: $\mathcal{D} = \{(s_i, y_i) \mid i \in \mathbb{N}, i \leq n\}$	▷ Sampled using LHS
Input: restarts n_r , ADAM iterations n_t , and MC samples n_l	
for $r \in \{1, 2, \dots, n_r\}$ do	▷ Multiple restarts
Initialise model hyperparameters Θ_0, Ψ_0	
for $t \in \{1, 2, \dots, n_t\}$ do	▷ Gradient descent
for $l \in \{1, 2, \dots, n_l\}$ do	▷ MC sampling
$\mathbf{Z}_l \sim q_\psi(\mathbf{Z})$	▷ Reparameterisation (44)
Compute $\mathcal{F}_t(y)$ and its gradient using $\{\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_{n_l}\}$, Θ_{t-1} , and Ψ_{t-1}	▷ ELBO (36)
$\Theta_t \setminus \{W_t\}, \Psi_t \leftarrow \text{ADAM}(\mathcal{F}_t(y))$	▷ Stochastic gradient method
$W_t \leftarrow \text{CAYLEYADAM}(\mathcal{F}_t(y))$	▷ Stiefel manifold optimisation
Result: optimum model hyperparameters Θ^*, Ψ^*	

and σ_{s_d} are vectors of lower limits, upper limits, and standard deviations corresponding to each dimension.

Slice sampling is used to infer the test output variables f_* corresponding to the design variables $\bar{s}_{d*}$. The approximate posterior probability density $q_{\theta,\omega}(f_*|\mathbf{z}_*)$ given by (38), is evaluated at the projected mean input variables $\bar{\mathbf{z}}_* = \mathbf{W}^\top \bar{\mathbf{s}}_*$ where $\bar{\mathbf{s}}_* = (\bar{s}_{d*}^\top \bar{s}_{f*}^\top)^\top$ and $\bar{s}_{d*} \in \{\bar{s}_{d*} \in \mathbb{R}^{d_d} \mid \bar{s}_d^{(l)} \leq \bar{s}_{d*} \leq \bar{s}_d^{(u)}\}$ (Figure 5).

3.5. Surrogate comparison

We verify the predictive accuracy of the RDVGP surrogate using two performance metrics; the coefficient of determination (COD) [62], and the maximum mean discrepancy (MMD) [63]. For any validation mean input variables $\bar{\mathbf{S}}_v$, corresponding true observations $\mathbf{y}_v$ of the objective or constraint function are obtained by sampling the input variable density $p(\mathbf{S}_v)$. Summary statistics of the true observations (such as mean and variance) can be compared to the inferred target output variables f_v . The COD values for the mean and variance

$$R_\mu^2 = 1 - \frac{\sum_{v=1}^{n_v} |\mathbb{E}(\mathbf{y}_v) - \mathbb{E}(f_v)|^2}{\sum_{v=1}^{n_v} |\mathbb{E}(\mathbf{y}_v) - \mu_{y_v}|^2}, \quad (45a)$$

$$R_\sigma^2 = 1 - \frac{\sum_{v=1}^{n_v} |\text{var}(\mathbf{y}_v) - \text{var}(f_v)|^2}{\sum_{v=1}^{n_v} |\text{var}(\mathbf{y}_v) - \sigma_{y_v}^2|^2}, \quad (45b)$$

of the marginal posterior probability density $q_{\theta,\omega,W}(f_v)$ in (39) are computed, where $\mu_{y_v} = \frac{1}{N_v} \sum_{v=1}^{n_v} \mathbb{E}(\mathbf{y}_v)$ and $\sigma_{y_v}^2 = \frac{1}{N_v} \sum_{v=1}^{n_v} \text{var}(\mathbf{y}_v)$.

While the COD provides a measure of the quality of fit over the entire model, the distance between probability densities at points of interest such as local minima is important for determining the reliability of the predicted robust optimum design variables. The MMD $D_{MMD} \in \mathbb{R}^+$ provides this measure in terms of

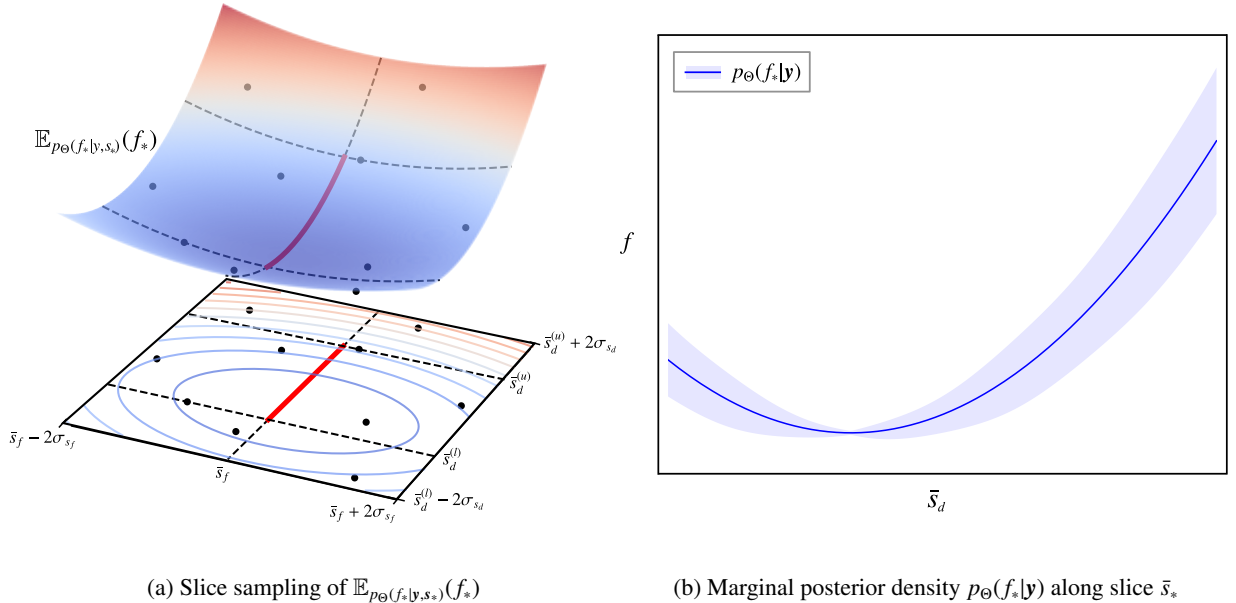


Figure 5: Schematic of slice sampling showing (a) the posterior probability density expected value $\mathbb{E}_{p_{\Theta}(f_*, \mathbf{y}, \mathbf{s}_*)}(f_*)$ along slice $\mathbf{s}_* = (s_{d*} \ s_{f*})^\top$, and (b) the marginal posterior probability density $p_{\Theta}(f_* | \mathbf{y})$ over mean test input variable $\bar{\mathbf{s}}_* = (\bar{s}_{d*} \ \bar{s}_{f*})^\top$.

a mean embedding in reproducing kernel Hilbert space with associated kernel $k : \mathbb{R}^{d_y} \times \mathbb{R}^{d_y} \rightarrow \mathbb{R}$. For two probability densities $p(\mathbf{w})$ and $q(\mathbf{v})$, the MMD is

$$D_{MMD} = \mathbb{E}_{p(\mathbf{w})} (k(\mathbf{w}, \mathbf{w}')) - 2 \mathbb{E}_{p(\mathbf{w}), q(\mathbf{v})} (k(\mathbf{w}, \mathbf{v})) + \mathbb{E}_{q(\mathbf{v})} (k(\mathbf{v}, \mathbf{v}')). \quad (46)$$

In all examples in this paper, we use a Gaussian kernel with unit amplitude and length scale to compute the MMD.

4. Examples

We demonstrate the accuracy of the proposed RDVGP surrogate using four examples of increasing complexity. In all examples, we use the sparse RDVGP surrogate, but refer to it as the RDVGP surrogate for brevity. Illustrative examples are used to compare the RDVGP surrogate to standard GP regression and examines its versatility to different levels of input variance. We also demonstrate the RDVGP surrogate on RDO examples.

The robust minimum as predicted by the RDVGP surrogate, is determined with genetic algorithms (NSGA-II), which use the emulated solutions inferred by the surrogate. To verify its accuracy, we directly MC sample the input variable and solve the corresponding FE model to estimate the true posterior probability density over the domain of the design variables. This is also optimised using genetic algorithms to yield an estimate of the true solution, for comparison.

4.1. One-dimensional illustrative example

The approximation of marginal posterior probability densities using the RDVGP surrogate is compared to the conventional approach of MC sampling from a standard GP surrogate as in (14) [64]. For both surrogates, the inferred variance depends on the random input variables $\mathbf{s}$, allowing variation over the input variable domain $\mathbf{s} \in D_s$. Consider the non-linear multi-modal objective function

$$J(s) = -0.5s \sin(3\pi s^2) + 0.25s, \quad (47)$$

where $s \sim \mathcal{N}(\bar{s}, \sigma^2)$, $\bar{s} \in \{\bar{s} \in \mathbb{R} \mid 0 \leq \bar{s} \leq 1\}$, and $\sigma = 0.025$. The training data set $\mathcal{D}$ is initialised using LHS with $n = 11$ random samples (as shown in Figures 6a and 6c) and the same number of inducing points (with $m = 11$). This example considers only a single input variable, hence no dimensionality reduction is required. The standard GP and RDVGP surrogates are trained by maximising the log marginal likelihood $\ln p_{\Theta}(\mathbf{y})$ from (10) and ELBO $\mathcal{F}(\mathbf{y})$ given by (36) respectively.

The standard GP overfits the training data (Figure 6a) due to the small length scale (hyperparameter of the covariance function), resulting in a posterior probability density $p(\mathbf{J}_* | \mathbf{y}, \mathbf{s}_*)$ with large epistemic uncertainty between sampled input variables. As mentioned in the introduction, epistemic uncertainty is caused by a lack of information and can in principle be reduced with additional training data. Conversely, aleatoric uncertainty is caused by the random input variable $\mathbf{s}$ itself [20]. When using MC sampling to approximate the true marginal posterior probability density $p(\mathbf{J})$ using the standard GP posterior probability density $p(\mathbf{J}_* | \mathbf{y}, \mathbf{s}_*)$ given by (14), the aleatoric and epistemic uncertainties are confounded and the total uncertainty is significantly overestimated (Figure 6b).

In contrast, the RDVGP surrogate does not overfit the training data (Figure 6c); due to the inclusion of the prescribed input variable probability density $p(s)$ in the ELBO $\mathcal{F}(\mathbf{y})$, adding a data-informed regularisation. The expectation in the ELBO $\mathcal{F}(\mathbf{y})$ formulation promotes smoothness near the sampled training data, since each MC sample $\mathbf{Z} \sim q_{\Psi}(\mathbf{Z})$ (following from (27)) is verified for a closeness of fit between the

true observations $\mathbf{y}$ and the inferred target output variables $\mathbf{f}$ using a likelihood formulation $\hat{\mathcal{F}}(\mathbf{y}, \mathbf{Z})$ given by (37), which promotes length scale hyperparameters that increase smoothness. This adds regularisation to the model and the use of pseudo target output variables $\tilde{\mathbf{f}}$ adds further flexibility when fitting the model (to prevent over-fitting). In comparison to the standard GP surrogate, this reduces epistemic uncertainty and more accurately estimates the total uncertainty, for improved selection of robust optimum design variables $\bar{s}^*$ (Figure 6d). In this example, the standard GP surrogate would require more training data than the RDVGP surrogate to reduce overfitting before accurately emulating the true marginal posterior probability density $p(\mathbf{J})$.

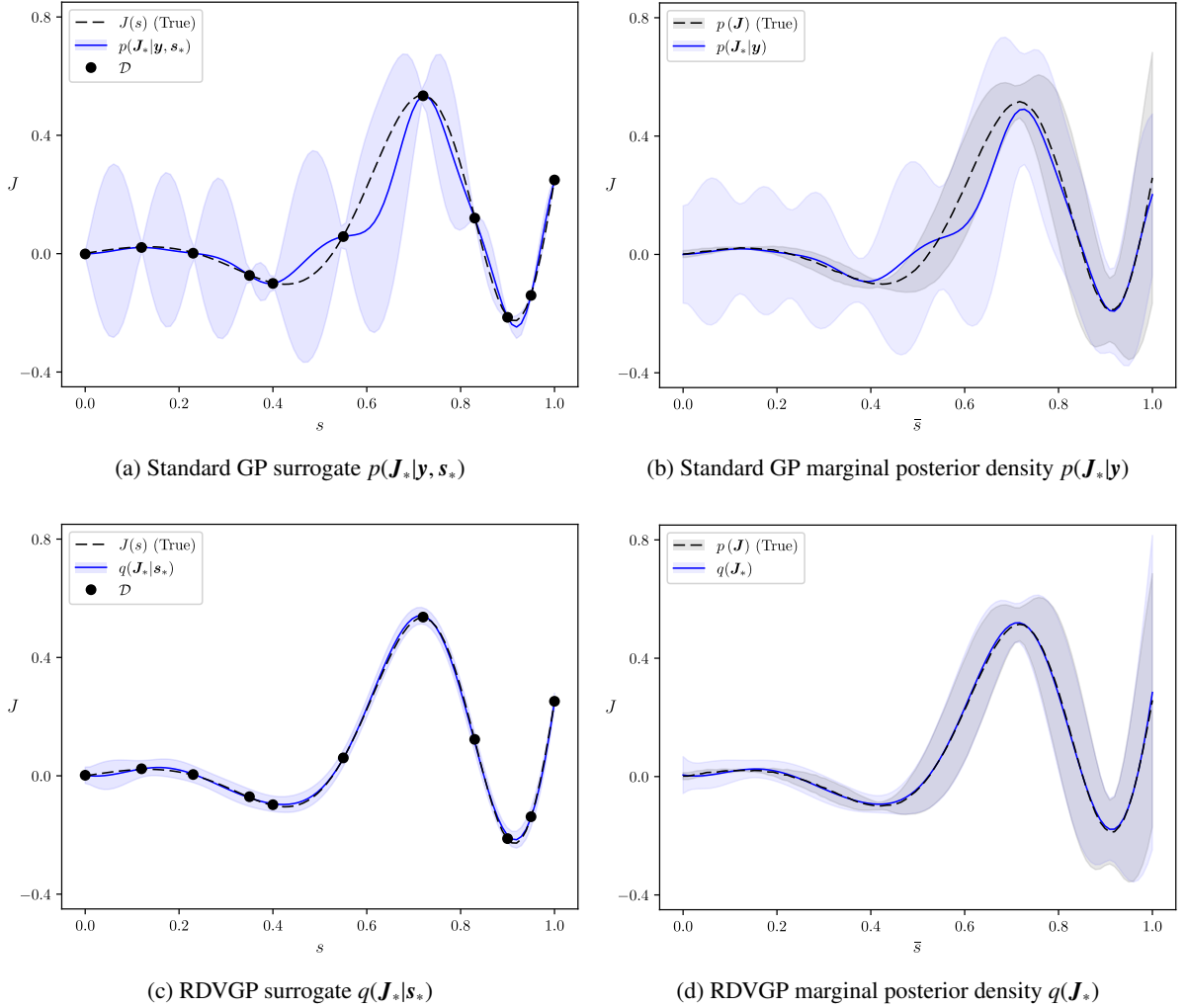


Figure 6: One-dimensional illustrative example. Comparing the (a) standard GP surrogate $p(\mathbf{J}_*|\mathbf{y}, s_*)$ given by (14) and its (b) marginal posterior probability density $p(\mathbf{J}_*|\mathbf{y}) = \mathbb{E}_{p(s_*)}(p(\mathbf{J}_*|\mathbf{y}, s_*))$, with the (c) RDVGP surrogate $q(\mathbf{J}_*|s_*)$ given by (38), and (d) RDVGP marginal posterior probability density $q(\mathbf{J}_*) = \mathbb{E}_{p(s_*)}(q(\mathbf{J}_*|s_*))$ from (39). The training data set $\mathcal{D}$ consists of $n = 11$ random input variable s and observation y pairs sampled from the true objective function $J(s)$, and $p(\mathbf{J})$ is the true marginal posterior probability density.

4.2. Three-dimensional illustrative example

This example investigates the accuracy and robustness of the RDVGP surrogate to different levels of input variable uncertainty and choices of parameters used to train the surrogate. Consider the minimisation of a non-linear multi-modal objective function

$$J(\mathbf{s}) = (s_2 s_1 - 2)^2 \sin(12s_1 - 4) + 8s_1 + s_3, \quad (48)$$

of three dimensional input variables $\mathbf{s} = (s_1 \ s_2 \ s_3)^\top$, subject to the constraint function

$$H(\mathbf{s}) = -\cos(2\pi s_1) - s_3 - 0.7, \quad (49)$$

where $s_1 \sim \mathcal{N}(\bar{s}_1, \sigma_1^2)$ is a design variable, and $s_2 \sim \mathcal{N}(6.0, \sigma_2^2)$ and $s_3 \sim \mathcal{N}(0, 0.1^2)$ are immutable variables. The mean of the design variable is sampled from the domain $\bar{s}_1 \in \{\bar{s}_1 \in \mathbb{R} \mid 0 \leq \bar{s}_1 \leq 1\}$. We perform RDO at three different choices for $\sigma_1 \in \{0.025, 0.05, 0.075\}$ and $\sigma_2 \in \{0.25, 0.5, 0.75\}$. According to the RDO problem stated in (1), a weighting factor of $\alpha = 0.25$ and feasibility index of $\beta = 2$ is chosen. No limit is imposed on the variance of the constraint function $H(\mathbf{s})$.

The training data set $\mathcal{D}$ is initialised using LHS with $n = 100$ samples and choosing $m = 25$ inducing points (reducing time complexity by a factor of 16 from $\mathcal{O}(n^3)$ to $\mathcal{O}(nm^2)$). The robust minimum variable $\bar{s}_1^*$ computed using the RDVGP surrogate is compared to the true minimum estimated using MC sampling $\langle \bar{s}_1^* \rangle$ with 10^4 samples. A total of $n_v = 30$ validation samples are used to compute COD values (given by (45)) of the inferred objective and constraint marginal posterior probability densities $q(\mathbf{J}_*)$ and $q(\mathbf{H}_*)$ respectively.

The immutable variable s_3 is substantially less influential than s_1 and s_2 , giving a subspace dimension of $d_z = 2$. An accurate robust minimum variable $\bar{s}_1^*$ is computed for all standard deviation combinations σ_1 and σ_2 , when compared to the true solution (Table 1). The largest computed percentage difference between the RDVGP $\bar{s}_1^*$ and true $\langle \bar{s}_1^* \rangle$ robust optimum design variables is 4%. Although the objective function $J(\mathbf{s})$ has two local minima, one is not robust since the objective function varies significantly in close proximity to that local minima (Figure 7). The approximate marginal posterior probability density $q(\mathbf{J}_*)$ computed using the RDVGP surrogate accurately emulates the true marginal posterior probability density $p(\mathbf{J})$, therefore the robust minimum $\bar{s}_1^*$ is well estimated. If the uncertainty in the random input variables becomes negligible, the global minimum would be $\bar{s}_1 = 0.749$, which does not coincide with the robust optimum design variable for any of the sampled statistics (Table 1).

Large COD values R_μ^2 are obtained from (45a) for the mean of both the objective $J(\mathbf{s})$ and constraint $H(\mathbf{s})$ functions for all combinations of σ_1 and σ_2 , indicating the mean is well predicted. However, the COD value R_σ^2 obtained from (45b) for the variance of the objective and constraint functions varies for different combinations of σ_1 and σ_2 , and is most affected by the choice of σ_1 . Furthermore, the MMD D_{MMD} , which compares the true $p(J(\langle \bar{s}_1^* \rangle))$ and approximate (RDVGP) $q(J(\langle \bar{s}_1^* \rangle))$ marginal posterior probability densities of the objective function at the true robust optimum design variable $\langle \bar{s}_1^* \rangle$ as given by (46), is affected by both σ_1 and σ_2 . At smaller values for σ_1 , the epistemic uncertainty makes up a greater proportion of the total uncertainty predicted by the RDVGP surrogate, increasing its discrepancy with the true solution. Since a slice of the full surrogate is taken at $s_2 = 6.0$, it becomes more difficult to approximate the aleatoric uncertainty at larger values of σ_2 .

As previously mentioned, the epistemic uncertainty can be reduced by increasing the size of training data set $\mathcal{D}$ (Figure 8). By including supplementary data sampled using LHS centred on the true robust optimum design variable $\langle \bar{s}_1^* \rangle$ and box bounded within two standard deviations of each random input variable, the approximate (RDVGP) marginal posterior probability density $q(J(\langle \bar{s}_1^* \rangle))$ converges towards the true marginal posterior probability density $p(J(\langle \bar{s}_1^* \rangle))$ with an increasing number of supplementary samples n_s .

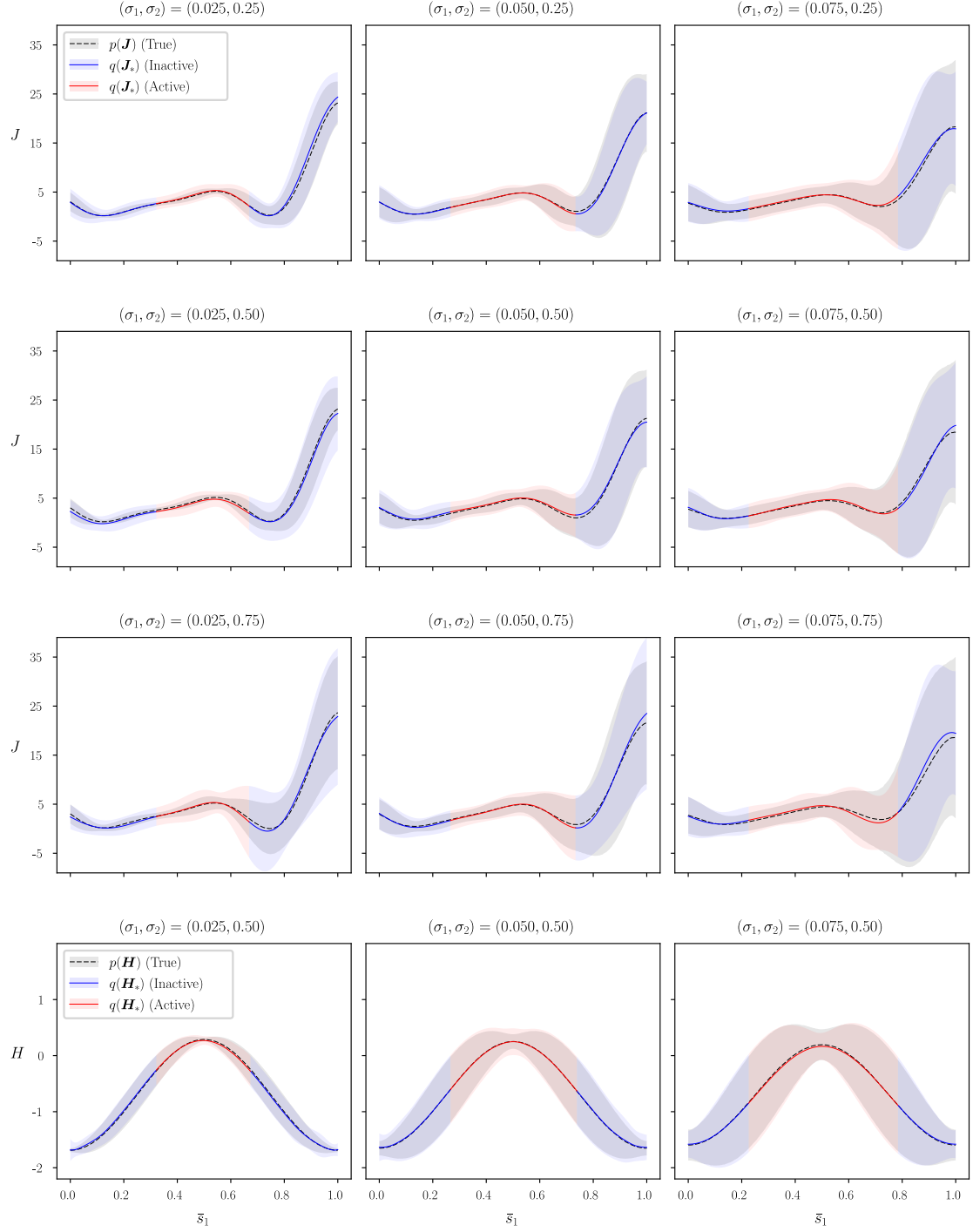
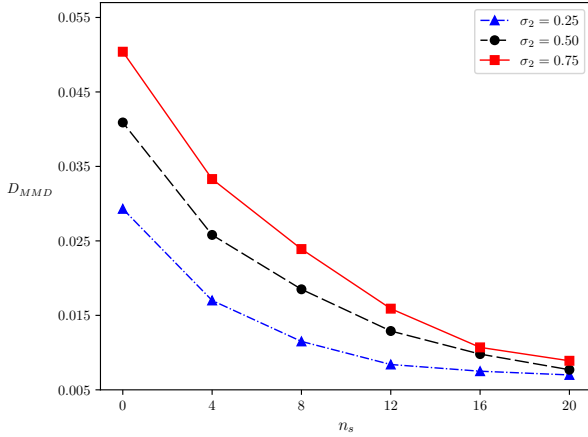


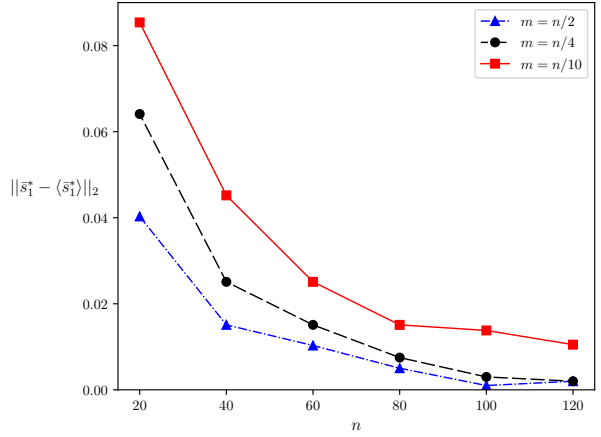
Figure 7: Three-dimensional illustrative example. Comparison of the true $p(\cdot)$ and approximate (RDVGP) $q(\cdot)$ marginal posterior probability densities of the objective and constraint functions $J(s)$ and $H(s)$ over the domain of the design variable s_1 . The design variable has density $\bar{s}_1 \sim \mathcal{N}(\bar{s}_1, \sigma_1^2)$ and the immutable variables have densities $\bar{s}_2 \sim \mathcal{N}(\bar{s}_2, \sigma_2^2)$ and $\bar{s}_3 \sim \mathcal{N}(\bar{s}_3, 0.1^2)$. The constraint is active only in the red region.

Table 1: Three-dimensional illustrative example. Robust minimum design variable as predicted by the RDVGP surrogate $\bar{s}_1^*$ and MC sampling $\langle \bar{s}_1^* \rangle$, with COD performance metrics computed using (45) for the objective $\left((R_\mu^J)^2, (R_\sigma^J)^2\right)$ and constraint $\left((R_\mu^H)^2, (R_\sigma^H)^2\right)$ functions, and MMD D_{MMD} (computed using (46)) between the approximate (RDVGP) $q(J(\langle \bar{s}_1^* \rangle))$ and true $p(J(\langle \bar{s}_1^* \rangle))$ marginal posterior probability densities evaluated at $\langle \bar{s}_1^* \rangle$.

σ_1	σ_2	$(R_\mu^J)^2$	$(R_\sigma^J)^2$	$(R_\mu^H)^2$	$(R_\sigma^H)^2$	D_{MMD}	$\bar{s}_1^*$	$\langle \bar{s}_1^* \rangle$
0.025	0.25	0.997	0.811	0.999	0.790	0.0293	0.128	0.123
	0.50	0.993	0.828	0.999	0.789	0.0409	0.127	0.124
	0.75	0.994	0.828	0.999	0.793	0.0504	0.127	0.124
0.050	0.25	0.991	0.923	0.998	0.883	0.0175	0.137	0.138
	0.50	0.985	0.924	0.999	0.888	0.0263	0.137	0.135
	0.75	0.988	0.929	0.998	0.887	0.0364	0.134	0.136
0.075	0.25	0.992	0.959	0.995	0.949	0.0039	0.152	0.148
	0.50	0.974	0.965	0.995	0.959	0.0076	0.151	0.155
	0.75	0.969	0.944	0.994	0.932	0.0108	0.151	0.150



(a) MMD D_{MMD} convergence



(b) Optimum solution $\bar{s}_1^*$ convergence

Figure 8: Three-dimensional illustrative example. Convergence plots showing the (a) effect of the number of supplementary training data points n_s sampled within the vicinity of the true robust optimum design variable $\langle \bar{s}_1^* \rangle$, on the MMD D_{MMD} computed between the true marginal posterior probability density $p(J(\langle \bar{s}_1^* \rangle))$ and the approximate (RDVGP) marginal posterior probability density $q(J(\langle \bar{s}_1^* \rangle))$ at $\sigma_1 = 0.025$ and varying values of σ_2 . In (b) the convergence of the predicted solution $\bar{s}_1^*$ to the true robust optimum design variable $\langle \bar{s}_1^* \rangle$ for different training data sample sizes n and choices of number of inducing points m , for $(\sigma_1, \sigma_2) = (0.050, 0.50)$ is shown.

This is demonstrated by the MMD D_{MMD} decreasing for an increasing number of supplementary samples, a trend seen across all values of σ_2 . Consequently, sequential or adaptive sampling strategies [65] would be effective for reducing the epistemic uncertainty when using the RDVGP surrogate.

4.3. Plate with elliptical hole

We study the use of the RDVGP surrogate on the RDO of a Poisson problem on a plate with an elliptical hole (Figure 9). The boundary value problem is given by

$$-\kappa \nabla^2 u(\mathbf{x}, \mathbf{s}) = b(\mathbf{x}, \mathbf{s}), \quad \mathbf{x} \in \Omega, \quad (50a)$$

$$u(\mathbf{x}, \mathbf{s}) = 0, \quad \mathbf{x} \in \partial\Omega_D, \quad (50b)$$

$$\nabla u(\mathbf{x}, \mathbf{s}) \cdot \mathbf{n} = 0, \quad \mathbf{x} \in \partial\Omega_N, \quad (50c)$$

where $\kappa \in \mathbb{R}^+$ is the diffusion coefficient, $b(\mathbf{x}, \mathbf{s}) \in \mathbb{R}$ is the source, $u(\mathbf{x}, \mathbf{s}) \in \mathbb{R}$ is the solution, $\partial\Omega_D$ is the Dirichlet boundary, $\partial\Omega_N$ is the Neumann boundary, and $\mathbf{n}$ is the normal to the Neumann boundary $\partial\Omega_N$. The domain Ω is discretised over spatial coordinates $\mathbf{x} = (x_1 \ x_2)^\top$, and consists of a unit square $\Omega_1 = \{x_1, x_2 \in \mathbb{R} \mid -0.5 \leq x_1 \leq 0.5, -0.5 \leq x_2 \leq 0.5\}$ with an elliptical void centred on the origin $\Omega_2 = \{x_1, x_2 \in \mathbb{R} \mid (x_1/s_1)^2 + (x_2/s_2)^2 < 1\}$, where $\mathbf{s}_d = (s_1 \ s_2)^\top$ are design variables and $\Omega = \Omega_1 \setminus \Omega_2$. The random source $b(\mathbf{x}, \mathbf{s})$ is a function of the fifth-order Chebyshev polynomial $T_5 : \mathbb{R} \rightarrow \mathbb{R}$, given by

$$b(\mathbf{x}, \mathbf{s}) = s_3 (T_5(x_1 - s_4))^2 + s_5 (T_5(x_2 - s_6))^2 + s_7, \quad (51)$$

where $\mathbf{s}_f = (s_3 \ s_4 \ \dots \ s_7)^\top$ are immutable variables and $\mathbf{s} = (\mathbf{s}_d^\top \ \mathbf{s}_f^\top)^\top$. The diffusion coefficient $\kappa = 1$ is constant throughout the domain Ω . All the random input variables s_i are Gaussian, where $s_i \sim \mathcal{N}(\bar{s}_i, \sigma_i^2)$ and $\bar{s}_i \in \mathbb{R}$ is a fixed scalar for $i \in \{3, 4, \dots, 7\}$, and box-constrained $\bar{s}_i \in \{\bar{s}_i \in \mathbb{R} \mid 0 \leq \bar{s}_i \leq 0.4\}$ for $i \in \{1, 2\}$. The goal is to compute the robust optimum design with $\bar{\mathbf{s}}_d^* = (\bar{s}_1^* \ \bar{s}_2^*)^\top$ that minimises the energy

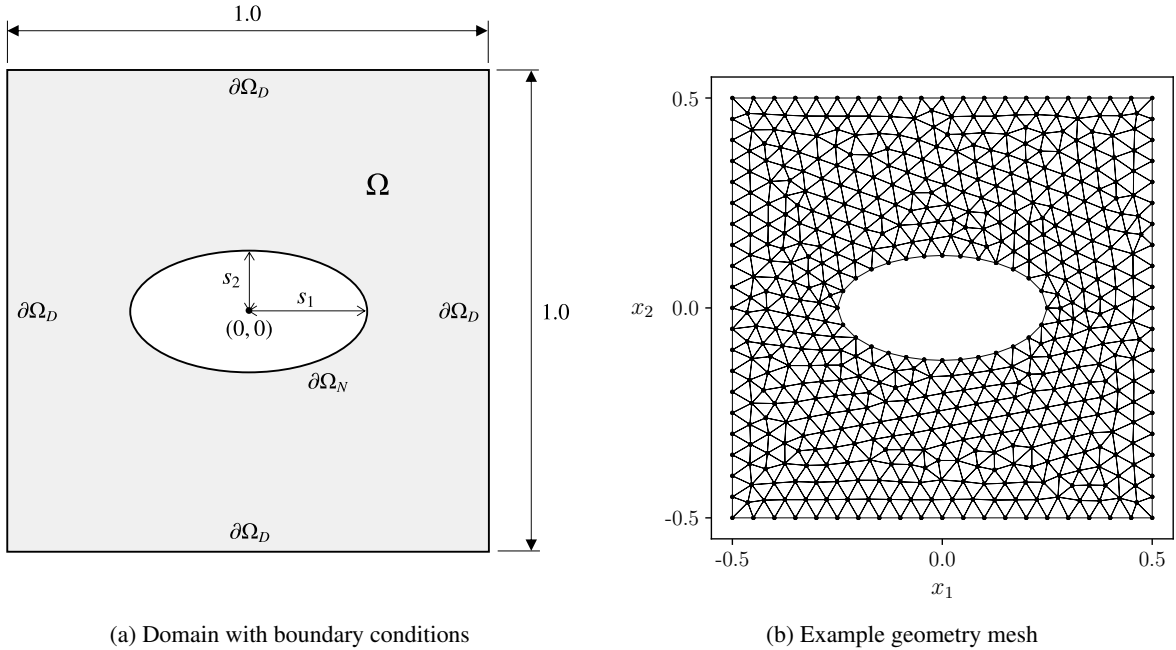


Figure 9: Plate with elliptical hole. Schematic showing (a) the specification of the domain Ω with the elliptical void parameterised by design variables s_1 and s_2 centred at the origin with boundaries $\partial\Omega_D$ and $\partial\Omega_N$ where the Dirichlet and Neumann boundary conditions are applied respectively, and (b) mesh of an example geometry containing 982 linear triangle elements.

$$J(s) = \frac{1}{2} \kappa \int_{\Omega} \|\nabla u(\mathbf{x}, s)\|_2^2 d\mathbf{x}, \quad (52)$$

subject to the area constraint

$$H(s) = A_0 - \int_{\Omega} d\mathbf{x}, \quad (53)$$

where $A_0 = 0.8$ is the limit on the domain area. The constraint function $H(s)$ is introduced to eliminate trivial solutions with a zero hole size. According to the RDO formulation given by (1), a weighting factor of $\alpha = 0.5$ and a feasibility index of $\beta = 2$ is used. No condition is applied to the variance of the constraint function $H(s)$.

The training data set $\mathcal{D}$ is initialised using LHS with $n = 200$ samples and $m = 100$ inducing points. The objective $J(s)$ and constraint $H(s)$ functions are sampled from a deterministic FE model (Figure 9). To test the reliability of the RDVGP surrogate, five source terms $b(\mathbf{x}, s)$ are considered, each with different immutable variables s_f (Table 2). Design variables s_1 and s_2 have standard deviations of $\sigma_1 = 0.010$ and $\sigma_2 = 0.025$ respectively. RDO normalisation constants $\bar{\mu}$ and $\bar{\sigma}$ are equal to the mean and standard deviation of the objective function $J(s)$ observations. RDVGP robust optimum design variables $\bar{s}_d^*$ are compared to the true solution $\langle \bar{s}_d^* \rangle$ evaluated using MC sampling with 10^5 evaluations of the FE model. A total of $n_v = 30$ validation samples $\mathcal{D}_v$ are used for computing COD values (given by (45)) of the objective and constraint functions.

The computed robust optimum design variables $\bar{s}_d^*$ maintain similarity to the true solution $\langle \bar{s}_d^* \rangle$ for all source terms, with the most and least similarity corresponding to the third $b^{(3)}(\mathbf{x}, s)$ and fourth $b^{(4)}(\mathbf{x}, s)$ source terms respectively (Table 3). In both cases, the source term varies uniaxially and is aligned perpendicular and parallel to the direction of greatest variance in the shape of the elliptical hole respectively (since σ_2 is greater than σ_1) (Figure 10). This effect propagates through to the solution $u(\mathbf{x}, s)$ (Figure 11).

If the random input variables s are treated as deterministic, the optimum design variables $\hat{s}_d^*$ would differ significantly from the robust optimum design variables $\bar{s}_d^*$ computed using the RDVGP surrogate (Table 3). For each source term, a significant difference between the deterministic $\hat{s}_d^*$ and robust $\bar{s}_d^*$ optimum design variables is observed. Since this example considers the uncertainty as equally important as the mean (since $\alpha = 0.5$), variance in the objective $J(s)$ and constraint $H(s)$ functions cannot be ignored. In the case of all source terms $b(\mathbf{x}, s)$, the location of the minimum mean and variance in the objective function $J(s)$ do not coincide, although the difference is more prominent with the fourth $b^{(4)}(\mathbf{x}, s)$ and fifth $b^{(5)}(\mathbf{x}, s)$ source terms (Figure 12). Even though the variance in the second source term $b^{(2)}(\mathbf{x}, s)$ is small, a relatively large variance in the objective $J(s)$ is obtained, since the variance in the design variables s_1 and s_2 contribute most to variance in the objective (as supported by ARD).

Large COD values are indicative of similarity between the true $p(\mathbf{J})$ and approximate (RDVGP) $q(\mathbf{J}_*)$ marginal posterior probability densities. In comparison to the COD values for the mean R_{μ}^2 , smaller COD values R_{σ}^2 are observed for the variance of both the objective $J(s)$ and constraint $H(s)$ functions. This is due to the presence of epistemic uncertainty in the approximate (RDVGP) marginal posterior probability density $q(\mathbf{J}_*)$, which increases the total uncertainty when compared to the true density $p(\mathbf{J})$. However, COD values greater than 0.7 indicate the RDVGP surrogate emulates the true objective $J(s)$ and constraint $H(s)$ functions well. Larger COD values for the standard deviation are observed where variance in the source $b(\mathbf{x}, s)$ is low. The relatively low MMD D_{MMD} values (as computed using (46)) show the similarity between the true $p(J(\langle \bar{s}_d^* \rangle))$ and approximate (RDVGP) $q(J(\langle \bar{s}_d^* \rangle))$ marginal posterior probability densities.

Next, we investigate the effect of the weighting factor α prescribed for the RDO problem (1). Depending on the choice of weighting factor, the robust optimum design variables $\bar{s}_d^*$ predicted from the RDVGP

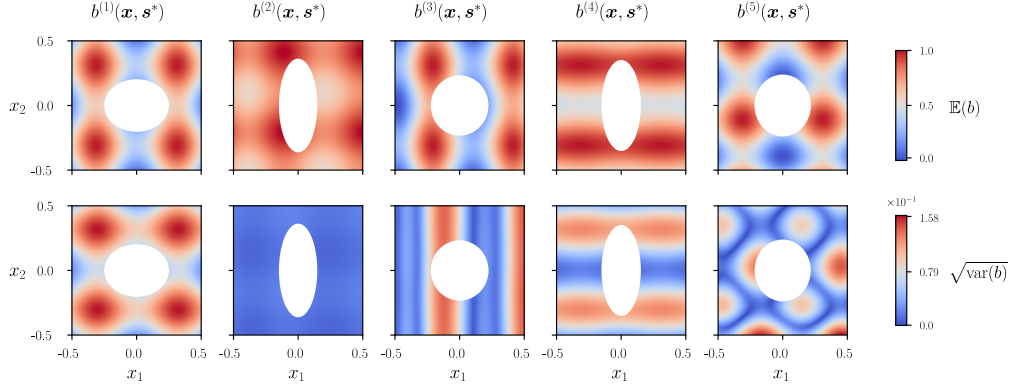


Figure 10: Plate with elliptical hole. Optimised geometries for different source terms $b^{(j)}(\mathbf{x}, \mathbf{s}^*)$ (with $j \in \{1, 2, \dots, 5\}$), and the respective empirical expectation $\mathbb{E}(b(\mathbf{x}, \mathbf{s}^*))$ and standard deviation $\sqrt{\text{var}(b(\mathbf{x}, \mathbf{s}^*))}$.

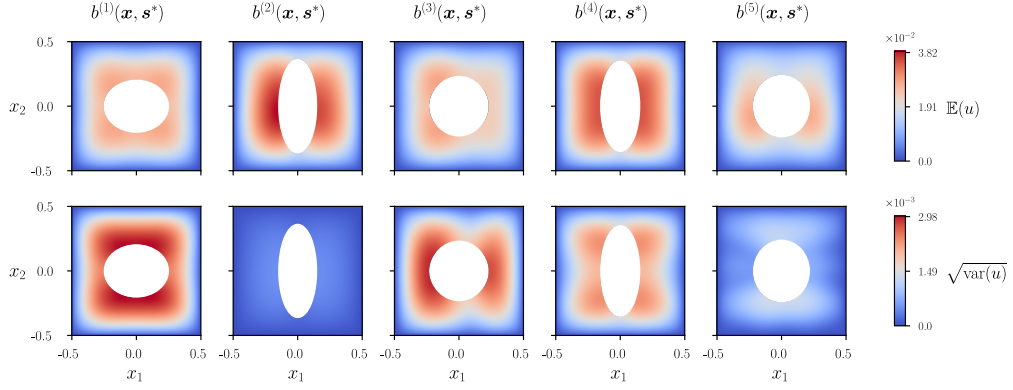


Figure 11: Plate with elliptical hole. Optimised geometries with solution fields $u(\mathbf{x}, \mathbf{s}^*)$ for each source term $b^{(j)}(\mathbf{x}, \mathbf{s})$ (with $j \in \{1, 2, \dots, 5\}$), and the respective empirical expectation $\mathbb{E}(u(\mathbf{x}, \mathbf{s}^*))$ and standard deviation $\sqrt{\text{var}(u(\mathbf{x}, \mathbf{s}^*))}$.

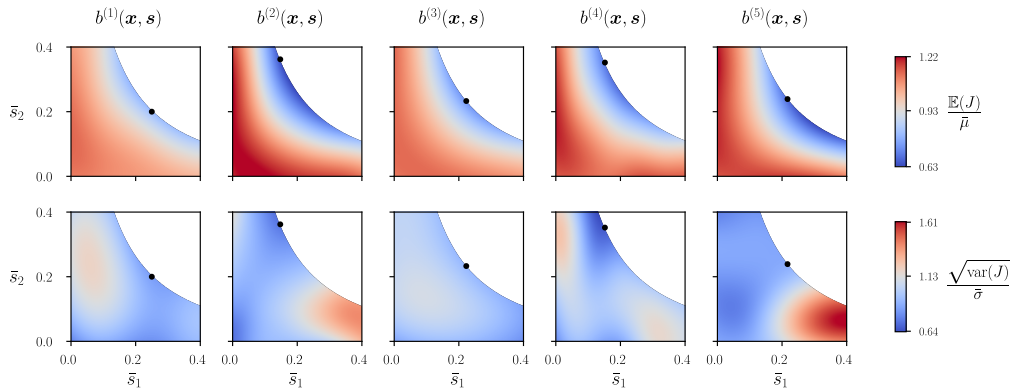


Figure 12: Plate with elliptical hole. Marginal posterior probability density $q(J_*)$ showing the empirical normalised expectation $\mathbb{E}(J(\bar{s}_d))/\bar{\mu}$ and standard deviation $\sqrt{\text{var}(J(\bar{s}_d))}/\bar{\sigma}$ for each of the source terms $b^{(j)}(\mathbf{x}, \mathbf{s})$ (with $j \in \{1, 2, \dots, 5\}$), showing the approximate (RDVGP) robust optimum design variables $\bar{s}_d^*$ (black dot).

Table 2: Plate with elliptical hole. Prescribed mean $\bar{s}_i$ and standard deviation σ_i of the random immutable variables s_i (with $i \in \{3, 4, \dots, 7\}$) of the j^{th} source term $b^{(j)}(\mathbf{x}, s)$ given by (51).

$b^{(j)}(\mathbf{x}, s)$	$\bar{s}_3$	σ_3	$\bar{s}_4$	σ_4	$\bar{s}_5$	σ_5	$\bar{s}_6$	σ_6	$\bar{s}_7$	σ_7
$b^{(1)}(\mathbf{x}, s)$	0.55	0.060	0.00	0.001	0.35	0.080	0.00	0.001	0.10	0.010
$b^{(2)}(\mathbf{x}, s)$	0.20	0.005	0.20	0.001	0.20	0.005	0.10	0.001	0.60	0.010
$b^{(3)}(\mathbf{x}, s)$	0.65	0.100	0.10	0.020	0.25	0.001	0.00	0.001	0.10	0.001
$b^{(4)}(\mathbf{x}, s)$	0.05	0.010	0.00	0.001	0.45	0.100	0.00	0.001	0.50	0.010
$b^{(5)}(\mathbf{x}, s)$	0.45	0.010	0.00	0.025	0.45	0.010	0.20	0.025	0.10	0.010

Table 3: Plate with elliptical hole. Robust optimum design variables obtained using the RDVGP surrogate $(\bar{s}_d^*)^T = (\bar{s}_1^* \bar{s}_2^*)$ and MC sampling $\langle \bar{s}_d^* \rangle^T = (\langle \bar{s}_1^* \rangle, \langle \bar{s}_2^* \rangle)$, compared to the deterministic prediction $(\hat{s}_d^*)^T = (\hat{s}_1^* \hat{s}_2^*)$, with COD performance metrics (45) for the objective $J(s) \left((R_\mu^J)^2, (R_\sigma^J)^2 \right)$ and constraint $H(s) \left((R_\mu^H)^2, (R_\sigma^H)^2 \right)$ functions, and MMD D_{MMD} between the true $p(J(\langle \bar{s}_d^* \rangle))$ and approximate (RDVGP) $q(J(\langle \bar{s}_d^* \rangle))$ marginal posterior probability densities.

$b^{(j)}(\mathbf{x}, s)$	d_z	$(R_\mu^J)^2$	$(R_\sigma^J)^2$	$(R_\mu^H)^2$	$(R_\sigma^H)^2$	D_{MMD}	$(\hat{s}_d^*)^T$	$(\bar{s}_d^*)^T$	$\langle \bar{s}_d^* \rangle^T$
$b^{(1)}(\mathbf{x}, s)$	4	0.957	0.821	0.997	0.924	0.000307	(0.326 0.196)	(0.250 0.203)	(0.270 0.183)
$b^{(2)}(\mathbf{x}, s)$	3	0.996	0.925	0.998	0.941	0.000434	(0.223 0.286)	(0.148 0.362)	(0.159 0.336)
$b^{(3)}(\mathbf{x}, s)$	4	0.959	0.801	0.995	0.915	0.000345	(0.269 0.237)	(0.223 0.233)	(0.228 0.227)
$b^{(4)}(\mathbf{x}, s)$	3	0.967	0.885	0.987	0.931	0.000314	(0.203 0.316)	(0.152 0.352)	(0.153 0.350)
$b^{(5)}(\mathbf{x}, s)$	4	0.997	0.744	0.992	0.902	0.000074	(0.397 0.161)	(0.217 0.239)	(0.196 0.268)

surrogate may not converge to the true solution $\langle \bar{s}_d^* \rangle$. Robust optimum design variables $\bar{s}_d^*$ predicted using large weighting factors α are unreliable (Figure 13), since the weighted sum being minimised is predominantly influenced by the variance, exacerbating the effect of the large epistemic uncertainty. The epistemic uncertainty is large when the sample size n in the training data set $\mathcal{D}$ is small, with confounding between the aleatoric and epistemic uncertainty causing discrepancy between the true $p(\mathbf{J})$ and approximate (RDVGP) $q(\mathbf{J}_*)$ marginal posterior probability densities. As expected, increasing the size n of the training data set $\mathcal{D}$ reduces the epistemic uncertainty and increases the reliability of the approximate (RDVGP) robust optimum design variables $\bar{s}_d^*$ (Figure 13). This emphasises the importance of using a training data set $\mathcal{D}$ with a sufficiently large sample size n when using surrogates to emulate the objective $J(s)$ and constraint $H^{(j)}(s)$ functions in RDO. Overestimating the uncertainty may result in a sub-optimal choice of design variables for complex engineering systems.

4.4. Bracket

Next, we demonstrate the use of the RDVGP surrogate on the RDO of a bracket. The three-dimensional equations of elasticity are given by

$$-\nabla \cdot \boldsymbol{\sigma}_s(\mathbf{u}(\mathbf{x}, s)) = \mathbf{b}(\mathbf{x}, s), \quad \mathbf{x} \in \Omega, \quad (54a)$$

$$\mathbf{u}(\mathbf{x}, s) = \mathbf{0}, \quad \mathbf{x} \in \partial\Omega_D, \quad (54b)$$

$$\boldsymbol{\sigma}_s(\mathbf{u}(\mathbf{x}, s)) \cdot \mathbf{n} = \mathbf{g}(\mathbf{x}, s), \quad \mathbf{x} \in \partial\Omega_N, \quad (54c)$$

where $\boldsymbol{\sigma}_s(\mathbf{u}(\mathbf{x}, s)) \in \mathbb{R}^{3 \times 3}$ is the stress tensor, $\mathbf{b}(\mathbf{x}, s) \in \mathbb{R}^3$ is the body force vector, $\mathbf{u}(\mathbf{x}, s) \in \mathbb{R}^3$ is the displacement vector, and $\mathbf{g}(\mathbf{x}, s) \in \mathbb{R}^3$ is the applied surface traction at the Neumann boundary $\partial\Omega_N$. For plane stress deformation, the stress tensor $\boldsymbol{\sigma}_s(\mathbf{u}(\mathbf{x}, s))$ consists of out-of-plane components $\sigma_{13} = \sigma_{23} = \sigma_{33} = 0$,

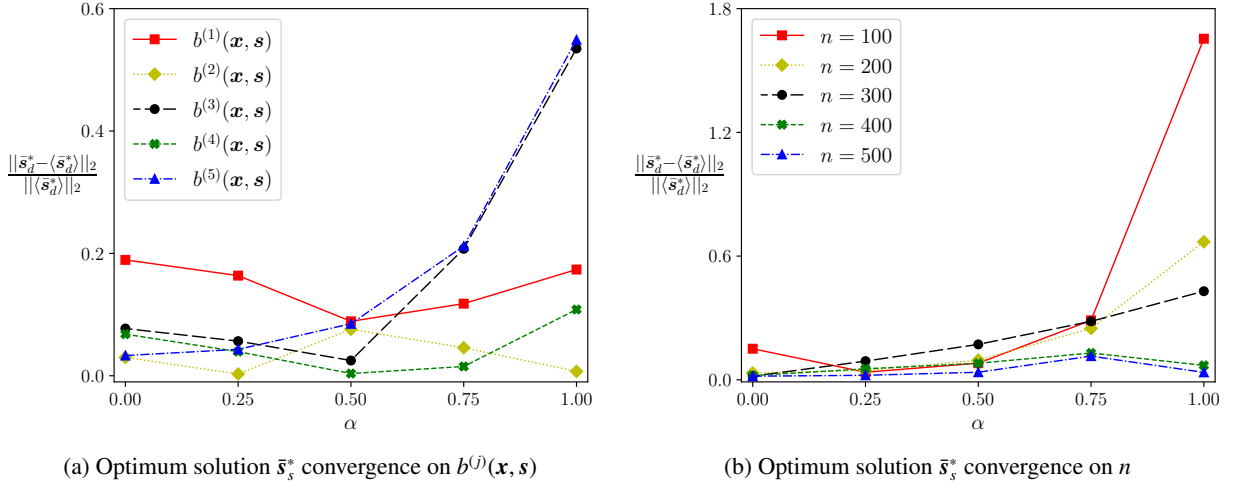


Figure 13: Plate with elliptical hole. Effect of the (a) weighting factor α for different source terms $b^{(j)}(x, s)$ (with $j \in \{1, 2, \dots, 5\}$) on the convergence (as measured by the relative error) between the approximate (RDVGP) $\bar{s}_d^*$ and true $\langle \bar{s}_d^* \rangle$ robust optimum design variables. The effect of (b) the weighting factor α and size n of the training data set (assuming $m = n/2$) for the fifth source term $b^{(5)}(x, s)$ on convergence.

and in-plane components given by the constitutive relation (Hooke's law)

$$\sigma_{ij} = \frac{E}{1+\nu} \left(\varepsilon_{ij} + \frac{\nu}{1-\nu} \text{Tr}(\boldsymbol{\varepsilon}(\mathbf{u}(\mathbf{x}, s))) \delta_{ij} \right), \quad i, j \in \{1, 2\}, \quad (55)$$

with the strain tensor given by

$$\boldsymbol{\varepsilon}(\mathbf{u}(\mathbf{x}, s)) = \frac{1}{2} \left(\nabla \mathbf{u}(\mathbf{x}, s) + (\nabla \mathbf{u}(\mathbf{x}, s))^T + (\nabla \mathbf{u}(\mathbf{x}, s))^T \nabla \mathbf{u}(\mathbf{x}, s) \right), \quad (56)$$

where $E, \nu \in \mathbb{R}^+$ are the Young's modulus and Poisson's ratio respectively, δ_{ij} is the Kronecker delta function, and $\boldsymbol{\varepsilon}(\mathbf{u}(\mathbf{x}, s)) \in \mathbb{R}^{3 \times 3}$ is the strain tensor with components ε_{ij} (with $i, j \in \{1, 2, 3\}$). As shown in Figure 14, the bracket has an outer boundary $\partial\Omega_o$, a polygonal void with boundary $\partial\Omega_v$, and a circular hole with boundary $\partial\Omega_c$. The Dirichlet boundary condition is applied along the left edge $\partial\Omega_D$, with the Neumann boundary condition applied along the remaining boundary $\partial\Omega_N$ where $\partial\Omega_N = \partial\Omega_o \cup \partial\Omega_v \cup \partial\Omega_c$.

The domain Ω is discretised over spatial coordinates $\mathbf{x} = (x_1 \ x_2 \ x_3)^T$, and is parameterised by five design variables $\mathbf{s}_d = (s_1 \ s_2 \ \dots \ s_5)^T$ and ten immutable variables $\mathbf{s}_f = (s_6 \ s_7 \ \dots \ s_{15})^T$ controlling the shape of the domain Ω , and the magnitude s_{14} and angle s_{15} of the resultant applied load. The body force $\mathbf{b}(\mathbf{x}, s)$ is constant over the geometry domain, where $\mathbf{b}(\mathbf{x}, s) = (0 \ -7.7 \times 10^{-5} \ 0)^T$. The surface traction $\mathbf{g}(\mathbf{x}, s)$ yields the resultant applied load (applied as uniformly distributed) along the circular hole boundary $\partial\Omega_c$ where $\int_{\partial\Omega_c} \mathbf{g}(\mathbf{x}, s) d\mathbf{x} = (s_{14} \cos(s_{15}) \ s_{14} \sin(s_{15}) \ 0)^T$. All random input variables $s_i \in \mathbf{s}$ are Gaussian, where $s_i \sim \mathcal{N}(\bar{s}_i, \sigma_i^2)$ and $\bar{s}_i \in \mathbb{R}$ is a fixed constant for $i \in \{6, 7, \dots, 15\}$, and box-constrained on the interval $\bar{s}_i \in \{\bar{s}_i \in \mathbb{R} \mid \bar{s}_i^{(l)} \leq \bar{s}_i \leq \bar{s}_i^{(u)}\}$ for $i \in \{1, 2, \dots, 5\}$ where $\bar{s}_i^{(l)}$ and $\bar{s}_i^{(u)}$ are the lower and upper limits respectively.

The goal is to compute the robust optimum design variables $\bar{\mathbf{s}}_d^* = (\bar{s}_1^* \ \bar{s}_2^* \ \dots \ \bar{s}_5^*)^T$ (1) that minimise structural compliance

$$J(\mathbf{s}) = \int_{\Omega} \boldsymbol{\sigma}_s(\mathbf{u}(\mathbf{x}, s)) : \boldsymbol{\varepsilon}(\mathbf{u}(\mathbf{x}, s)) d\mathbf{x}, \quad (57)$$

subject to the volume constraint function

$$H^{(1)}(\mathbf{s}) = \int_{\Omega} d\mathbf{x} - V_0, \quad (58)$$

where $V_0 = 1.25 \times 10^5$ is the chosen limit on the volume. A constraint is also applied on the stress (von Mises), where

$$H^{(2)}(\mathbf{s}) = \max_{\mathbf{x} \in \Omega} \sigma_v(\mathbf{u}(\mathbf{x}, \mathbf{s})) - \sigma_{v0}, \quad (59)$$

$\sigma_v(\mathbf{u}(\mathbf{x}, \mathbf{s}))$ is the von Mises stress, and $\sigma_{v0} = 125$ is the upper limit on the stress. According to the RDO formulation (1), a weighting factor of $\alpha = 0.5$ and feasibility indices of $\beta^{(1)} = \beta^{(2)} = 2$ are used. The standard deviation in the maximum von Mises stress $\sigma_v(\mathbf{u}(\mathbf{x}, \mathbf{s}))$ must also be less than 35.

The training data set $\mathcal{D}$ is initialised using LHS with $n = 500$ samples and $m = 250$ inducing points. The objective $J(\mathbf{s})$ and stress constraint $H^{(2)}(\mathbf{s})$ are sampled from a deterministic FE model (Figure 15). To avoid sampling stress singularities, the maximum von Mises stress $\max \sigma_v(\mathbf{u}(\mathbf{x}, \mathbf{s}))$ is not sampled within a distance of 5 from the Dirichlet $\partial\Omega_D$ and Neumann $\partial\Omega_N$ boundary condition surfaces. The bracket is modelled as a thin plate, using the plane stress assumption. A Young's modulus of $E = 2.1 \times 10^5$ and Poisson's ratio of $\nu = 0.3$ are used. The mean and standard deviation of the immutable variables $\mathbf{s}_f$ are prescribed. Similarly, the lower and upper bounds on the mean, and standard deviation of the design variables $\mathbf{s}_d$ are also prescribed (Table 4). Normalisation constants $\bar{\mu}$ and $\bar{\sigma}$ in the RDO weighted sum (1) are set equal to the mean and standard deviation of the objective function $J(\mathbf{s})$ observations.

The robust optimum design variables $\bar{\mathbf{s}}_d^*$ computed using the RDVGP surrogate are compared to the true

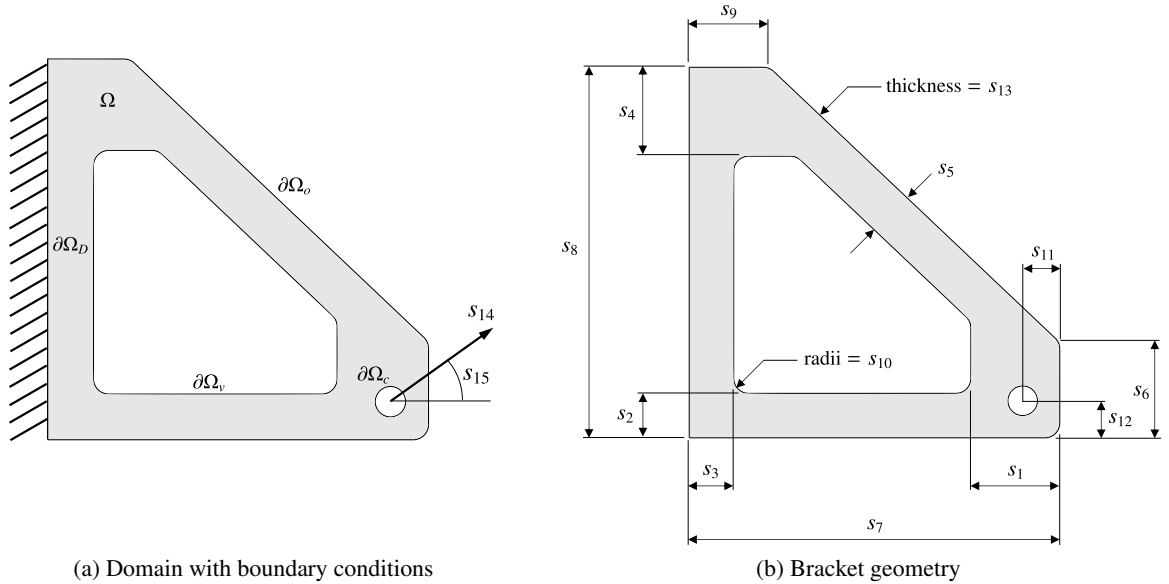


Figure 14: Bracket. Schematic showing (a) the domain of the bracket with a Dirichlet boundary condition applied on the left edge $\partial\Omega_D$, the Neumann boundary condition applied along the union of the outer boundary $\partial\Omega_o$, void boundary $\partial\Omega_v$, and circular hole boundary $\partial\Omega_c$ and (b) bracket geometry with random design variables $\mathbf{s}_d = (s_1 \ s_2 \ \dots \ s_5)^T$ and random immutable variables $\mathbf{s}_f = (s_6 \ s_7 \ \dots \ s_{15})^T$.

optimum $\langle \bar{s}_d^* \rangle$ evaluated using MC sampling with 10^5 samples. A total of $n_v = 30$ validation samples are used for computing COD values given by (45).

The immutable variables s_f with negligible variance when propagated through the objective $J(s)$ and constraint $H^{(j)}(s)$ functions, contribute insignificantly to variance in the observations. This is a reasonable representation of many complex engineering systems manufactured with high-precision equipment. As expected, the subspace is primarily composed of a linear combination of the design variables s_d and the magnitude s_{14} and angle s_{15} of the applied load, where the projection matrix $\mathbf{W}$ has dimension $d_z = 7$.

The true robust optimum design variables $\langle \bar{s}_d^* \rangle = (88.6 \ 41.1 \ 10.5 \ 69.6 \ 29.2)^T$ (as estimated with MC sampling) are well approximated by the robust optimum design variables $\bar{s}_d^* = (84.8 \ 40.2 \ 10.8 \ 75.2 \ 28.5)^T$ proposed by performing RDO (1) with the RDVGP surrogate. Due to the uncertainty in the magnitude s_{14} and angle s_{15} of the applied load (Figure 14), the marginal posterior probability densities for the compliance $q(\mathbf{J}_*)$ and the stress constraint $q(\mathbf{H}_*^{(2)})$ yield a large aleatoric uncertainty prediction (Figure 16). Consequently, permissible COD values are obtained for the mean R_μ^2 (0.801, 0.907) and variance R_σ^2 (0.757, 0.716) of the objective $J(s)$ and stress constraint $H^{(2)}(s)$ functions respectively. Since the aleatoric uncertainty in the volume constraint $H^{(1)}(s)$ is negligible due to the small standard deviation in the random input variables s (Table 4), the epistemic uncertainty constitutes most of the total variance, and the RDVGP surrogate be-

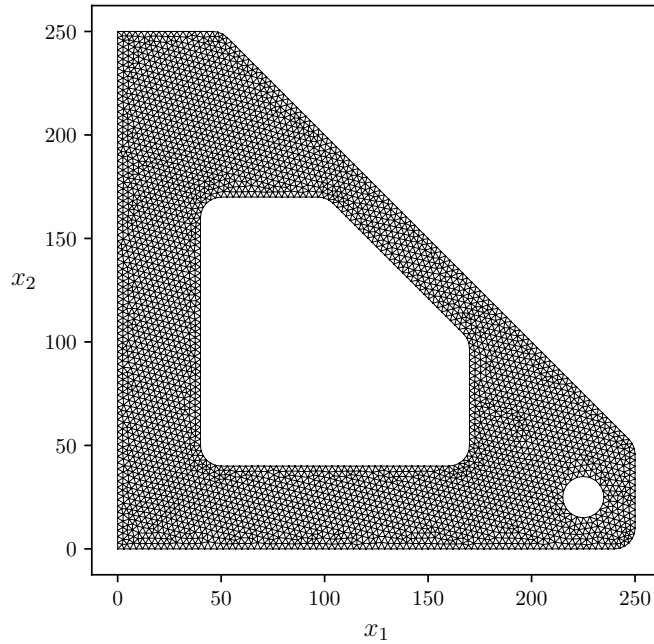


Figure 15: Bracket. FE mesh of an example geometry containing 7668 linear triangle elements, discretised over spatial coordinates $\mathbf{x} = (x_1 \ x_2 \ x_3)^T$ according to the plane stress assumption.

Table 4: Bracket. Mean $\bar{s}_i$ and standard deviation σ_i of the random input variables $\mathbf{s} = (s_1 \ s_2 \ \dots \ s_{15})^T$, where box constraints $\bar{s}_i \in [\bar{s}_i^{(l)}, \bar{s}_i^{(u)}]$ are given for the design variables $\mathbf{s}_d = (s_1 \ s_2 \ \dots \ s_5)^T$.

	s_1	s_2	s_3	s_4	s_5	s_6	s_7	s_8	s_9	s_{10}	s_{11}	s_{12}	s_{13}	s_{14}	s_{15}
$\bar{s}_i$	[60,100]	[10,50]	[10,50]	[60,100]	[0,30]	50	250	250	50	10	25	25	5	5000	0
σ_i	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	1000	$\pi/6$

comes a poor predictor of the aleatoric uncertainty in the volume constraint $H^{(1)}(\mathbf{s})$. This is demonstrated by COD values of 0.864 and 0.331 for the mean and variance respectively. The probability density near the robust optimum design variables is accurate, as demonstrated by an MMD D_{MMD} (computed using (46)) of 0.000983 between the true $p(J(\langle \bar{\mathbf{s}}_d^* \rangle))$ and approximate (RDVGP) $q(J(\langle \bar{\mathbf{s}}_d^* \rangle))$ marginal posterior probability densities. This example is a good demonstration of the RDVGP surrogate accuracy in uncertainty quantification with application to RDO.

5. Conclusion

We introduced the RDVGP surrogate, which presents a statistical surrogate that can be used to emulate the quantities of interest in complex engineering systems, for application in solving RDO problems. Due to the availability of only a finite number of deterministic PDE observations that may depend on high-dimensional random input variables, we generalised a standard GP by postulating a graphical model with intrinsic assumptions relating the random input variables and observations to latent variables sampled from a low-dimensional subspace. Both the surrogate hyperparameters and the linear projection matrix describ-

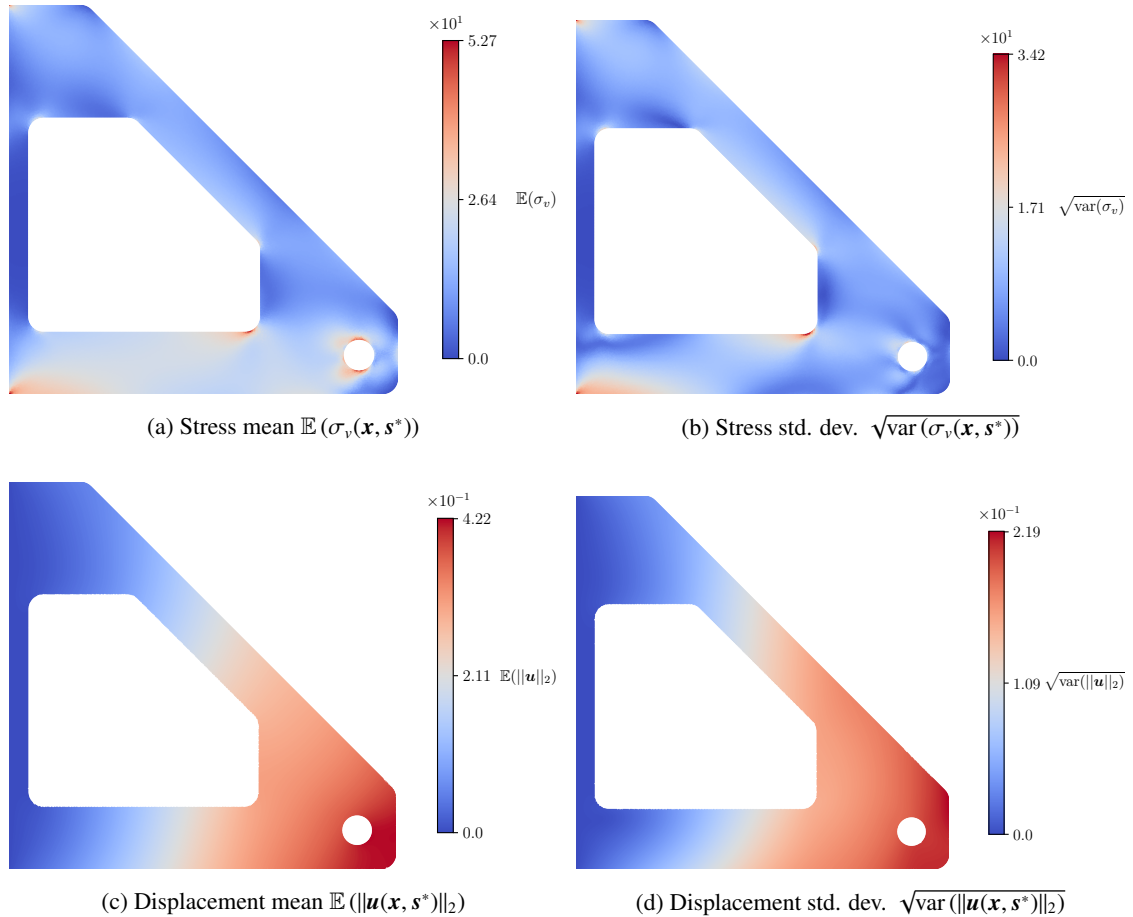


Figure 16: Bracket. Deformation showing for the optimum geometry $\mathbf{s}^* = ((\bar{\mathbf{s}}_d^*)^\top \mathbf{s}_f^\top)^\top$ (a) the von Mises stress mean $\mathbb{E}(\sigma_v(\mathbf{x}, \mathbf{s}^*))$, (b) the von Mises stress standard deviation $\sqrt{\text{var}(\sigma_v(\mathbf{x}, \mathbf{s}^*))}$, (c) the displacement mean $\mathbb{E}(\|\mathbf{u}(\mathbf{x}, \mathbf{s}^*)\|_2)$, and (d) the displacement standard deviation $\sqrt{\text{var}(\|\mathbf{u}(\mathbf{x}, \mathbf{s}^*)\|_2)}$.

ing the subspace are learned jointly by optimising a single log marginal likelihood formulation, which is computed using VB. The computational simplicity of the RDVGP surrogate provides a favourable alternative to other algorithmic approaches, which do not consider input variable uncertainty when training a standard GP, leading to over-fitting.

Compared to the conventional approach of fitting a standard GP and computing the marginal posterior probability density by MC sampling over the prescribed input variable probability density, the RDVGP surrogate is more accurate while remaining efficient, since the induced sparsity reduces the computational expense of sampling over the random input variables. The approximate marginal posterior probability density predicted using the RDVGP surrogate sufficiently emulates the true marginal posterior probability density, yielding accurate robust optimum design variables in RDO applications.

In closing, we mention promising research extensions to the proposed approach. The statistical surrogate model is implemented using a GP; however, since our approach uses a graphical model, this can be further generalised to consider alternative surrogates such as Bayesian physics-informed deep neural networks [66, 67], extending to grey-box applications [68]. The linear subspace could also be discovered using Riemannian manifolds that retain knowledge of intrinsic geometric constraints, for use in geometry-aware Bayesian optimisation [69]. The efficiency, accuracy, and scalability of the proposed model makes it promising for use in the RDO of complex engineering systems, such as construction-aware optimisation of concrete-based structures [70]. Although we assumed a linear dimensionality reduction in the statistical observation model, non-linear dimensionality reduction [71] may further improve efficiency and accuracy for some types of problems. The use of the RDVGP surrogate when used with adaptive sampling algorithms such as Bayesian optimisation, is also a promising direction for further research.

Appendix A : Derivation of the KL divergence terms

The KL divergence terms in the ELBO $\mathcal{F}(\mathbf{y})$ given by (36) are analytically tractable and can be computed in exact form. For example, consider the second KL divergence term

$$\begin{aligned}
D_{KL}(q_\psi(\mathbf{Z}) \| p_W(\mathbf{Z})) &= \int q_\psi(\mathbf{Z}) \ln \left(\frac{q_\psi(\mathbf{Z})}{p_W(\mathbf{Z})} \right) d\mathbf{Z} \\
&= \sum_{i=1}^n \mathbb{E}_{q_\psi(z_i)} (\ln q_\psi(z_i) - \ln p_W(z_i)) \\
&= \frac{1}{2} \sum_{i=1}^n \mathbb{E}_{q_\psi(z_i)} \left(\ln \left(\frac{|\mathbf{W}^T \Sigma_s \mathbf{W}|}{|\tilde{\Sigma}_{z,i}|} \right) + (z_i - \mathbf{W}^T \tilde{s}_i)^T (\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} (z_i - \mathbf{W}^T \tilde{s}_i) - (z_i - \tilde{\mu}_{z,i})^T (\tilde{\Sigma}_{z,i})^{-1} (z_i - \tilde{\mu}_{z,i}) \right) \\
&= \frac{1}{2} \sum_{i=1}^n \left(\ln \left(\frac{|\mathbf{W}^T \Sigma_s \mathbf{W}|}{|\tilde{\Sigma}_{z,i}|} \right) + \mathbb{E}_{q_\psi(z_i)} \left(\text{Tr} \left((\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} (z_i - \mathbf{W}^T \tilde{s}_i) (z_i - \mathbf{W}^T \tilde{s}_i)^T \right) - \text{Tr} \left((\tilde{\Sigma}_{z,i})^{-1} (z_i - \tilde{\mu}_{z,i}) (z_i - \tilde{\mu}_{z,i})^T \right) \right) \right) \quad (\text{A.1}) \\
&= \frac{1}{2} \sum_{i=1}^n \left(\ln \left(\frac{|\mathbf{W}^T \Sigma_s \mathbf{W}|}{|\tilde{\Sigma}_{z,i}|} \right) + \mathbb{E}_{q_\psi(z_i)} \left(\text{Tr} \left((\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} (z_i z_i^T - 2z_i \tilde{s}_i^T \mathbf{W} + \mathbf{W}^T \tilde{s}_i \tilde{s}_i^T \mathbf{W}) \right) - \text{Tr} \left((\tilde{\Sigma}_{z,i})^{-1} \tilde{\Sigma}_{z,i} \right) \right) \right) \\
&= \frac{1}{2} \sum_{i=1}^n \left(\ln \left(\frac{|\mathbf{W}^T \Sigma_s \mathbf{W}|}{|\tilde{\Sigma}_{z,i}|} \right) - d_z + \text{Tr} \left((\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} (\tilde{\Sigma}_{z,i} + \tilde{\mu}_{z,i} (\tilde{\mu}_{z,i})^T - 2\mathbf{W}^T \tilde{s}_i (\tilde{\mu}_{z,i})^T + \mathbf{W}^T \tilde{s}_i \tilde{s}_i^T \mathbf{W}) \right) \right) \\
&= \frac{1}{2} \sum_{i=1}^n \left(\ln \left(\frac{|\mathbf{W}^T \Sigma_s \mathbf{W}|}{|\tilde{\Sigma}_{z,i}|} \right) - d_z + \text{Tr} \left((\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} \tilde{\Sigma}_{z,i} \right) + (\mathbf{W}^T \tilde{s}_i - \tilde{\mu}_{z,i})^T (\mathbf{W}^T \Sigma_s \mathbf{W})^{-1} (\mathbf{W}^T \tilde{s}_i - \tilde{\mu}_{z,i}) \right),
\end{aligned}$$

and after following a similar derivation for the first KL divergence term

$$D_{KL}(q_\omega(\tilde{\mathbf{f}}) \| p_\theta(\tilde{\mathbf{f}})) = \frac{1}{2} \left(\ln \left(\frac{|C_{\tilde{Z}\tilde{Z}}|}{|\tilde{\Sigma}_{\tilde{f}}|} \right) - m + \text{Tr} (C_{\tilde{Z}\tilde{Z}}^{-1} \tilde{\Sigma}_{\tilde{f}}) + (\tilde{\mu}_{\tilde{f}})^T C_{\tilde{Z}\tilde{Z}}^{-1} \tilde{\mu}_{\tilde{f}} \right). \quad (\text{A.2})$$

Appendix B: Multiple observation vectors

The d_y observation vectors and corresponding target output variables can be collected into the matrices $\mathbf{Y} \in \mathbb{R}^{n \times d_y}$ and $\mathbf{F} \in \mathbb{R}^{n \times d_y}$ respectively, where $\mathbf{Y} = (\mathbf{y}^{(1)} \mathbf{y}^{(2)} \dots \mathbf{y}^{(d_y)})$ and $\mathbf{F} = (\mathbf{f}^{(1)} \mathbf{f}^{(2)} \dots \mathbf{f}^{(d_y)})$. The likelihood

$$p_{\sigma_y}(\mathbf{Y}|\mathbf{F}) = \prod_{i=1}^{d_y} p_{\sigma_y}(\mathbf{y}^{(i)}|\mathbf{f}^{(i)}) = \prod_{i=1}^{d_y} \mathcal{N}(\mathbf{f}^{(i)}, \sigma_y^{(i)} \mathbf{I}), \quad (\text{B.1})$$

and GP prior probability density

$$p_{\theta}(\mathbf{F}|\mathbf{Z}) = \prod_{i=1}^{d_y} p_{\theta}(\mathbf{f}^{(i)}|\mathbf{Z}) = \prod_{i=1}^{d_y} \mathcal{N}(\mathbf{0}, \mathbf{C}_{ZZ}^{(i)}), \quad (\text{B.2})$$

can be expanded using the mean field approximation. Similarly, the ELBO $\mathcal{F}(\mathbf{y})$ in (28) for the statistical latent variable model may be expanded as the summation of the ELBO over each observation vector

$$\mathcal{F}(\mathbf{y}) = \sum_{i=1}^{d_y} \mathbb{E}_{q_{\psi}(\mathbf{Z})} \left(\mathbb{E}_{p_{\theta}(\mathbf{f}^{(i)}|\mathbf{Z})} (\ln p_{\sigma_y}(\mathbf{y}^{(i)}|\mathbf{f}^{(i)})) \right) - D_{KL}(q_{\psi}(\mathbf{Z}) \| p_W(\mathbf{Z})). \quad (\text{B.3})$$

Similarly for the sparse statistical latent variable model, the sparse GP prior

$$p_{\theta}(\mathbf{F}|\tilde{\mathbf{F}}, \mathbf{Z}) = \prod_{i=1}^{d_y} p_{\theta}(\mathbf{f}^{(i)}|\tilde{\mathbf{f}}^{(i)}, \mathbf{Z}) = \prod_{i=1}^{d_y} \mathcal{N}(\mathbf{C}_{ZZ}^{(i)} (\mathbf{C}_{\tilde{Z}\tilde{Z}}^{(i)})^{-1} \tilde{\mathbf{f}}^{(i)}, \mathbf{C}_{ZZ}^{(i)} - \mathbf{C}_{ZZ}^{(i)} (\mathbf{C}_{\tilde{Z}\tilde{Z}}^{(i)})^{-1} \mathbf{C}_{\tilde{Z}\tilde{Z}}^{(i)}), \quad (\text{B.4})$$

and trial probability density

$$q_{\omega}(\tilde{\mathbf{F}}) := \prod_{i=1}^{d_y} q_{\omega}(\tilde{\mathbf{f}}^{(i)}) := \prod_{i=1}^{d_y} \mathcal{N}(\tilde{\boldsymbol{\mu}}_{\tilde{\mathbf{f}}}^{(i)}, \tilde{\boldsymbol{\Sigma}}_{\tilde{\mathbf{f}}}^{(i)}), \quad (\text{B.5})$$

can be expanded, where $\tilde{\mathbf{F}} \in \mathbb{R}^{m \times d_y}$ is the matrix of pseudo target output variables.

References

- [1] D. Xiu, G. E. Karniadakis, Modeling uncertainty in steady state diffusion problems via generalized polynomial chaos, *Comput. Methods Appl. Mech. Engrg.* 191 (43) (2002) 4927–4948. [doi:10.1016/S0045-7825\(02\)00421-8](https://doi.org/10.1016/S0045-7825(02)00421-8).
- [2] C. Soize, R. Ghanem, Physical systems with random uncertainties: chaos representations with arbitrary probability measure, *SIAM J. Sci. Comput.* 26 (2) (2004) 395–410. [doi:10.1137/S1064827503424505](https://doi.org/10.1137/S1064827503424505).
- [3] R. Zhang, H. Dai, Stochastic analysis of structures under limited observations using kernel density estimation and arbitrary polynomial chaos expansion, *Comput. Methods Appl. Mech. Engrg.* 403 (2023) 115689. [doi:10.1016/j.cma.2022.115689](https://doi.org/10.1016/j.cma.2022.115689).
- [4] K. D. Kantarakias, G. Papadakis, Sensitivity-enhanced generalized polynomial chaos for efficient uncertainty quantification, *J. Comput. Phys.* 491 (2023) 112377. [doi:10.1016/j.jcp.2023.112377](https://doi.org/10.1016/j.jcp.2023.112377).
- [5] M. A. Nielsen, *Neural networks and deep learning*, Determination Press, 2015.
- [6] M. Papadrakakis, N. D. Lagaros, Reliability-based structural optimization using neural networks and Monte Carlo simulation, *Comput. Methods Appl. Mech. Engrg.* 191 (32) (2002) 3491–3507. [doi:10.1016/S0045-7825\(02\)00287-6](https://doi.org/10.1016/S0045-7825(02)00287-6).
- [7] C. M. Bishop, H. Bishop, *Deep learning: foundations and concepts*, Springer, 2024.

- [8] E. Haghighat, M. Raissi, A. Moure, H. Gomez, R. Juanes, A physics-informed deep learning framework for inversion and surrogate modeling in solid mechanics, *Comput. Methods Appl. Mech. Engrg.* 379 (2021) 113741. doi:10.1016/j.cma.2021.113741.
- [9] C. K. I. Williams, C. E. Rasmussen, *Gaussian processes for machine learning*, MIT Press, 2006.
- [10] T. J. Santner, B. J. Williams, W. I. Notz, B. J. Williams, *The design and analysis of computer experiments*, Springer, 2003.
- [11] J. D. Martin, T. W. Simpson, Use of kriging models to approximate deterministic computer models, *AIAA J.* 43 (4) (2005) 853–863. doi:10.2514/1.8650.
- [12] A. Forrester, A. Sobester, A. Keane, *Engineering design via surrogate modelling: a practical guide*, John Wiley & Sons, 2008.
- [13] K. J. Koh, F. Cirak, Stochastic PDE representation of random fields for large-scale Gaussian process regression and statistical finite element analysis, *Comput. Methods Appl. Mech. Engrg.* 417 (2023) 116358. doi:10.1016/j.cma.2023.116358.
- [14] G. I. Schuëller, H. A. Jensen, Computational methods in optimization considering uncertainties—an overview, *Comput. Methods Appl. Mech. Engrg.* 198 (1) (2008) 2–13. doi:10.1016/j.cma.2008.05.004.
- [15] I. Doltsinis, Z. Kang, Robust design of structures using optimization methods, *Comput. Methods Appl. Mech. Engrg.* 193 (23) (2004) 2221–2237. doi:10.1016/j.cma.2003.12.055.
- [16] A. Asadpoure, M. Tootkaboni, J. K. Guest, Robust topology optimization of structures with uncertainties in stiffness—application to truss structures, *Comput. Struct.* 89 (2011) 1131–1141. doi:10.1016/j.compstruc.2010.11.004.
- [17] G. Da Silva, E. Cardoso, Stress-based topology optimization of continuum structures under uncertainties, *Comput. Methods Appl. Mech. Engrg.* 313 (2017) 647–672. doi:10.1016/j.cma.2016.09.049.
- [18] I. Ben-Yelun, A. O. Yuksel, F. Cirak, Robust topology optimisation of lattice structures with spatially correlated uncertainties, *Struct. Multidiscip. Optim.* 67 (2) (2024) 16. doi:10.1007/s00158-023-03716-4.
- [19] J. T. Oden, R. Moser, O. Ghattas, Computer predictions with quantified uncertainty, Part I, *SIAM News* 43 (2010) 1–3.
- [20] C. J. Roy, W. L. Oberkampf, A comprehensive framework for verification, validation, and uncertainty quantification in scientific computing, *Comput. Methods Appl. Mech. Engrg.* 200 (25–28) (2011) 2131–2144. doi:10.1016/j.cma.2011.03.016.
- [21] A. O’Hagan, Bayesian analysis of computer code outputs: A tutorial, *Reliab. Eng. Syst. Saf.* 91 (10–11) (2006) 1290–1300. doi:10.1016/j.ress.2005.11.025.
- [22] A. Der Kiureghian, O. Ditlevsen, Aleatory or epistemic? Does it matter?, *Struct. Saf.* 31 (2) (2009) 105–112. doi:10.1016/j.strusafe.2008.06.020.
- [23] S. Sakata, F. Ashida, M. Zako, Structural optimization using kriging approximation, *Comput. Methods Appl. Mech. Engrg.* 192 (7–8) (2003) 923–939. doi:10.1016/S0045-7825(02)00617-5.
- [24] J. Zhang, A. A. Taflanidis, J. Medina, Sequential approximate optimization for design under uncertainty problems utilizing kriging metamodeling in augmented input space, *Comput. Methods Appl. Mech. Engrg.* 315 (2017) 369–395. doi:10.1016/j.cma.2016.10.042.
- [25] M. Tootkaboni, A. Asadpoure, J. K. Guest, Topology optimization of continuum structures under uncertainty—a polynomial chaos approach, *Comput. Methods Appl. Mech. Engrg.* 201 (2012) 263–275. doi:10.1016/j.cma.2011.09.009.
- [26] M. A. Bessa, S. Pellegrino, Design of ultra-thin shell structures in the stochastic post-buckling range using Bayesian machine learning and optimization, *Int. J. Solids Struct.* 139 (2018) 174–188. doi:10.1016/j.ijsolstr.2018.01.035.
- [27] T. Zafar, Y. Zhang, Z. Wang, An efficient kriging based method for time-dependent reliability based robust design optimization via evolutionary algorithm, *Comput. Methods Appl. Mech. Engrg.* 372 (2020) 113386. doi:10.1016/j.cma.2020.113386.
- [28] V. Keshavarzadeh, R. M. Kirby, A. Narayan, Stress-based topology optimization under uncertainty via simulation-based Gaussian process, *Comput. Methods Appl. Mech. Engrg.* 365 (2020) 112992. doi:10.1016/j.cma.2020.112992.
- [29] K. Tian, Z. Li, L. Huang, K. Du, L. Jiang, B. Wang, Enhanced variable-fidelity surrogate-based optimization framework by Gaussian process regression and fuzzy clustering, *Comput. Methods Appl. Mech. Engrg.* 366 (2020) 113045. doi:10.1016/j.cma.2020.113045.
- [30] M. C. Kennedy, A. O’Hagan, Bayesian calibration of computer models, *J. R. Stat. Soc. Series B Stat. Methodol.* 63 (2001) 425–464. doi:10.1111/1467-9868.00294.
- [31] K. P. Murphy, *Machine learning: A probabilistic perspective*, MIT Press, 2012.
- [32] A. Gelman, J. B. Carlin, H. S. Stern, D. B. Dunson, A. Vehtari, D. B. Rubin, *Bayesian data analysis*, 3rd Edition, CRC Press, 2014.
- [33] C. Grigo, P.-S. Koutsourelakis, Bayesian model and dimension reduction for uncertainty propagation: applications in random media, *SIAM-ASA J. Uncertain. Quantif.* 7 (2019) 292–323. doi:10.1137/17M1155867.
- [34] M. Girolami, E. Febrianto, G. Yin, F. Cirak, The statistical finite element method (statFEM) for coherent synthesis of observation data and model predictions, *Comput. Methods Appl. Mech. Engrg.* 375 (2021) 113533. doi:10.1016/j.cma.2020.113533.
- [35] P. G. Constantine, E. Dow, Q. Wang, *Active subspace methods in theory and practice: applications to kriging surfaces*, SIAM

- J. Sci. Comput. 36 (4) (2014) A1500–A1524. [doi:10.1137/130916138](#).
- [36] M. A. Bouhlel, N. Bartoli, A. Otsmane, J. Morlier, Improving kriging surrogates of high-dimensional design models by partial least squares dimension reduction, *Struct. Multidiscip. Optim.* 53 (2016) 935–952. [doi:10.1007/s00158-015-1395-9](#).
 - [37] D. Gaudrie, R. Le Riche, V. Picheny, B. Enaux, V. Herbert, Modeling and optimization with Gaussian processes in reduced eigenbases, *Struct. Multidiscip. Optim.* 61 (2020) 2343–2361. [doi:10.1007/s00158-019-02458-6](#).
 - [38] R. R. Lam, O. Zahm, Y. M. Marzouk, K. E. Willcox, Multifidelity dimension reduction via active subspaces, *SIAM J. Sci. Comput.* 42 (2) (2020) A929–A956. [doi:10.1137/18M1214123](#).
 - [39] F. Romor, M. Tezzele, M. Mrosek, C. Othmer, G. Rozza, Multi-fidelity data fusion through parameter space reduction with applications to automotive engineering, *Int. J. Numer. Methods. Eng.* (2021). [doi:10.1002/nme.7349](#).
 - [40] R. Tripathy, I. Bilonis, M. Gonzalez, Gaussian processes with built-in dimensionality reduction: applications to high-dimensional uncertainty propagation, *J. Comput. Phys.* 321 (2016) 191–223. [doi:10.1016/j.jcp.2016.05.039](#).
 - [41] M. Guo, J. S. Hesthaven, Reduced order modeling for nonlinear structural analysis using Gaussian process regression, *Comput. Methods Appl. Mech. Engrg.* 341 (2018) 807–826. [doi:10.1016/j.cma.2018.07.017](#).
 - [42] P. Tsilifis, P. Pandita, S. Ghosh, V. Andreoli, T. Vandeputte, L. Wang, Bayesian learning of orthogonal embeddings for multi-fidelity Gaussian processes, *Comput. Methods Appl. Mech. Engrg.* 386 (2021) 114147. [doi:10.1016/j.cma.2021.114147](#).
 - [43] M. I. Jordan, Z. Ghahramani, T. S. Jaakkola, L. K. Saul, An introduction to variational methods for graphical models, *Mach. Learn.* 37 (1999) 183–233. [doi:10.1023/A:1007665907178](#).
 - [44] D. M. Blei, A. Kucukelbir, J. D. McAuliffe, Variational inference: A review for statisticians, *J. Am. Stat. Assoc.* 112 (518) (2017) 859–877. [doi:10.1080/01621459.2017.1285773](#).
 - [45] J. Povala, I. Kazlauskaitė, E. Febrianto, F. Cirak, M. Girolami, Variational Bayesian approximation of inverse problems using sparse precision matrices, *Comput. Methods Appl. Mech. Engrg.* 393 (2022) 114712. [doi:10.1016/j.cma.2022.114712](#).
 - [46] A. Vadeboncoeur, Ö. D. Akyildiz, I. Kazlauskaitė, M. Girolami, F. Cirak, Fully probabilistic deep models for forward and inverse problems in parametric PDEs, *J. Comput. Phys.* 491 (2023) 112369. [doi:10.1016/j.jcp.2023.112369](#).
 - [47] M. D. Hoffman, D. M. Blei, C. Wang, J. Paisley, Stochastic variational inference, *J. Mach. Learn. Res.* 14 (2013) 1303–1347.
 - [48] J. Li, L. Fuxin, S. Todorovic, Efficient Riemannian optimization on the Stiefel manifold via the Cayley transform, *arXiv preprint arXiv:2002.01113* (2020). [doi:10.48550/arXiv.2002.01113](#).
 - [49] J. Quinonero-Candela, C. E. Rasmussen, A unifying view of sparse approximate Gaussian process regression, *J. Mach. Learn. Res.* 6 (2005) 1939–1959.
 - [50] M. Titsias, Variational learning of inducing variables in sparse Gaussian processes, in: *Artificial Intelligence and Statistics*, PMLR, 2009, pp. 567–574.
 - [51] N. D. Lawrence, Learning for larger datasets with the Gaussian process latent variable model, in: *Artificial Intelligence and Statistics*, PMLR, 2007, pp. 243–250.
 - [52] A. C. Damianou, M. K. Titsias, N. D. Lawrence, Variational inference for latent variables and uncertain inputs in Gaussian processes, *J. Mach. Learn. Res.* 17 (2016) 1–62.
 - [53] F. Leibfried, V. Dutordoir, S. John, N. Durrande, A tutorial on sparse Gaussian processes and variational inference, *arXiv preprint arXiv:2012.13962* (2020). [doi:10.48550/arXiv.2012.13962](#).
 - [54] J. Hensman, N. Fusi, N. D. Lawrence, Gaussian processes for big data, *arXiv preprint arXiv:1309.6835* (2013). [doi:10.48550/arXiv.1309.6835](#).
 - [55] D. P. Kingma, J. Ba, Adam: A method for stochastic optimization, *arXiv preprint arXiv:1412.6980* (2014). [doi:10.48550/arXiv.1412.6980](#).
 - [56] D. P. Kingma, M. Welling, Auto-encoding variational Bayes, *arXiv preprint arXiv:1312.6114* (2013). [doi:10.48550/arXiv.1312.6114](#).
 - [57] D. P. Kingma, M. Welling, et al., An introduction to variational autoencoders, *Trends Mach. Learn.* 12 (4) (2019) 307–392.
 - [58] D. Wipf, S. Nagarajan, A new view of automatic relevance determination, *Adv. Neural. Inf. Process. Syst.* 20 (2007).
 - [59] R. Horst, P. M. Pardalos, *Handbook of global optimization*, Vol. 2, Springer, 2013.
 - [60] G. H. Golub, C. F. Van Loan, *Matrix computations*, JHU Press, 2013.
 - [61] M. D. McKay, R. J. Beckman, W. J. Conover, A comparison of three methods for selecting values of input variables in the analysis of output from a computer code, *Technometrics* 42 (1) (2000) 55–61. [doi:10.1080/00401706.2000.10485979](#).
 - [62] D. Zhang, A coefficient of determination for generalized linear models, *Am. Stat.* 71 (4) (2017) 310–316. [doi:10.1080/00031305.2016.1256839](#).
 - [63] A. Gretton, K. M. Borgwardt, M. J. Rasch, B. Schölkopf, A. Smola, A kernel two-sample test, *J. Mach. Learn. Res.* 13 (1) (2012) 723–773.
 - [64] T. Zhou, Y. Peng, Structural reliability analysis via dimension reduction, adaptive sampling, and Monte Carlo simulation, *Struct. Multidiscip. Optim.* 62 (5) (2020) 2629–2651. [doi:10.1007/s00158-020-02633-0](#).
 - [65] J. N. Fuhg, A. Fau, U. Nackenhorst, State-of-the-art and comparative review of adaptive sampling methods for kriging, *Arch.*

- Comput. Methods Eng. 28 (2021) 2689–2747. [doi:10.1007/s11831-020-09474-6](https://doi.org/10.1007/s11831-020-09474-6).
- [66] M. Rixner, P.-S. Koutsourelakis, A probabilistic generative model for semi-supervised training of coarse-grained surrogates and enforcing physical constraints through virtual observables, *J. Comput. Phys.* 434 (2021) 110218. [doi:10.1016/j.jcp.2021.110218](https://doi.org/10.1016/j.jcp.2021.110218).
 - [67] K. Linka, A. Schäfer, X. Meng, Z. Zou, G. E. Karniadakis, E. Kuhl, Bayesian physics informed neural networks for real-world nonlinear dynamical systems, *Comput. Methods Appl. Mech. Engrg.* 402 (2022) 115346. [doi:10.1016/j.cma.2022.115346](https://doi.org/10.1016/j.cma.2022.115346).
 - [68] T. P. Bohlin, *Practical grey-box process identification: theory and applications*, Springer, 2006.
 - [69] N. Jaquier, L. Rozo, High-dimensional Bayesian optimization via nested Riemannian manifolds, *Adv. Neural. Inf. Process. Syst.* 33 (2020) 20939–20951.
 - [70] Y. Zelickman, J. K. Guest, Construction aware optimization of concrete plate thicknesses, *Eng. Struct.* 296 (2023) 116889. [doi:10.1016/j.engstruct.2023.116889](https://doi.org/10.1016/j.engstruct.2023.116889).
 - [71] J. A. Lee, M. Verleysen, et al., *Nonlinear dimensionality reduction*, Vol. 1, Springer, 2007.