

Solving Bayesian Inverse Problems via Variational Autoencoders

Hwan Goh
Sheroze Sherifdeen
Jonathan Wittmer
Tan Bui-Thanh

Oden Institute for Computational Engineering and Sciences
The University of Texas at Austin
Austin, TX 78712

HWAN.GOH@GMAIL.COM
SHEROZE@ODEN.UTEXAS.EDU
JONATHAN.WITTMER@UTEXAS.EDU
TANBUI@ODEN.UTEXAS.EDU

Editors: Joan Bruna, Jan S Hesthaven, Lenka Zdeborova

Abstract

In recent years, the field of machine learning has made phenomenal progress in the pursuit of simulating real-world data generation processes. One notable example of such success is the variational autoencoder (VAE). In this work, with a small shift in perspective, we leverage and adapt VAEs for a different purpose: uncertainty quantification in scientific inverse problems. We introduce UQ-VAE: a flexible, adaptive, hybrid data/model-constrained framework for training neural networks capable of rapid modelling of the posterior distribution representing the unknown parameter of interest. Specifically, from divergence-based variational inference, our framework is derived such that most of the information usually present in scientific inverse problems is fully utilized in the training procedure. Additionally, this framework includes an adjustable hyperparameter that allows selection of the notion of distance between the posterior model and the target distribution. This introduces more flexibility in controlling how optimization directs the learning of the posterior model. Further, this framework possesses an inherent adaptive optimization property that emerges through the learning of the posterior uncertainty. Numerical results for an elliptic PDE-constrained Bayesian inverse problem are provided to verify the proposed framework.

Keywords: Machine Learning, Uncertainty Quantification, Bayesian Inverse Problems, Variational Autoencoders

1. Introduction

The challenge of generative modelling is often approached through optimization of a parameterized likelihood function $p_{\theta}(\mathbf{x})$ to learn the optimal parameters θ . Latent variable models augment this likelihood function using an unobservable latent variable in hopes that one can also learn useful latent representations that facilitate reconstruction and generation of high-dimensional distributions. Variational autoencoders (VAE) were introduced in [Kingma and Welling \(2013\)](#) to train a generating function that maps from the latent space to the target data distribution. The approach in [Kingma and Welling \(2013\)](#) assumes the latent distribution to simply be an isotropic Gaussian $p(\mathbf{z}) = \mathcal{N}(\mathbf{z}|\mathbf{0}, \mathbf{I})$; thereby shifting all the burden of generative modelling to the learning of the generating function. To address the infeasibly large number of draws from the isotropic Gaussian required to ensure significant contribution to the likelihood function, [Kingma and Welling \(2013\)](#) proposes to instead draw from the posterior distribution $p_{\theta}(\mathbf{z}|\mathbf{x})$. This introduces the subtask of

variational inference; the modelling of the posterior with a distribution $q_\phi(\mathbf{z}|\mathbf{x})$ parameterized by an additional set of learnable parameters ϕ . The learning of the parameters ϕ can be achieved through the minimization of the Evidence Lower Bound (ELBO) [Blei et al. \(2017\)](#); a quantity obtainable through the Kullback-Leibler divergence (KLD) between the model posterior and the target posterior.

Although the goal of generative modelling fundamentally differs from that of scientific inverse problems, we claim that VAEs are also well suited for the latter challenge. In the setting of physical systems, scientific inverse problems task us with determining parameters-of-interest (PoI) given observations of a state variable. Often, these parameters are coefficients in a partial differential equation (PDE) governing the system and the observations are discrete measurements of the state variables of the PDE. Assuming that our observations are corrupted by additive noise, then the equation

$$\mathbf{y}_{\text{obs}} = \mathcal{F}(\mathbf{u}) + \mathbf{e} \quad (1)$$

represents the connection between the PoI and observational data. Here, $\mathbf{u}$ denotes a discretized representation of the PoI, $\mathcal{F}$ denotes the parameter-to-observable (PtO) map and $\mathbf{y}_{\text{obs}}$ denotes our observational data.

To bridge the two contexts, we consider the latent space and generating function in latent variable models to be the PoI space and the PtO map in inverse problems. Therefore, while the former uses inference to facilitate generative modelling, the latter uses generative modelling to facilitate inference. We contrast the two settings under this viewpoint in order to expose the utility of VAEs for scientific inverse problems:

1. For generative modelling, the focus of training VAEs is to learn the generating function. In contrast, for scientific inverse problems, the PtO map need not be learnt as it can often be modelled using physics-based numerical methods.
2. For generative modelling, the latent space of VAEs is an unspecified quantity with properties emerging almost entirely through the process of learning the generating function. In contrast, for scientific inverse problems, by considering the latent space to be the PoI space, the latent space therefore possesses structure that represents the physical properties of the PoI.

Therefore, the context of scientific inverse problems provides more information that can be exploited to improve the training procedure of a VAE. Indeed, in reference to the first point, one can surrogate the generating function with a physics-based numerical model of the PtO map. In doing so, the focus of the training is shifted completely towards inference. In reference to the second point, knowledge of the physical properties of the PoI can be encoded in the prior model which, therefore, allows for a more informative latent distribution than the isotropic Gaussian used in [Kingma and Welling \(2013\)](#). Further, through scientific experimentation or simulation, one may have on hand a dataset of PoI and observation pairs $\left\{ \left(\mathbf{u}^{(m)}, \mathbf{y}_{\text{obs}}^{(m)} \right) \right\}_{m=1}^M$. To draw an analogy within the context of generative modelling, the datapoints $\mathbf{u}^{(m)}$ are analogous to latent data that is usually not available. This advantage provides the main motivation for our proposed UQ-VAE framework: a mathematically justified training procedure for a VAE that can utilize the latent data we possess.

2. Preliminaries and Motivation

In scientific inverse problems, the solution process for obtaining estimates of the PoI often involves the optimization of a functional of the form

$$\min_{\mathbf{u}} \|\mathbf{y}_{\text{obs}} - \mathcal{F}(\mathbf{u})\|_2^2. \quad (2)$$

Whilst the observations are finite-dimensional entities, the PoI often exists in an infinite-dimensional function space. Thus, our task is complicated by the inherent ill-posed nature of inverse problems; there are many possible estimates of the PoI that are consistent with the observations. To alleviate this issue, regularization can be introduced to essentially reduce the size of the solution space. With the introduction of a regularization term, the optimization problem (2) becomes

$$\min_{\mathbf{u}} \|\mathbf{y}_{\text{obs}} - \mathcal{F}(\mathbf{u})\|_2^2 + \mathcal{R}(\mathbf{u}). \quad (3)$$

The optimization of (3) is often a time-consuming process; for example when using derivative-based iterative optimization algorithms. This motivates the need for an inverse problems solver that can rapidly output PoI estimates from given observational data.

2.1. Learning a Solver for Deterministic Inverse Problems

We briefly discuss data-driven deterministic inverse problem solvers to motivate both the concept of uncertainty quantification as well as our proposed training procedure.

Suppose that we are tasked with learning an inverse problem solver using a training dataset of M PoI and observation pairs $\left\{ \left(\mathbf{u}^{(m)}, \mathbf{y}_{\text{obs}}^{(m)} \right) \right\}_{m=1}^M$. If we elect to use a neural network Ψ , then our solver is parameterized by the weights $\mathbf{W}$ of the network. Analogous to (3), the training of this solver requires the optimization of the following functional:

$$\min_{\mathbf{W}} \frac{1}{M} \sum_{m=1}^M \left\| \mathbf{u}^{(m)} - \Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right\|_2^2 + \mathcal{R}(\mathbf{W}). \quad (4)$$

Following the offline training stage of the neural network, we obtain a rapid online inverse problem solver that can output an estimate of our PoI given observation data. However, the ill-posed nature of inverse problems is inherited by this task; there are many possible weights that can be used to parameterize a solver that outputs PoI estimates consistent with the input observational data. Further, while the regularization term in (3) can be designed based on some prior knowledge of the physical properties of the PoI, the physical interpretation of the neural network weights is not so clear. Therefore, there is no natural choice for the regularization of the optimization problem (4).

Instead of regularizing the weights of the network directly, one possible approach is to regularize the output of the network. For example, we could make use of the PtO map so that (4) becomes

$$\min_{\mathbf{W}} \frac{1}{M} \sum_{m=1}^M \left\| \mathbf{u}^{(m)} - \Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right\|_2^2 + \left\| \mathbf{y}_{\text{obs}}^{(m)} - \mathcal{F} \left(\Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right) \right\|_2^2. \quad (5)$$

In doing so, the PtO map indirectly informs the weights of the neural network of the inversion task. An additional benefit with the inclusion of the PtO map is that we are also able to include noise

regularization. That is, if we possess some knowledge about the properties of the noise $\mathbf{e}$ afflicting our observational data, we may quantify it and include it in the second term to yield

$$\min_{\mathbf{W}} \frac{1}{M} \sum_{m=1}^M \left\| \mathbf{u}^{(m)} - \Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right\|_2^2 + \left\| \mathcal{M} \left(\mathbf{y}_{\text{obs}}^{(m)} - \mathcal{F} \left(\Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right) \right) \right\|_2^2 \quad (6)$$

where $\mathcal{M}$ represents some noise regularization mapping. Furthermore, if we possess some prior knowledge about the physical properties of the parameter of interest, we may also include it as follows:

$$\min_{\mathbf{W}} \frac{1}{M} \sum_{m=1}^M \left\| \mathbf{u}^{(m)} - \Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right\|_2^2 + \left\| \mathcal{M} \left(\mathbf{y}_{\text{obs}}^{(m)} - \mathcal{F} \left(\Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) \right) \right) \right\|_2^2 \quad (7a)$$

$$+ \left\| \mathcal{P} \left(\Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right) - \boldsymbol{\mu}_{\text{pr}} \right) \right\|_2^2 \quad (7b)$$

where $\mathcal{P}$ is some prior regularization mapping and $\boldsymbol{\mu}_{\text{pr}}$ is the prior reference value of parameters.

Training a neural network through the optimization problem (7) yields a learned inverse problems solver that outputs a point estimate of our PoI. As it is, this deterministic solver is unable to provide information about the accuracy/uncertainty of the estimate. It would be more ideal to have a probabilistic interpretation of our learned solver that facilitates uncertainty quantification (UQ). With this in mind, we are motivated to view inverse problems under the framework of Bayesian statistics. *In this setting, we instead work towards a UQ-VAE solver for Bayesian inverse problems which, in turn, allows us to formally establish the regularization terms in (7). Conversely, by ignoring the uncertainty estimate, the UQ-VAE solver recovers the deterministic setting (7).*

2.2. Learning a Solver for Bayesian Inverse Problems

Under the statistical framework, the PoI of an inverse problem is considered to be a random variable instead of an unknown value. Consequently, the solution of the statistical inverse problem is a probability distribution instead of a single estimated value. By adopting a statistical framework instead of a deterministic one, the question asked by the inverse problem essentially changes from “what is the value of our parameter?” to “how accurate is the estimate of our parameter?”. Furthermore, the statistical framework attempts to remove the ill-posedness of inverse problems by restating the inverse problem as a well-posed extension in a larger space of probability distributions [Kaipio and Somersalo \(2006\)](#); [Dashti and Stuart \(2013\)](#).

With the additive noise assumption we consider the following observational model

$$Y = \mathcal{F}(U) + E \quad (8)$$

where $\mathcal{F}$ is the PtO map and Y , U and E are the random variables representing the observational data, PoI and noise model respectively. Since U is unknown to us, we represent its uncertainty with a probability distribution constructed from prior knowledge. However, we do not use equation (8) directly when working under a statistical framework. Instead, we primarily consider the relations between the probability distributions of the random variables involved to answer the question: “given the observational data $\mathbf{y}_{\text{obs}}$, what is the distribution of the PoI U responsible for our measurement?”. Therefore, the conditional density $p_{U|Y}(\mathbf{u}|\mathbf{y}_{\text{obs}})$ is the solution to the statistical parameter estimation problem under the Bayesian framework. To approximate this conditional density, we utilize

Bayes' Theorem to form a model of $p_{U|Y}(\mathbf{u}|Y = \mathbf{y}_{\text{obs}})$ called the posterior distribution which we denote as p_{post} :

$$p_{\text{post}}(\mathbf{u}|\mathbf{y}_{\text{obs}}) \propto p_{\text{lkhd}}(\mathbf{y}_{\text{obs}}|\mathbf{u}) p_{\text{pr}}(\mathbf{u}). \quad (9)$$

where p_{lkhd} is the likelihood model and p_{pr} is the prior model.

Formulating the parameter estimation problem under a Bayesian framework requires us to obtain (9). This challenges us with the completion of three tasks:

1. construct the likelihood model p_{lkhd} interrelating the data and the unknown,
2. use prior information to construct a prior probability density p_{pr}
3. develop methods to extract meaningful information from the posterior density p_{post} .

To address these three tasks, two assumptions are often made. The first assumption supposes that the noise E is mutually independent with respect to the PoI U . Then, using our observation model (8) and marginalization of the noise E , we obtain the following likelihood model:

$$p_{\text{lkhd}} = p_E(\mathbf{y}_{\text{obs}} - \mathcal{F}(\mathbf{u})). \quad (10)$$

The second assumption supposes that all are random variables are Gaussian with $\mathcal{N}(\boldsymbol{\mu}_E, \boldsymbol{\Gamma}_E)$ and $\mathcal{N}(\boldsymbol{\mu}_{\text{pr}}, \boldsymbol{\Gamma}_{\text{pr}})$. With this, our posterior model becomes

$$p_{\text{post}}(\mathbf{u}|\mathbf{y}_{\text{obs}}) \propto p_E(\mathbf{y}_{\text{obs}} - \mathcal{F}(\mathbf{u})) p_{\text{pr}}(\mathbf{u}) \quad (11a)$$

$$= \exp\left(-\frac{1}{2}\left(\|\mathbf{y}_{\text{obs}} - \mathcal{F}(\mathbf{u}) - \boldsymbol{\mu}_E\|_{\boldsymbol{\Gamma}_E^{-1}}^2 + \|\mathbf{u} - \boldsymbol{\mu}_{\text{pr}}\|_{\boldsymbol{\Gamma}_{\text{pr}}^{-1}}^2\right)\right). \quad (11b)$$

In order to perform UQ with given observational data $\mathbf{y}_{\text{obs}}$, one often seeks an approximation of the posterior covariance of (11). A commonly used approximation is the Laplace approximation [Evans and Swartz \(2000\)](#); [Press \(2009\)](#); [Stigler \(1986\)](#); [Tierney and Kadane \(1986\)](#); [Wong \(2001\)](#)

$$\boldsymbol{\Gamma}_{\text{post}} = (\mathbf{J}_{\mathcal{F}}(\mathbf{u}_{\text{map}})^T \boldsymbol{\Gamma}_E^{-1} \mathbf{J}_{\mathcal{F}}(\mathbf{u}_{\text{map}}) + \boldsymbol{\Gamma}_{\text{pr}}^{-1})^{-1} \quad (12)$$

where $\mathbf{u}_{\text{map}}$ is the maximum a posteriori (MAP) estimate and $\mathbf{J}_{\mathcal{F}}(\mathbf{u}_{\text{map}})$ is the Jacobian of $\mathcal{F}$ at the MAP estimate. The MAP estimate can be obtained through optimization of the exponentiated functional in (11b); which often requires a potentially expensive derivative-based iterative optimization procedure. Markov Chain Monte Carlo (MCMC) methods can also be employed but such methods notably suffer the curse of dimensionality.

This desire for efficient UQ motivates our work. Although the training procedure of a Bayesian inverse problems solver may be computationally demanding, it represents an offline stage. Once trained, the online stage of UQ has a computational cost governed only by the forward propagation of the observational data through the trained neural network; which is often a very computationally efficient procedure.

Additionally, the PoI data $\mathbf{u}^{(m)}$ in the training dataset $\left\{\left(\mathbf{u}^{(m)}, \mathbf{y}_{\text{obs}}^{(m)}\right)\right\}_{m=1}^M$ —encoding information of the underlying physics via the PtO map—usually plays no role in VAE methods. Information about the PoI is often only encoded in the prior model. This motivates the data-driven aspect of our

approach in that it not only uses prior information but also takes advantage of, and hence constrained by, the physics encoded in the underlying PtO map.

With this motivation in mind, we now detail a framework for training our solver. When comparing with (7) from Section 2.1, the exponentiated functional in (11b) looks like an appealing candidate for the regularization terms. We formalize the inclusion of these terms as regularizers through the derivation of our proposed model-constrained VAE framework for UQ.

3. UQ-VAE: Flexible, Adaptive, Data/Model-Constrained Framework for UQ

We begin by detailing the derivation of our UQ-VAE framework before discussing from the more tractable perspective of optimization. While the former centers the discussion more in the setting of variational inference, the latter is centered more towards regularization of the training procedure for our solver. Furthermore, whilst the flexibility of our framework emerges through the viewpoint of variational inference, the optimization viewpoint uncovers its inherent adaptive properties.

3.1. Derivation of the UQ-VAE Framework

Let $p(\mathbf{u}|\mathbf{y})$ denote the target posterior density we wish to estimate and let $q_\phi(\mathbf{u}|\mathbf{y})$ denote our model (approximate) of the target density with statistics parameterized by ϕ . To optimize for the parameters ϕ , we require some notion of distance between our model posterior and target posterior. In this work, instead of using the KLD to obtain the ELBO, we elect to use the following family of Jensen-Shannon divergences (JSD) [Nielsen \(2010\)](#):

$$\text{JS}_\alpha(q||p) = \alpha \text{KL}(q||(1-\alpha)q + \alpha p) + (1-\alpha) \text{KL}(p||(1-\alpha)q + \alpha p). \quad (13)$$

With this, we have the following result.

Lemma 1 *Let $\alpha \in (0, 1)$. Then*

$$\frac{1}{\alpha} \text{JS}_\alpha(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) = -\mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (14a)$$

$$+ \log(p(\mathbf{y})) \quad (14b)$$

$$- L_{\text{JS}_\alpha}(\phi, \mathbf{y}) \quad (14c)$$

where

$$L_{\text{JS}_\alpha}(\phi, \mathbf{y}) = \frac{1-\alpha}{\alpha} \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y})} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] + \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{y}, \mathbf{u})}{q_\phi(\mathbf{u}|\mathbf{y})} \right) \right]. \quad (15)$$

The proof is given in Section A.1 of the Appendix. From an implementation perspective, Lemma 1 does not offer much insight. In particular, we would like to have direct access to the quotients within the expectations. Therefore, we offer the following corollary which yields useful KLD terms:

Corollary 2 *Let $\alpha \in (0, 1)$ and consider again (14). Equation (14) is bounded above such that:*

$$\frac{1}{\alpha} \text{JS}_\alpha(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) \leq -\text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) \quad (16a)$$

$$+ \log(p(\mathbf{y})) - \log(1 - \alpha) - \frac{(1 - \alpha) \log(1 - \alpha)}{\alpha} \quad (16b)$$

$$+ \frac{1 - \alpha}{\alpha} \text{KL}(p(\mathbf{u}|\mathbf{y})||q_\phi(\mathbf{u}|\mathbf{y})) \quad (16c)$$

$$- \mathbb{E}_{\mathbf{u} \sim q_\phi} [\log(p(\mathbf{y}|\mathbf{u}))] + \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u})). \quad (16d)$$

In particular, we have that

$$-L_{\text{JS}_\alpha}(\phi, \mathbf{y}) \leq -\frac{(1 - \alpha) \log(1 - \alpha)}{\alpha} + \frac{1 - \alpha}{\alpha} \text{KL}(p(\mathbf{u}|\mathbf{y})||q_\phi(\mathbf{u}|\mathbf{y})) \quad (17a)$$

$$- \mathbb{E}_{\mathbf{u} \sim q_\phi} [\log(p(\mathbf{y}|\mathbf{u}))] + \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u})). \quad (17b)$$

Proof *The bounds are obtained simply by dropping the α from within the logarithms in (14a) and (15). From there, the assertions are readily available. ■*

The significance of Corollary 2 is that the minimization of the upper bound

$$\frac{1 - \alpha}{\alpha} \text{KL}(p(\mathbf{u}|\mathbf{y})||q_\phi(\mathbf{u}|\mathbf{y})) - \mathbb{E}_{\mathbf{u} \sim q_\phi} [\log(p(\mathbf{y}|\mathbf{u}))] + \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u})) \quad (18)$$

with respect to ϕ could minimize

$$\frac{1}{\alpha} \text{JS}_\alpha(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) + \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) \quad (19)$$

which would, therefore, perform the task of variational inference. From a VAE perspective, notice that the second and third term of (18) together form the negative of the ELBO. Relating back to the exponentiated functional (11b), from a perspective of Bayesian inverse problems, the second term in (18) is the likelihood model containing the PtO map and the third term contains the prior model $p(\mathbf{u})$ representing information about the PoI. The first term in (18) represents information regarding the target posterior. This is the key achievement of utilizing the JSD as the notion of distance between the model posterior and target posterior as it is this term that allows for the inclusion of paired PoI-observation datasets, and hence encoding the underlying physics in variational inference.

Through the choice of $\alpha \in (0, 1)$, the family of JSDs (13) allows for interpolation between the zero-forcing $\text{KL}(p||q)$ as $\alpha \rightarrow 0$ and the zero-avoiding $\text{KL}(q||p)$ as $\alpha \rightarrow 1$ Bishop (2006); Murphy (2012). With this in mind, we make the following observations:

- In (13), it is clear that if $\alpha = 1$ then we recover the usual zero-avoiding $\text{KL}(q||p)$ used in Kingma and Welling (2013). Indeed, as $\alpha \rightarrow 1$, then the first term in (18) tends to 0 which recovers the negative of the ELBO.
- In (19), the presence of the KLD term ensures that our model posterior will inherently be zero-avoiding. However, the $\frac{1}{\alpha}$ scaling factor ensures that as $\alpha \rightarrow 0$, the consequently zero-forcing JSD dominates the KLD.

Therefore, our UQ-VAE framework essentially retains the full flexibility of the JSD family. Thus, with only an adjustment of a single scalar value, our framework allows the selection of the notion of distance used by the optimization routine to direct the model posterior towards the target posterior. This, in turn, translates to control of the balance of data-fitting and regularization used in the training procedure as discussed in the following section.

3.2. Regularized Optimization Problem

We now move our discussion towards implementation by detailing the optimization problem implied by the UQ-VAE framework. In particular, we employ approximations of (18) from Section 3.1 to yield an implementable loss functional.

We begin by forming an approximation to the first term in (18). First, we apply the well-known fact that the minimization of the KLD between the empirical and model distributions is equivalent to maximization of the likelihood function with respect to ϕ . Next, we form a Monte-Carlo estimation using our PoI data. However, since every observation datapoint $\mathbf{y}_{\text{obs}}^{(m)}$ is associated with only one PoI datapoint $\mathbf{u}^{(m)}$, our approximation is somewhat crude.¹ Finally, we adopt a Gaussian model for our model posterior: $q_\phi(\mathbf{u}|\mathbf{y}_{\text{obs}}^{(m)}) = \mathcal{N}(\mathbf{u}|\boldsymbol{\mu}_{\text{post}}^{(m)}, \boldsymbol{\Gamma}_{\text{post}}^{(m)})$. These assumptions can be summarized by the following chain of equations:

$$\begin{aligned} \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y}_{\text{obs}}^{(m)})} \left[-\log \left(q_\phi(\mathbf{u}|\mathbf{y}_{\text{obs}}^{(m)}) \right) \right] &\lesssim \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u})} \left[-\log \left(q_\phi(\mathbf{u}|\mathbf{y}_{\text{obs}}^{(m)}) \right) \right] \approx -\log \left(q_\phi(\mathbf{u}^{(m)}|\mathbf{y}_{\text{obs}}^{(m)}) \right) \\ &= \frac{D}{2} \log(2\pi) + \frac{1}{2} \log |\boldsymbol{\Gamma}_{\text{post}}^{(m)}| + \frac{1}{2} \left\| \boldsymbol{\mu}_{\text{post}}^{(m)} - \mathbf{u}^{(m)} \right\|_{\boldsymbol{\Gamma}_{\text{post}}^{(m)-1}}^2. \end{aligned}$$

Now we incorporate deep learning into our framework. We consider a neural network that takes in our observation data $\mathbf{y}_{\text{obs}}^{(m)}$ as an input to output the statistics $(\boldsymbol{\mu}_{\text{post}}^{(m)}, \boldsymbol{\Gamma}_{\text{post}}^{(m)})$ of our posterior model. In doing so, we are reparameterizing the statistics of our Gaussian posterior model represented by ϕ with the weights $\mathbf{W}$ of the neural network Ψ ; thereby increasing the learning capacity of our model.

As mentioned in Section 2.2, it is common to adopt Gaussian noise and prior models $\mathcal{N}(\boldsymbol{\mu}_E, \boldsymbol{\Gamma}_E)$ and $\mathcal{N}(\boldsymbol{\mu}_{\text{pr}}, \boldsymbol{\Gamma}_{\text{pr}})$. With this and our approximation above, we obtain the following optimization problem:

$$\min_{\mathbf{W}} \frac{1}{M} \sum_{m=1}^M \frac{1-\alpha}{\alpha} \left(\log |\boldsymbol{\Gamma}_{\text{post}}^{(m)}| + \left\| \boldsymbol{\mu}_{\text{post}}^{(m)} - \mathbf{u}^{(m)} \right\|_{\boldsymbol{\Gamma}_{\text{post}}^{(m)-1}}^2 \right) \quad (21a)$$

$$+ \left\| \mathbf{y}_{\text{obs}}^{(m)} - \mathcal{F} \left(\mathbf{u}_{\text{draw}}^{(m)}(\mathbf{W}) \right) - \boldsymbol{\mu}_E \right\|_{\boldsymbol{\Gamma}_E^{-1}}^2 \quad (21b)$$

$$+ \text{tr} \left(\boldsymbol{\Gamma}_{\text{pr}}^{-1} \boldsymbol{\Gamma}_{\text{post}}^{(m)} \right) + \left\| \boldsymbol{\mu}_{\text{post}}^{(m)} - \boldsymbol{\mu}_{\text{pr}} \right\|_{\boldsymbol{\Gamma}_{\text{pr}}^{-1}}^2 + \log \frac{|\boldsymbol{\Gamma}_{\text{pr}}|}{|\boldsymbol{\Gamma}_{\text{post}}^{(m)}|} \quad (21c)$$

$$\text{where } \left(\boldsymbol{\mu}_{\text{post}}^{(m)}, \boldsymbol{\Gamma}_{\text{post}}^{\frac{1}{2}(m)} \right) = \Psi \left(\mathbf{y}_{\text{obs}}^{(m)}, \mathbf{W} \right), \quad (21d)$$

$$\mathbf{u}_{\text{draw}}^{(m)}(\mathbf{W}) = \boldsymbol{\mu}_{\text{post}}^{(m)} + \boldsymbol{\Gamma}_{\text{post}}^{\frac{1}{2}(m)} \boldsymbol{\epsilon}^{(m)}, \quad (21e)$$

$$\boldsymbol{\epsilon}^{(m)} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_D) \quad (21f)$$

1. Note that more than one prior samples can be used here for more accurate approximation of the expectation.

Drawing connections with our discussion in Section 2.1, our data-misfit (21a) containing our PoI data is regularized by the likelihood model (21b) containing the PtO map (and thus the underlying mathematical model) and the prior model (21c) containing information on our PoI space.

In (21a), the $\frac{1-\alpha}{\alpha}$ factor allows α to adjust the influence of the PoI data and regularization over the optimization procedure. However, the onus of balancing data-fitting and regularization is not completely on the hyperparameter α . The presence of the matrix $\left(\Gamma_{\text{post}}^{(m)}\right)^{-1}$ in the weighted norm of (21a) acts as an adaptive penalty for the data-misfit term. Indeed, during training, this matrix changes as the network weights $\mathbf{W}$, on which the matrix depends on, are optimized. With this in mind, we make the following observations about the two terms in (21a):

- Since $\lim_{\alpha \rightarrow 0} \frac{1-\alpha}{\alpha} = \infty$, then a choice of $\alpha \approx 0$ emphasizes the $\log \left| \Gamma_{\text{post}}^{(m)} \right|$ term. This causes a preference for a small posterior variance which, in turn, creates a large penalization of the data-misfit term by the inverse of the posterior covariance.
- In contrast, a choice of $\alpha \approx 1$ relieves the requirement of a small posterior variance to promote the influence of the PoI data on the optimization problem.

We postulate that this counterbalancing dynamic between the two terms of (21a) induced by α prevents the posterior variance from converging to too low or too high values as the optimization procedure progresses. A schematic of the UQ-VAE framework is displayed in Figure 1.

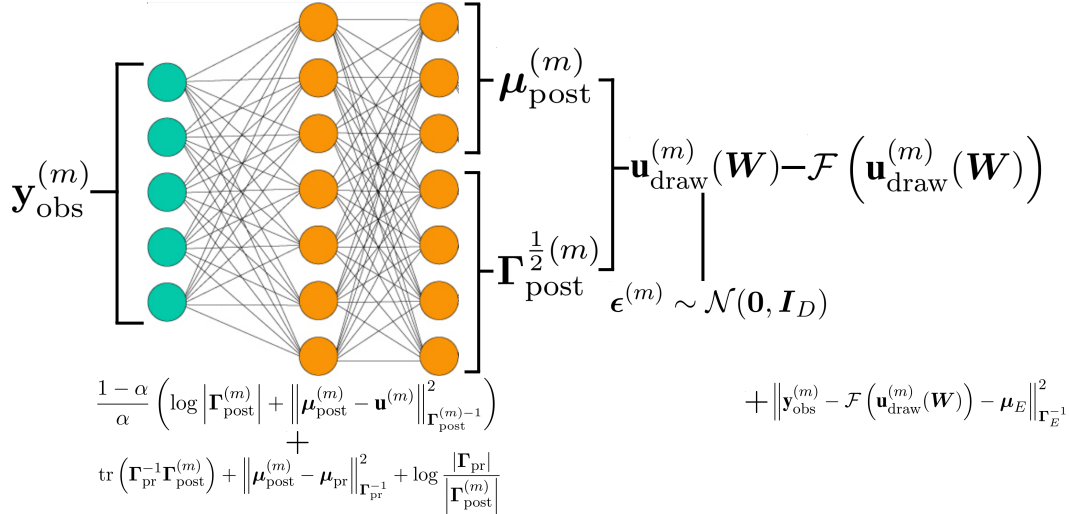


Figure 1: Schematic of the UQ-VAE framework with the exact PtO map.

During the training procedure, the repeated operation of the PtO map on our dataset $\left\{ \left(\mathbf{u}^{(m)}, \mathbf{y}_{\text{obs}}^{(m)} \right) \right\}_{m=1}^M$ may incur a significant computational cost. To alleviate this, we also consider an alternative for the PtO map in the optimization problem (21) where we simultaneously learn the PtO map along with the inverse problem solver. To this end, we replace the PtO map $\mathcal{F}$ with another neural network Ψ_d parameterized by weights $\mathbf{W}_d$. With this, the term (21b) becomes

$$\left\| \mathbf{y}_{\text{obs}}^{(m)} - \Psi_d \left(\mathbf{u}_{\text{draw}}^{(m)}(\mathbf{W}), \mathbf{W}_d \right) - \mu_E \right\|_{\Gamma_E^{-1}}^2 \quad (22)$$

and (21) becomes an optimization problem over both sets of weights $(\mathbf{W}, \mathbf{W}_d)$. Notice that the resulting optimization problem resembles those typically used to train an autoencoder. It is for this reason that we elect to use the subscript ‘d’ to indicate the decoder. Since the learned model Ψ_d of the PtO map $\mathcal{F}$ is a neural network, the evaluation of this map during the training phase is computationally inexpensive. However, this efficiency comes at the expense of utilizing knowledge of the governing physics of the inverse problems. Thus, the modified framework is almost purely data-driven and resembles more the original utility of VAEs for generative modelling. A schematic of the UQ-VAE framework with learned PtO map is displayed in Figure 2.

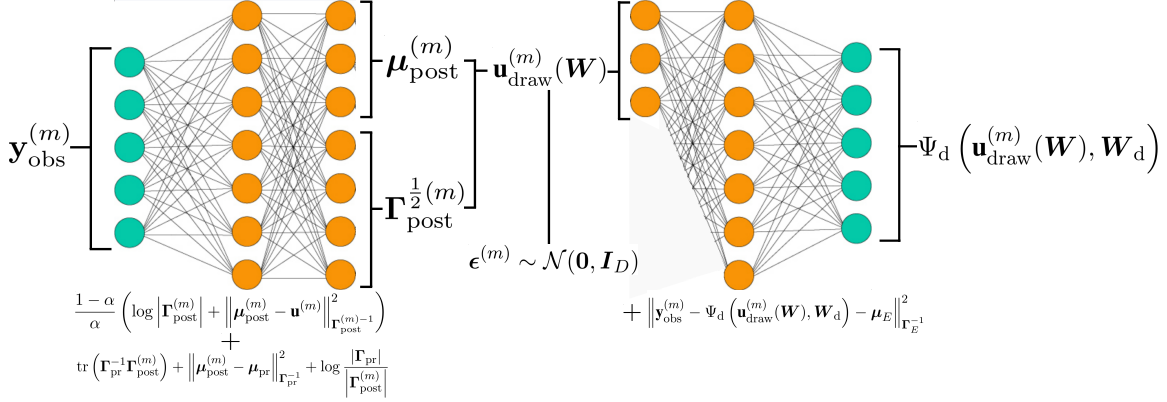


Figure 2: Schematic of the UQ-VAE framework with learned PtO map.

3.3. Theoretical Result

We now present a theoretical justification for the model-constrained UQ-VAE framework for linear PtO map $\mathcal{F}$. In this case, the target posterior $p(\mathbf{u}|\mathbf{y}_{\text{obs}})$ is Gaussian provided that the prior and noise are Gaussians. We show that for linear, one-layer neural networks the model-constrained UQ-VAE solution is exactly the target posterior. In particular, the UQ-VAE posterior mean and Cholesky factor of the UQ-VAE posterior covariance are exactly those of the target posterior. The proof is given in Section A.2 of the Appendix.

Theorem 3 *Let $\mathbf{y}_{\text{obs}} \in \mathbb{R}^O$. Consider a Gaussian prior model $\mathcal{N}(\boldsymbol{\mu}_{\text{pr}}, \boldsymbol{\Gamma}_{\text{pr}})$ and Gaussian noise model $\mathcal{N}(\boldsymbol{\mu}_E, \boldsymbol{\Gamma}_E)$. Suppose the PtO map is linear, i.e., $\mathcal{F}(\mathbf{u}) = \mathbf{F}\mathbf{u}$ with $\mathbf{F} \in \mathbb{R}^{O \times D}$, and consequently the target posterior $p(\mathbf{u}|\mathbf{y}_{\text{obs}})$ is a Gaussian $\mathcal{N}(\boldsymbol{\mu}_{\text{true}}, \boldsymbol{\Gamma}_{\text{true}})$ where*

$$\boldsymbol{\Gamma}_{\text{true}} = (\mathbf{F}^T \boldsymbol{\Gamma}_E^{-1} \mathbf{F} + \boldsymbol{\Gamma}_{\text{pr}}^{-1})^{-1} \quad (23a)$$

$$\boldsymbol{\mu}_{\text{true}} = \boldsymbol{\Gamma}_{\text{true}} (\mathbf{F}^T \boldsymbol{\Gamma}_E^{-1} (\mathbf{y}_{\text{obs}} - \boldsymbol{\mu}_E) + \boldsymbol{\Gamma}_{\text{pr}}^{-1} \boldsymbol{\mu}_{\text{pr}}) \quad (23b)$$

Suppose the model posterior $q_\phi(\mathbf{u}|\mathbf{y}_{\text{obs}}) = \mathcal{N}(\boldsymbol{\mu}_{\text{post}}, \boldsymbol{\Gamma}_{\text{post}})$ is such that

$$\boldsymbol{\mu}_{\text{post}} = \mathbf{W}_\mu \mathbf{y}_{\text{obs}} + \mathbf{b}_\mu \quad (24a)$$

$$\boldsymbol{\Gamma}_{\text{post}}^{\frac{1}{2}} = \mathbf{L} \odot \mathbf{L}_1 + \text{diag}(\boldsymbol{\sigma}) \quad (24b)$$

$$\log(\boldsymbol{\sigma}) = \mathbf{W}_\sigma \mathbf{y}_{\text{obs}} + \mathbf{b}_\sigma \quad (24c)$$

$$\text{vec}(\mathbf{L}) = \mathbf{W}_L \mathbf{y}_{\text{obs}} + \mathbf{b}_L \quad (24d)$$

where $\mathbf{L}_1$ is a lower triangular matrix of ones with zeros on the diagonal. Let $\alpha = \frac{1}{2}$. Then the optimization problem

$$\min_{\mathbf{W}_\mu, \mathbf{b}_\mu, \mathbf{W}_\sigma, \mathbf{b}_\sigma, \mathbf{W}_L, \mathbf{b}_L} \frac{1-\alpha}{\alpha} \left(\log |\mathbf{\Gamma}_{\text{post}}| + \text{tr}(\mathbf{\Gamma}_{\text{post}}^{-1} \mathbf{\Gamma}_{\text{true}}) + \|\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}}\|_{\mathbf{\Gamma}_{\text{post}}^{-1}}^2 \right) \quad (25a)$$

$$+ \text{tr}(\mathbf{\Gamma}_E^{-1} \mathbf{F} \mathbf{\Gamma}_{\text{post}} \mathbf{F}^T) + \|\mathbf{y}_{\text{obs}} - \mathbf{F} \boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_E\|_{\mathbf{\Gamma}_E^{-1}}^2 \quad (25b)$$

$$+ \text{tr}(\mathbf{\Gamma}_{\text{pr}}^{-1} \mathbf{\Gamma}_{\text{post}}) + \|\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{pr}}\|_{\mathbf{\Gamma}_{\text{pr}}^{-1}}^2 + \log \frac{|\mathbf{\Gamma}_{\text{pr}}|}{|\mathbf{\Gamma}_{\text{post}}|} \quad (25c)$$

achieves its minimum if and only if $\mathbf{W}_\mu, \mathbf{b}_\mu, \mathbf{W}_\sigma, \mathbf{b}_\sigma, \mathbf{W}_L, \mathbf{b}_L$ are such that $\boldsymbol{\mu}_{\text{post}} = \boldsymbol{\mu}_{\text{true}}$ and $\mathbf{\Gamma}_{\text{post}} = \mathbf{\Gamma}_{\text{true}}$.

Although this result provides theoretical guarantees, the requirement of learning a full matrix in (24d) can be computational infeasible. For this reason, it is often the case that only the diagonal of the model posterior covariance is learned Kingma and Welling (2013).

4. Related Work

In generative modelling, it is beneficial to learn better representations of the latent space. To this end, the use of the JSD for VAEs is explored in Deasy et al. (2020); Sutter et al. (2020). These works consider using the JSD in place of the KLD in the ELBO between the posterior and the prior. Our work differs from these approaches as the JSD is applied not directly to the ELBO. Rather, the JSD is applied at the inference level to yield a different bound that includes an expression of the latent data.

In scientific inverse problems, neural networks have been used to augment the solution process in Adler and Öktem (2017); Jin et al. (2017); Li et al. (2018); Patel and Oberai (2019); Peng et al. (2020); Seo et al. (2019). Some of these works use deep learning to assist a more traditional solution algorithm. In Adler and Öktem (2017), a convolutional neural network is used in a partially learned gradient descent scheme. In Jin et al. (2017), the component of the iterative solution algorithm requiring convolutional operators is surrogated with convolutional neural networks. Other approaches use deep learning directly for regularization of an ill-posed inverse problem. In Li et al. (2018), neural networks are used in place of a Tikhonov regularizer. The work in Patel and Oberai (2019); Gonzalez et al. (2019), explored prior modelling in Bayesian inverse problems using generative adversarial networks (GANs) and VAEs respectively. In Seo et al. (2019), VAEs are used to learn a low dimensional manifold representation of the PoI which introduces regularization for ill-posed inverse problems. Similarly, in Peng et al. (2020), an autoencoder is used to learn a latent representation of the unknown of interest which thereby allows the inversion task to be reframed into solving for the latent representation. In summary, even with improvements through deep learning, the computational cost of these approaches are still inherited by the cost of the base solution algorithms that were augmented. In our approach, the solution process is itself propagation through a trained neural network; a process largely unrivalled in terms of computational efficiency. Additionally, most of the work mentioned above do not quantify uncertainty.

Broadly speaking, there are three types of approaches for performing uncertainty quantification with deep learning. In an analogous manner to above, the first type utilizes deep learning to augment traditional methods for uncertainty quantification. In Jiang et al. (2019), an autoencoder

is trained to reduce the dimensionality of the PoI space. MCMC sampling is then conducted on the low-dimensional representation and estimation in the full PoI space is recovered by propagation through the trained decoder. Similarly, in [Hou et al. \(2019\)](#), a class of network training methods was proposed that can be combined with sample-based Bayesian inference algorithms. For the second type of approach, deep learning models that inherently possess some stochastic property are, themselves, used to represent uncertainty. In [Caldeira and Nord \(2020\)](#), an investigation into Bayesian Neural Networks, Concrete Dropout and Deep Ensembles was provided in the setting of scientific problems. However, the discussion in the paper does not directly pertain to inverse problems.

Our proposed method to uncertainty quantification falls under the third type of approach. Here, deep learning is used to directly learn and model the uncertainty. In [Gabbard et al. \(2019\)](#); [Tonolini et al. \(2020\)](#); [Chua and Vallisneri \(2020\)](#); [Adler and Öktem \(2018\)](#), purely deep learning posterior modelling approaches were proposed. In [Gabbard et al. \(2019\)](#) and [Tonolini et al. \(2020\)](#), conditional VAEs (CVAE) [Sohn et al. \(2015\)](#) were trained to perform variational inference for gravitational-wave astronomy problems and computational imaging problems respectively. Unlike our UQ-VAE approach, CVAEs do not marry the latent space with the PoI space; they require a separate latent variable for its derivation. The work in [Chua and Vallisneri \(2020\)](#) trains a neural network to rapidly produce one or two dimensional projections of Bayesian posteriors for gravitational-wave astronomy. For more general and higher-dimensional settings, the work in [Adler and Öktem \(2018\)](#) introduces Deep Posterior Sampling which trains a Wasserstein GAN to sample from the posterior. However, the training of GANs is notoriously unstable; a limitation also met and acknowledged by the authors of [Adler and Öktem \(2018\)](#) in their paper. Compared to GANs, VAEs possess a mathematical rigor that is more directly related to the data distribution it attempts to generate from. Therefore, with the VAE at the core of our proposed framework, we postulate that there is more potential for mathematically rigorous extension.

5. Results

For our numerical experiments, we consider the two dimensional steady state heat conduction problem. The details of this problem setup can be found in Section B of the Appendix. Although the non-linear forward mapping is quite simple, elliptic inverse problems are severely ill-posed and serve as an excellent benchmark inference problems [Kabanikhin \(2008\)](#); [Kirsch \(2011\)](#); [Beskos et al. \(2017\)](#); [Cui et al. \(2016\)](#); [Chen et al. \(2019\)](#). Details of the neural network architectures and optimization procedure can be found in Section C of the Appendix. The codes implementing the UQ-VAE framework used to produce these results can be obtained from [Goh \(2020\)](#). For comparison, we use the hIPPYlib library [Villa et al. \(2019\)](#) to compute the Laplace approximation using gradients and Hessians derived via the adjoint method, inexact Newton-CG, Armijo line search and randomized eigensolvers.

For our numerical experiments, we consider the following cases:

1. four noise levels $\eta = 0, 0.01, 0.05, 0.1$
2. four sizes of training datasets $M = 50, 500, 1000, 5000$
3. four choices of the JSD family $\alpha = 0.00001, 0.001, 0.1, 0.5$
4. exact PtO map and learned PtO map.

For our numerical simulations, in some cases, selecting $\alpha \geq 0.5$ may yield an unstable training procedure with exploding gradients and so we only display results with α as high as $\alpha = 0.5$. The computational cost of training is detailed in Section C.3 of the Appendix. The table of relative errors and figures are found in Section D of the Appendix with each figure set and subsection representing a noise level and dataset size. Note that we have elected to only display the cross-sectional predictions and not the predictions over the full domain since it is harder to visualize the uncertainty in the latter.

Before continuing to a detailed discussion, we address the comment in Section 3.2 regarding the adoption of a crude Monte-Carlo approximation using our PoI data. The overall feasibility of our obtained uncertainty estimates indicate that it is sufficient to use only one PoI datapoint. We postulate that any inaccuracies resulting from such a crude approximation is alleviated by the variety of information embedded in the training dataset. Indeed, since the weights are used to capture the statistics of a whole training dataset, the amortized nature of our inference method perhaps reduces the requirement for accurate Monte-Carlo approximations. Moreover, in the original VAE paper Kingma and Welling (2013), the authors observed that, providing the number of training data M was large enough, using only one draw from the model posterior to approximate the expected value of the likelihood term was sufficient to obtain good results. This finding regarding the approximation of the likelihood term mirrors ours regarding the posterior data misfit term.

5.1. Comparison with the Laplace Approximation

We summarize our results beginning with a comparison between uncertainty quantification with UQ-VAE and with the Laplace approximation. *It is important to point out that UQ-VAE provides a global Gaussian approximation while the Laplace approximation in this case is a local Gaussian approximation at the MAP point.* From the results, we can see that the posterior mean estimated by the UQ-VAE is closer to the true value than the MAP estimate. However, the Laplace approximation provides larger uncertainty estimates which ensures they are still feasible despite the underestimation of values by the MAP. In contrast, especially for larger values of α , the UQ-VAE uncertainty estimates are smaller and, in some cases, the true parameter values at the location of the anomaly lie slightly outside the uncertainty estimates.

Recall that the statistics of the noise are accurately modelled. Therefore, the larger the value of η , the more diminished the influence of the likelihood model (21b) on the training procedure; through penalization by the weighting matrix Γ_E^{-1} . For the Laplace approximation (12), this translates to an increase in the approximate posterior variance. The approximate posterior variance obtained from the UQ-VAE also exhibits this response to the increased noise regularization.

Further, the Laplace approximation uncertainty estimates are clearly responsive to the sensor locations; providing smaller uncertainty where the sensors are located. Whilst the UQ-VAE uncertainty estimates do exhibit a similar behaviour, it is less pronounced. The difference is due to the aforementioned global and local approximation of the UQ-VAE and Laplace approximations, respectively. Furthermore, we postulate that this is because the UQ-VAE uncertainty estimates represent two sources of information; the observational data fed into the encoder as well as the PoI and observational pairs in the training dataset. Whilst the former provides information only at the locations at the sensors, the latter provides information at all areas of the domain. Indeed, the process of training under the UQ-VAE framework works to embed information about the training dataset into the neural network. Uncertainty quantification is then performed with this embedded information as

well as the input observational data. Despite these differences, we believe that it is still instructive to compare the two approximations.

Finally, the computational efficiency of inference by propagation through a trained neural network Ψ is, on average, more than 2750 times faster than that of the Laplace approximation.

See Section C.3 of the Appendix for more details on the computational cost of inference. An example of our computational results is displayed in Figure 3 which compares the uncertainty estimates obtained using the Laplace approximation with those obtained with our UQ-VAE framework.

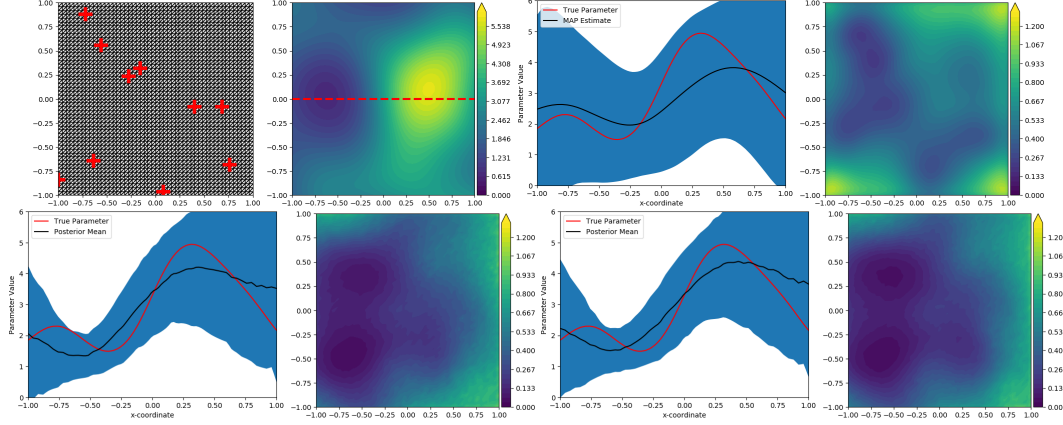


Figure 3: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second row: UQ-VAE uncertainty estimates with $\alpha = 0.001$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: exact PtO map. Third and fourth columns: learned PtO map.

5.2. Comparison Within the UQ-VAE Framework

Now, we discuss the results within the UQ-VAE framework by comparing between noise levels, dataset size, modelling and learning of the PtO map and choices of α :

1. Comparing between noise levels, feasible estimates were obtained for all noise levels. However, for $\eta = 0.05$, feasible estimates were only obtained for $M = 50, 500, 1000$. With the exception of $\eta = 0.05$ and $M = 5000$, the uncertainty estimates are larger for larger values of η . For the noise level of $\eta = 0.1$, no feasible estimates were obtained for any choice of dataset size M .
2. Comparing between training dataset sizes, the uncertainty estimates are larger for smaller dataset sizes. When $\eta = 0.05$, larger dataset sizes appear to be detrimental to the accuracy of the estimation.
3. Comparing between the case where the PtO is modelled and when the PtO is learned, for the cases of $\eta = 0$, $M = 500$, $\alpha = 0.5$ and $\eta = 0.01$, $M = 500$, $\alpha = 0.1$, using the modelled PtO clearly outperforms using the learned PtO.
4. Comparing between choices of α , for smaller values of α the estimated posterior mean underestimates the true PoI values and so a larger uncertainty estimate is required to ensure feasible estimates. The reverse occurs for larger values of α .

We now offer some key interpretations our above observations. For the first observation, the reasoning for larger uncertainty estimates observed with larger values of η was discussed in Section 5.1.

For the second observation, we see that the posterior uncertainty is responsive to the size of the dataset. With less data, our uncertainty estimates are larger. This behaviour is ideal as more information intuitively implies less uncertainty. For the case of $\eta = 0.05$, we conjecture that the use of large datasets of heavily corrupted data yields an unfavourable influence on our optimization procedure. We believe that the limitation of our framework lies in its potential dependence on an accurate training dataset. *For problem with exact PtO map such as the one in this work, the accurate training dataset is available and this corresponds to the case with zero noise $\eta = 0$.* One of the key merits of our proposed method is that we have introduced, in a mathematically justified manner, a data-driven component to the inversion process. This expands the toolkit at one's disposal when solving inverse problems. Indeed, in addition to constructing models for the prior and the PtO map, information on the properties of the PoI as well as the governing physics can be encoded into the training dataset. However, it would be incomplete to view this added option purely as an advantage. Although the utilization of datasets alleviates the burden on accurate prior and physics modelling, poorly constructed or highly corrupted datasets could completely sabotage the inversion process regardless of any accuracy achieved by the prior and physics models.

For the third observation, it can be reasoned that the similarity of results for the two cases is due to the success of the decoder in learning the PtO. The elliptic forward problem is quite simple and so it is not unexpected that the decoder learns it with ease. Thus, simultaneously learning an accurate PtO may be a reasonable task that aids and does not detract from the main task of learning the Bayesian inverse problems solver. This supports the strategy of using a learned PtO to reduce the offline cost of training the neural network. However, the two cases where use of the modelled PtO outperforms using the learned PtO with the small training dataset size of $M = 500$ suggests that including the physics of the problem allows for feasible uncertainty estimates when large amounts of training data may not be readily available.

For the fourth observation, in alignment with our discussion in Section 3.2, a lower value of α creates a larger penalty on the posterior terms (21a). Since our results suggest that a lower value of α yields a larger posterior variance, then it is implied that the minimization of $\log \left| \Gamma_{\text{post}}^{(m)} \right|$ is not a priority during the training procedure; so long as the posterior data $\mathbf{u}^{(m)}$ has enough influence. One can also reason that since lower values of α correspond to the zero-forcing KLD, one would expect a larger model variance; for example in the case of approximating a multimodal distribution with a unimodal model Bishop (2006). The reverse occurs with higher values of α which, due to tendency towards the zero-avoiding KLD, would result in a smaller model variance.

Finally, we address the instability of the training procedure as $\alpha \rightarrow 1$ which diminishes the influence of the PoI data on our optimization problem and moves towards recovering the zero-avoiding KLD. Recall the more traditional method of optimizing the exponentiated functional in (11) where the target optimization parameter is the physically meaningful $\mathbf{u}$. In contrast, our method optimizes for the more arbitrary and often more plentiful weights $\mathbf{W}$. Intuitively, we are therefore faced with a more ill-posed problem. Moreover, from a perspective centered around the traditional posterior (11), one can instead view the posterior data terms (21a) as regularizers of the traditional loss functional involving only the likelihood and prior term. This provides the much needed regularization for our

more ill-posed inverse problem; which therefore necessitates that $\alpha \ll 1$ for this regularization to take effect. We leave concrete theoretical analysis for future work.

6. Conclusion

In this paper, we propose a framework for training of neural networks capable of rapid uncertainty quantification. Although the VAE, on which this framework is based, was originally motivated by methods in generative modelling, we have shown that it is also well suited for scientific inverse problems. Indeed, it utilizes three sources of information usually available for such problems:

1. the physical laws governing the problem through the modelled PtO map
2. the physical properties of the PoI through the prior model
3. simulation or experimentation procedures through the paired PoI-observation datasets.

Furthermore, this framework is derived from a solid mathematical foundation and possesses a complex, dynamic interplay of factors from variational inference as well as regularization. Despite this complexity, this framework requires only a few design decisions. Indeed, aside from the usual adjustment of hyperparameters associated with neural network architecture, essentially only the tuning of the hyperparameter α is required. The selection of the PtO map and prior model can be guided by the underlying physical properties of the problem. Our results also show that feasible estimates are achievable and, moreover, that these estimates exhibit behaviour similar to that of existing inversion methods. Additionally, our uncertainty estimates show a favourable response to the size of our training dataset; the uncertainty is inversely proportional to the amount of training data used which reflects the availability or lack of information. The preliminary investigation offered in this paper employs only standard optimization techniques and simple deep learning architectures with a very minimal amount of tuning performed. This decision was made in order to showcase the robustness as well as the limitations of our framework in its most basic form. Moreover, we utilize somewhat crude statistical approximations. We believe that the results presented in this paper can be improved with the employment of more sophisticated optimization procedures, neural network architecture and statistical machinery.

Acknowledgments

This research was partially funded by the National Science Foundation awards NSF-1808576 and NSF-CAREER-1845799; by the Defense Threat Reduction Agency award DTRA-M1802962; by the Department of Energy award DE-SC0018147; by KAUST; by 2018 ConTex award; and by 2018 UT-Portugal CoLab award. The authors acknowledge the Texas Advanced Computing Center (TACC) at The University of Texas at Austin for providing HPC resources that have contributed to the research results reported within this paper. URL: <http://www.tacc.utexas.edu>. The authors would like to thank Jari Kaipio, Ruanui Nicholson and Rory Wittmer for the insightful discussions.

References

Jonas Adler and Ozan Öktem. Solving ill-posed inverse problems using iterative deep neural networks. *Inverse Problems*, 33(12):124007, 2017.

- Jonas Adler and Ozan Öktem. Deep bayesian inversion. *arXiv preprint arXiv:1811.05910*, 2018.
- Alexandros Beskos, Mark Girolami, Shiwei Lan, Patrick E Farrell, and Andrew M Stuart. Geometric mcmc for infinite-dimensional inverse problems. *Journal of Computational Physics*, 335: 327–351, 2017.
- Christopher M Bishop. *Pattern recognition and machine learning*. springer, 2006.
- David M Blei, Alp Kucukelbir, and Jon D McAuliffe. Variational inference: A review for statisticians. *Journal of the American statistical Association*, 112(518):859–877, 2017.
- Tan Bui-Thanh, Omar Ghattas, James Martin, and Georg Stadler. A computational framework for infinite-dimensional bayesian inverse problems part i: The linearized case, with application to global seismic inversion. *SIAM Journal on Scientific Computing*, 35(6):A2494–A2523, 2013.
- João Caldeira and Brian Nord. Deeply uncertain: comparing methods of uncertainty quantification in deep learning algorithms. *Machine Learning: Science and Technology*, 2(1):015002, dec 2020. doi: 10.1088/2632-2153/aba6f3. URL <https://doi.org/10.1088/2632-2153/aba6f3>.
- Peng Chen, Keyi Wu, Joshua Chen, Tom O’Leary-Roseberry, and Omar Ghattas. Projected stein variational newton: A fast and scalable bayesian inference method in high dimensions. In *Advances in Neural Information Processing Systems*, pages 15130–15139, 2019.
- Alvin JK Chua and Michele Vallisneri. Learning bayesian posteriors with neural networks for gravitational-wave inference. *Physical Review Letters*, 124(4):041102, 2020.
- Tiangang Cui, Kody JH Law, and Youssef M Marzouk. Dimension-independent likelihood-informed mcmc. *Journal of Computational Physics*, 304:109–137, 2016.
- Yair Daon and Georg Stadler. Mitigating the influence of the boundary on pde-based covariance operators. *arXiv preprint arXiv:1610.05280*, 2016.
- Masoumeh Dashti and Andrew M Stuart. The bayesian approach to inverse problems. *arXiv preprint arXiv:1302.6989*, 2013.
- Jacob Deasy, Nikola Simidjievski, and Pietro Liò. Constraining variational inference with geometric jensen-shannon divergence. *arXiv preprint arXiv:2006.10599*, 2020.
- Phoebus J Dhrymes. *Mathematics for econometrics*. Technical report, Springer, 1978.
- Michael Evans and Timothy Swartz. *Approximating integrals via Monte Carlo and deterministic methods*, volume 20. OUP Oxford, 2000.
- Hunter Gabbard, Chris Messenger, Ik Siong Heng, Francesco Tonolini, and Roderick Murray-Smith. Bayesian parameter estimation using conditional variational autoencoders for gravitational-wave astronomy. *arXiv preprint arXiv:1909.06296*, 2019.
- Hwan Goh. *uq-vae*, 2020. <https://github.com/hwangoh/uq-vae>.

- Mario Gonzalez, Andrés Almansa, Mauricio Delbracio, Pablo Musé, and Pauline Tan. Solving inverse problems by joint posterior maximization with a vae prior. *arXiv preprint arXiv:1911.06379*, 2019.
- Thomas Y Hou, Ka Chun Lam, Pengchuan Zhang, and Shumao Zhang. Solving bayesian inverse problems from the perspective of deep generative networks. *Computational Mechanics*, 64(2): 395–408, 2019.
- Zhenjiao Jiang, Siyu Zhang, Chris Turnadge, and Tianfu Xu. Combining autoencoder neural network and bayesian inversion algorithms to estimate heterogeneous fracture permeability in enhanced geothermal reservoirs. *essoar*, 2019.
- Kyong Hwan Jin, Michael T McCann, Emmanuel Froustey, and Michael Unser. Deep convolutional neural network for inverse problems in imaging. *IEEE Transactions on Image Processing*, 26(9): 4509–4522, 2017.
- Sergei Igorevich Kabanikhin. Definitions and examples of inverse and ill-posed problems. *Journal of Inverse and Ill-Posed Problems*, 16(4):317–357, 2008.
- Jari Kaipio and Erkki Somersalo. *Statistical and computational inverse problems*, volume 160. Springer Science & Business Media, 2006.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Diederik P Kingma and Max Welling. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- Andreas Kirsch. *An introduction to the mathematical theory of inverse problems*, volume 120. Springer Science & Business Media, 2011.
- Housen Li, Johannes Schwab, Stephan Antholzer, and Markus Haltmeier. Nett: Solving inverse problems with deep neural networks. *arXiv preprint arXiv:1803.00092*, 2018.
- Anders Logg, Kent-Andre Mardal, and Garth Wells. *Automated solution of differential equations by the finite element method: The FEniCS book*, volume 84. Springer Science & Business Media, 2012.
- Kevin P Murphy. *Machine learning: a probabilistic perspective*. MIT press, 2012.
- Frank Nielsen. A family of statistical symmetric divergences based on jensen’s inequality. *arXiv preprint arXiv:1009.4004*, 2010.
- Dhruv Patel and Assad A Oberai. Bayesian inference with generative adversarial network priors. *arXiv preprint arXiv:1907.09987*, 2019.
- Pei Peng, Shirin Jalali, and Xin Yuan. Solving inverse problems via auto-encoders. *IEEE Journal on Selected Areas in Information Theory*, 1(1):312–323, 2020.
- S James Press. *Subjective and objective Bayesian statistics: Principles, models, and applications*, volume 590. John Wiley & Sons, 2009.

- Lassi Roininen, Janne MJ Huttunen, and Sari Lasanen. Whittle-matérn priors for bayesian statistical inversion with applications in electrical impedance tomography. *Inverse Problems & Imaging*, 8(2):561, 2014.
- Jin Keun Seo, Kang Cheol Kim, Ariungerel Jargal, Kyounghun Lee, and Bastian Harrach. A learning-based method for solving ill-posed nonlinear inverse problems: A simulation study of lung eit. *SIAM journal on Imaging Sciences*, 12(3):1275–1295, 2019.
- Kihyuk Sohn, Honglak Lee, and Xinchun Yan. Learning structured output representation using deep conditional generative models. In *Advances in neural information processing systems*, pages 3483–3491, 2015.
- Stephen M Stigler. Laplace’s 1774 memoir on inverse probability. *Statistical Science*, pages 359–363, 1986.
- Andrew M Stuart. Inverse problems: a bayesian perspective. *Acta numerica*, 19:451, 2010.
- Thomas M Sutter, Imant Daunhawer, and Julia E Vogt. Multimodal generative learning utilizing jensen-shannon-divergence. *arXiv preprint arXiv:2006.08242*, 2020.
- Luke Tierney and Joseph B Kadane. Accurate approximations for posterior moments and marginal densities. *Journal of the american statistical association*, 81(393):82–86, 1986.
- Francesco Tonolini, Jack Radford, Alex Turpin, Daniele Faccio, and Roderick Murray-Smith. Variational inference for computational imaging inverse problems. *Journal of Machine Learning Research*, 21(179):1–46, 2020.
- Umberto Villa, Noemi Petra, and Omar Ghattas. hIPPYlib: An Extensible Software Framework for Large-Scale Inverse Problems Governed by PDEs; Part I: Deterministic Inversion and Linearized Bayesian Inference. *arXiv e-prints*, 2019.
- Roderick Wong. *Asymptotic approximations of integrals*. SIAM, 2001.

Appendix A. Proofs

A.1. Proof of Theorem 1

Proof The clearest path forward is to manipulate

$$\text{JS}_\alpha(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) = \alpha \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})) \quad (26a)$$

$$+ (1-\alpha) \text{KL}(p(\mathbf{u}|\mathbf{y})||(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})) \quad (26b)$$

term-by-term. Beginning with the first term,

$$\alpha \text{KL}(q_\phi(\mathbf{u}|\mathbf{y})||(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})) \quad (27a)$$

$$= \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{q_\phi(\mathbf{u}|\mathbf{y})}{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (27b)$$

$$= -\alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})}{q_\phi(\mathbf{u}|\mathbf{y})} \right) \right] \quad (27c)$$

$$= -\alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{u}|\mathbf{y})}{q_\phi(\mathbf{u}|\mathbf{y})} \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right) \right] \quad (27d)$$

$$= -\alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{u}|\mathbf{y})}{q_\phi(\mathbf{u}|\mathbf{y})} \right) \right] - \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (27e)$$

$$= -\alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{y}, \mathbf{u})}{p(\mathbf{y})q_\phi(\mathbf{u}|\mathbf{y})} \right) \right] - \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (27f)$$

$$= \alpha \log(p(\mathbf{y})) - \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{y}, \mathbf{u})}{q_\phi(\mathbf{u}|\mathbf{y})} \right) \right] - \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right]. \quad (27g)$$

Similarly, the second term can be decomposed as

$$(1-\alpha) \text{KL}(p(\mathbf{u}|\mathbf{y})||(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})) \quad (28a)$$

$$= (1-\alpha) \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y})} \left[\log \left(\frac{p(\mathbf{u}|\mathbf{y})}{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (28b)$$

$$= -(1-\alpha) \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y})} \left[\log \left(\frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y}) + \alpha p(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (28c)$$

$$= -(1-\alpha) \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y})} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right]. \quad (28d)$$

Combining equations (27) and (28), we arrive at

$$\text{JS}_\alpha(q_\phi(\mathbf{u}|\mathbf{y})||p(\mathbf{u}|\mathbf{y})) = \alpha \log(p(\mathbf{y})) \quad (29a)$$

$$- \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\frac{p(\mathbf{y}, \mathbf{u})}{q_\phi(\mathbf{u}|\mathbf{y})} \right) \right] \quad (29b)$$

$$- \alpha \mathbb{E}_{\mathbf{u} \sim q_\phi} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right] \quad (29c)$$

$$- (1-\alpha) \mathbb{E}_{\mathbf{u} \sim p(\mathbf{u}|\mathbf{y})} \left[\log \left(\alpha + \frac{(1-\alpha)q_\phi(\mathbf{u}|\mathbf{y})}{p(\mathbf{u}|\mathbf{y})} \right) \right]. \quad (29d)$$

Finally, the expressions in Theorem 1 follow through division by α and rearrangement. ■

A.2. Proof of Theorem 3

Proof Set $\alpha = \frac{1}{2}$. Forcing the variations with respect to weights $\mathbf{W}_\mu$ and biases $\mathbf{b}_\mu$ to vanish gives

$$\mathbf{\Gamma}_{\text{post}}^{-1} (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}}) \quad (30a)$$

$$- \mathbf{F}^T \mathbf{\Gamma}_E^{-1} (\mathbf{y}_{\text{obs}} - \mathbf{F} \boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_E) + \mathbf{\Gamma}_{\text{pr}}^{-1} (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{pr}}) \quad (30b)$$

$$= \mathbf{0} \quad (30c)$$

which is true if and only if $\boldsymbol{\mu}_{\text{post}} = \boldsymbol{\mu}_{\text{true}}$. Forcing the variations with respect to $\mathbf{W}_L$ and biases $\mathbf{b}_L$ to vanish gives

$$- \left(\mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}} \mathbf{\Gamma}_{\text{true}} \otimes \mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}T} + \mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}} (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}}) (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}})^T \otimes \mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}T} \right) \text{vec} \left(\mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}} \right) \quad (31a)$$

$$+ (\mathbf{I}_D \otimes \mathbf{F}^T \mathbf{\Gamma}_E^{-1} \mathbf{F}) \text{vec} \left(\mathbf{\Gamma}_{\text{post}}^{\frac{1}{2}} \right) + (\mathbf{I}_D \otimes \mathbf{\Gamma}_{\text{pr}}^{-1}) \text{vec} \left(\mathbf{\Gamma}_{\text{post}}^{\frac{1}{2}} \right) \quad (31b)$$

$$= \mathbf{0} \quad (31c)$$

where $\otimes$ denotes the Kronecker product. Using the two identities [Dhrymes \(1978\)](#)

$$\text{vec}(\mathbf{ABCD}) = (\mathbf{D}^T \mathbf{C}^T \otimes \mathbf{A}) \text{vec}(\mathbf{B}) \quad (32a)$$

$$\text{vec}(\mathbf{AB}) = (\mathbf{I} \otimes \mathbf{A}) \text{vec}(\mathbf{B}) \quad (32b)$$

along with some basic manipulation gives

$$- \mathbf{\Gamma}_{\text{post}}^{-1} \left(\mathbf{\Gamma}_{\text{true}} + (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}}) (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}})^T \right) + (\mathbf{F}^T \mathbf{\Gamma}_E^{-1} \mathbf{F} + \mathbf{\Gamma}_{\text{pr}}^{-1}) \mathbf{\Gamma}_{\text{post}} = \mathbf{0}. \quad (33)$$

Forcing the variations with respect to $\mathbf{W}_\sigma$ and biases $\mathbf{b}_\sigma$ to vanish gives

$$- \mathbf{\Gamma}_{\text{post}}^{-1} \odot \left(\mathbf{\Gamma}_{\text{true}} + (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}}) (\boldsymbol{\mu}_{\text{post}} - \boldsymbol{\mu}_{\text{true}})^T \right) \mathbf{\Gamma}_{\text{post}}^{-\frac{1}{2}T} \quad (34a)$$

$$+ \left(\mathbf{F}^T \mathbf{\Gamma}_E^{-1} \mathbf{F} \odot \mathbf{\Gamma}_{\text{post}}^{\frac{1}{2}} \right) + \left(\mathbf{\Gamma}_{\text{pr}}^{-1} \odot \mathbf{\Gamma}_{\text{post}}^{\frac{1}{2}} \right) \quad (34b)$$

$$= \mathbf{0} \quad (34c)$$

where $\odot$ denotes the entrywise product. Here, (33) and (34) with $\boldsymbol{\mu}_{\text{post}} = \boldsymbol{\mu}_{\text{true}}$ are true if and only if $\mathbf{\Gamma}_{\text{post}} = \mathbf{\Gamma}_{\text{true}}$ as required. $\blacksquare$

Appendix B. Two Dimensional Steady State Heat Conduction Problem

We considered the heat equation with heat conductivity as the PoI and temperature as the state. The governing PDE and associated boundary conditions are displayed below

$$-\nabla \cdot u \nabla y = 0 \quad \text{in } \Omega \quad (35a)$$

$$-u(\nabla y \cdot \hat{\mathbf{n}}) = \text{Bi } y \quad \text{on } \Omega^{\text{ext}} \setminus \Omega^{\text{root}} \quad (35b)$$

$$-u(\nabla y \cdot \hat{\mathbf{n}}) = -1 \quad \text{on } \Omega^{\text{root}} \quad (35c)$$

where u denotes the thermal heat conductivity, Bi is the Biot number set to $\text{Bi} = 0.5$, Ω is the physical domain, Ω^{root} is the bottom edge of the domain and Ω^{ext} is the remaining edges of the domain.

The prior model has mean $\mu_{\text{pr}} = 2$ and covariance Γ_{pr} that is a discretization of the infinite dimensional covariance operator $\mathcal{C}_{\text{pr}} = \mathcal{A}^{-2}$ where $\mathcal{A}$ is a differential operator such that

$$\mathcal{A}u = \begin{cases} -\gamma\Delta u + \delta u & \text{in } \Omega \\ \nabla u \cdot \hat{\mathbf{n}} + \beta u & \text{on } \partial\Omega. \end{cases} \quad (36)$$

Here, $\delta, \gamma > 0$ controls the correlation length and variance of the prior operator. We set $\gamma = 0.1$, $\delta = 0.5$ and β is chosen as in [Daon and Stadler \(2016\)](#); [Roininen et al. \(2014\)](#) to reduce boundary artifacts. Priors of this type ensure that $\mathcal{C}_{\text{pr}}$ is a trace-class operator which guarantees bounded pointwise variance and a well-posed infinite-dimensional Bayesian inverse problem [Stuart \(2010\)](#); [Bui-Thanh et al. \(2013\)](#).

The testing dataset of PoI values is drawn from a Gaussian autocorrelation smoothness prior [Kaipio and Somersalo \(2006\)](#) with mean $\mu = 2$ and covariance

$$\Gamma_{ij} = \sigma^2 \exp\left(-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|_2^2}{2\rho^2}\right) \quad (37)$$

with $\sigma^2 = 2$, $\rho = 0.5$. The training set for the UQ-VAE network is a set of separate draws from the same distribution. From this, the corresponding state observations are computed using the solvers from the FEniCS library [Logg et al. \(2012\)](#) on the PoI dataset. The same solver is used as our numerical model of the PtO map during the training procedure.

We consider a computational domain consisting of $D = 2601$ degrees of freedom. The observation data corresponds to sensor measurements from 10 randomly selected locations and is afflicted with Gaussian distributed additive noise on all sensors with zero mean and standard deviation $\sigma = \eta \max |\mathbf{y}_{\text{obs}}|$. For our noise model, we set $(\mu_E, \Gamma_E) = (\mathbf{0}, \sigma^2 \mathbf{I}_O)$ which corresponds to the scenario where the statistics of the noise are accurately modelled. Finally, we assume a diagonal matrix for our posterior model covariance.

Appendix C. Neural Network Architecture and Training Properties

C.1. Architecture

The architecture of our neural network Ψ consists of 5 hidden layers of 500 nodes with the ReLU activation function. No activation function was used at the output layer. The input layer has O number of nodes to match the dimension of the observational data which comprises of O number of measurement points. The output layer has $2D$ nodes, with D the dimension of the PoI, to represent the estimated posterior mean μ_{post} and diagonal of the posterior covariance Γ_{post} .

When the parameter-to-observable map is learnt, the corresponding decoder network Ψ_d has 2 hidden layers also with 500 nodes and the ReLU activation function. Again, no activation function was used at the output layer. The input layer has D number of nodes to represent a draw from the learned posterior and the output layer has O number of nodes to match the dimension of the observational data.

C.2. Training

For optimization, we use a batch size of 100. Therefore, the loss in (21) is averaged over the number of PoI and observation pairs and the gradient is calculated for each batch. Optimization is conducted using the Adam optimizer which performs mini-batch momentum-based stochastic gradient descent Kingma and Ba (2014). This training procedure was repeated for 400 epochs.

The metric used to measure accuracy is the averaged relative error $\frac{1}{L} \sum_{l=1}^L \frac{\|\mathbf{u}^{(l)} - \boldsymbol{\mu}_{\text{post}}(\mathbf{y}_{\text{obs}}^{(l)})\|_2^2}{\|\mathbf{u}^{(l)}\|_2^2}$ where $\left\{ \left(\mathbf{u}^{(l)}, \mathbf{y}_{\text{obs}}^{(l)} \right) \right\}_{l=1}^L$ is a dataset unseen by the training procedure and $\boldsymbol{\mu}_{\text{post}}(\mathbf{y}_{\text{obs}}^{(l)})$ is the estimated posterior mean from the neural network Ψ taking a datapoint $\mathbf{y}_{\text{obs}}^{(l)}$ as an input. An estimate is said to be *feasible* if and only if true value lies within the estimated uncertainty bounds. The uncertainty bounds displayed represent three standard deviations. Similarly, when the PtO map is learned by Ψ_d , we use the relative error $\frac{1}{L} \sum_{l=1}^L \frac{\|\mathbf{y}_{\text{obs}}^{(l)} - \Psi_d(\mathbf{u}(\mathbf{y}_{\text{obs}}^{(l)}))\|_2^2}{\|\mathbf{y}_{\text{obs}}^{(l)}\|_2^2}$ where $\mathbf{u}(\mathbf{y}_{\text{obs}}^{(l)})$ is a draw model posterior output by the trained encoder with a datapoint $\mathbf{y}_{\text{obs}}^{(l)}$ as an input.

C.3. Computational Cost

When training a neural network with the (21b) term in the loss functional (21) using the PtO map $\mathcal{F}$, the solving of a linear system is required. Because of this, training on CPUs is more efficient. In contrast, when the (21b) term uses the learned PtO map Ψ_d instead, no linear system solves are required and so the training procedure is most efficient on GPUs. We monitor the cost of training each batch as wall-clock time for the whole optimization procedure. The offline cost for training the networks using $\mathcal{F}$ is, on average, 4.5 seconds per batch on a dual-socket node with two Intel Xeon E5-2690 CPUs for a total of 24 cores. The offline cost for training the networks using Ψ_d is, on average, 0.35 seconds per batch on a NVIDIA 1080-TI GPU.

We compare the computational efficiency of inference between the Laplace approximation and propagation through the neural network Ψ trained under the UQ-VAE framework on a Intel Core i7-8550U CPU using an average wall-clock time over 20 evaluations. The time taken to form the MAP estimate with hIPPYlib is, on average, 40 seconds with a maximum number of descent iterations of 25. The time taken to form the low-rank Gaussian approximation of the posterior covariance with 50 requested eigenvectors is, on average, 70 seconds. Therefore, in total, it takes on average 110 seconds to form the Laplace approximation. In contrast, forming the model posterior by propagation through a trained neural network takes, on average, 0.04 seconds; more than 2750 times faster.

Appendix D. Results: Two Dimensional Steady State Heat Conduction Problem

D.1. $\eta = 0, M = 50$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	28.88%	30.86%	24.61%
0.001	30.46%	30.76%	22.80%
0.1	30.70%	28.95%	49.38%
0.5	29.75%	32.73%	45.34%

Table 1: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 15.15%.

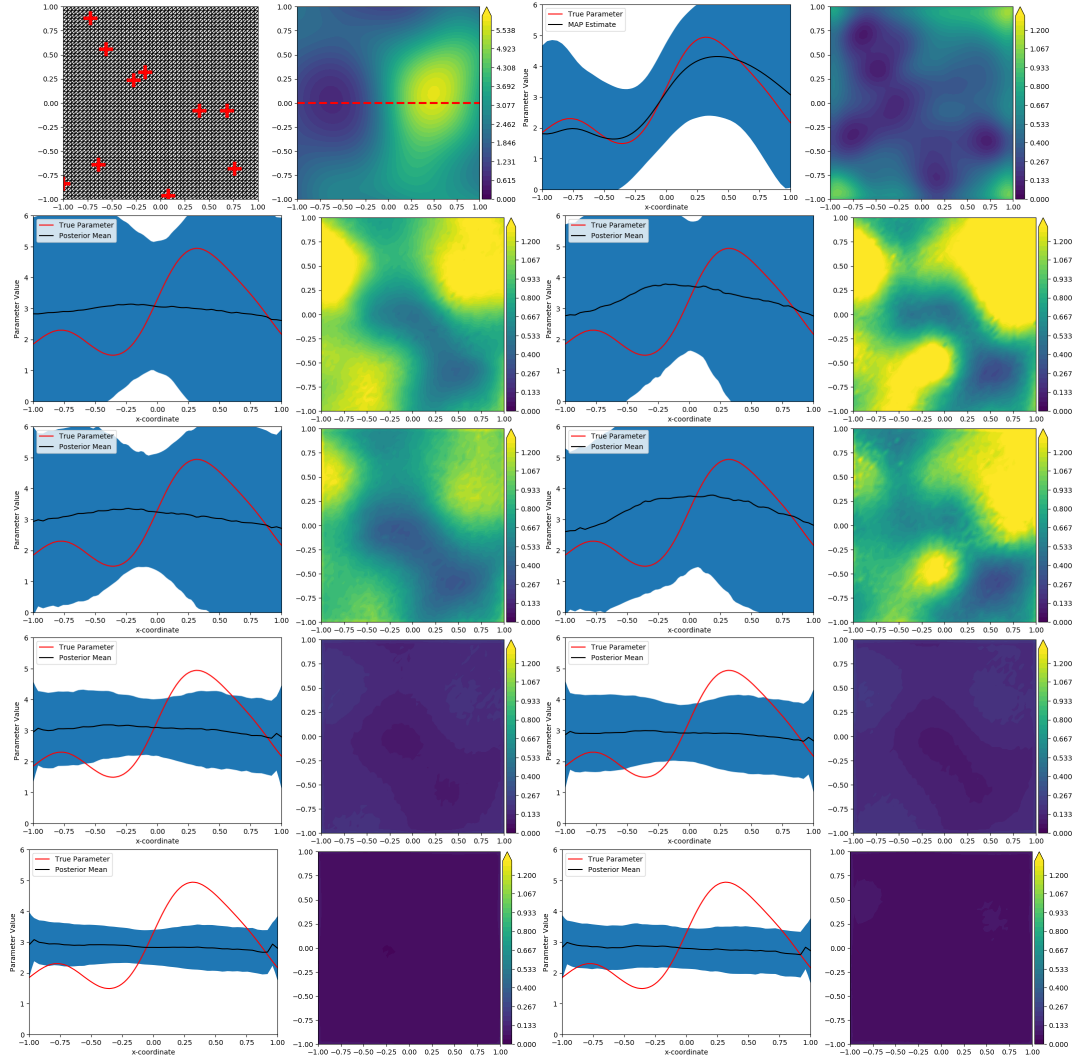
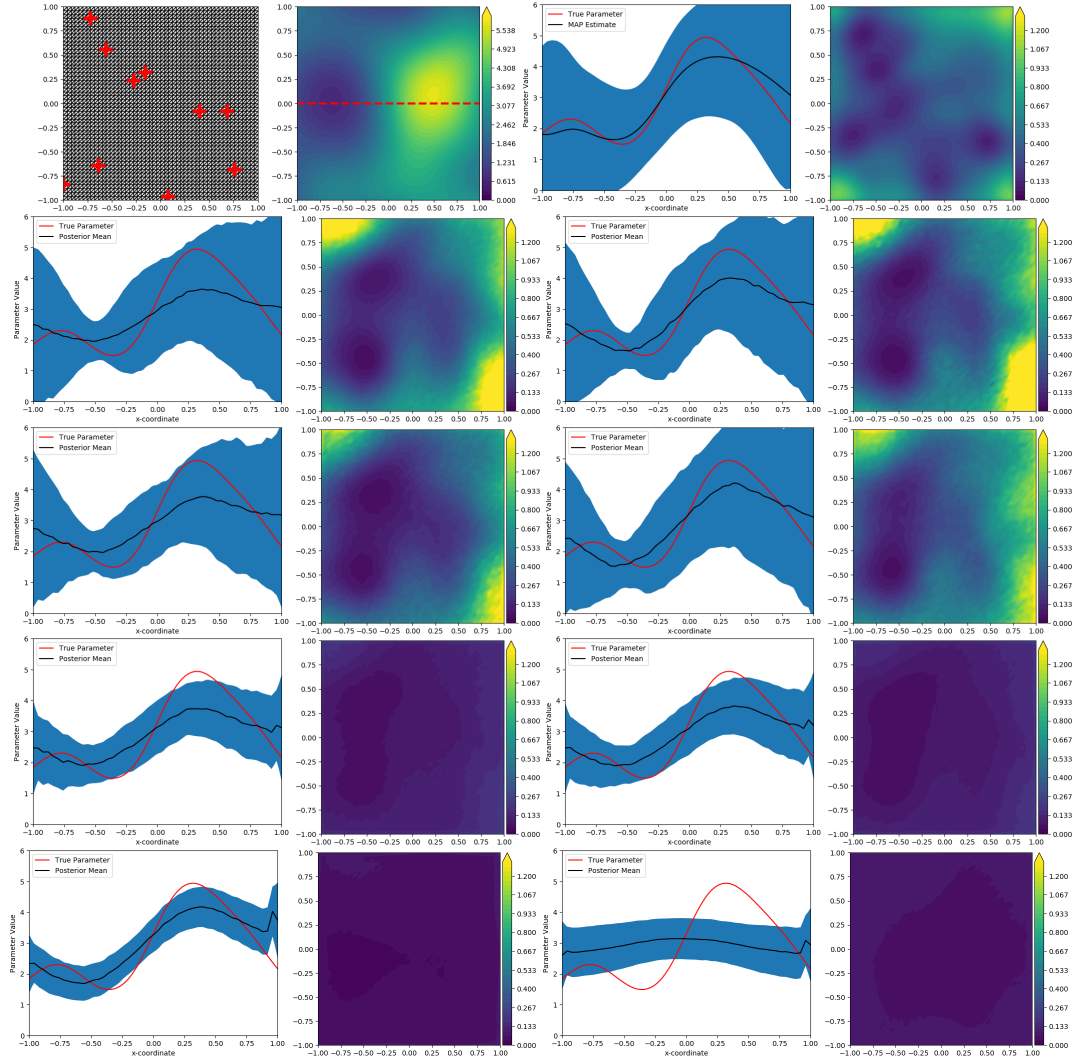


Figure 4: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.2. $\eta = 0, M = 500$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	21.72%	21.70%	22.49%
0.001	21.33%	22.13%	20.75%
0.1	21.84%	21.91%	24.08%
0.5	26.75%	29.47%	35.26%

Table 2: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 15.15%.


 Figure 5: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.3. $\eta = 0, M = 1000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	20.65%	18.15%	18.55%
0.001	18.54%	19.04%	15.02%
0.1	18.74%	19.28%	22.72%
0.5	21.84%	22.74%	21.74%

Table 3: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 15.15%.

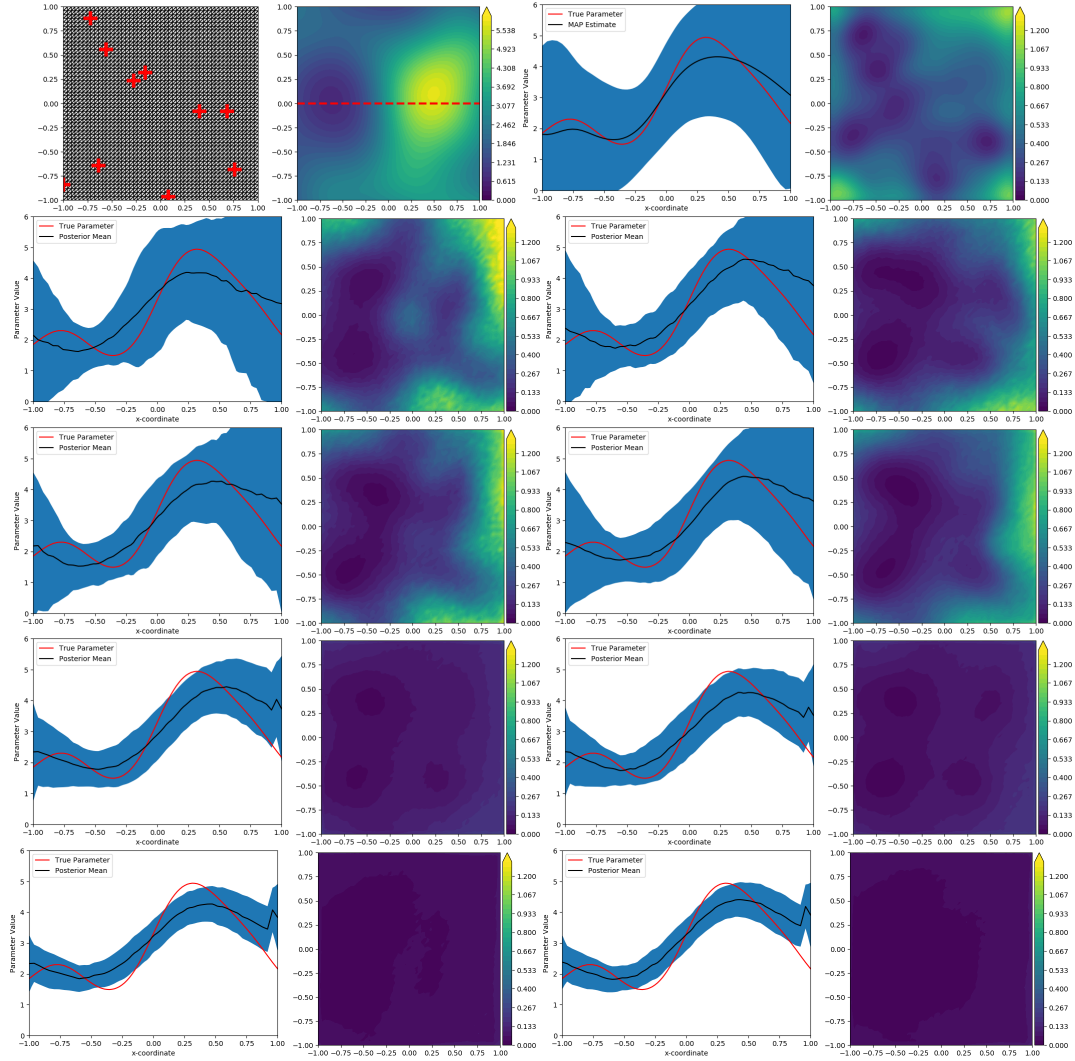


Figure 6: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.4. $\eta = 0, M = 5000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	14.49%	15.07%	3.21%
0.001	14.28%	14.07%	3.13%
0.1	14.72%	14.48%	3.19%
0.5	14.89%	15.47%	4.72%

Table 4: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 15.15%.

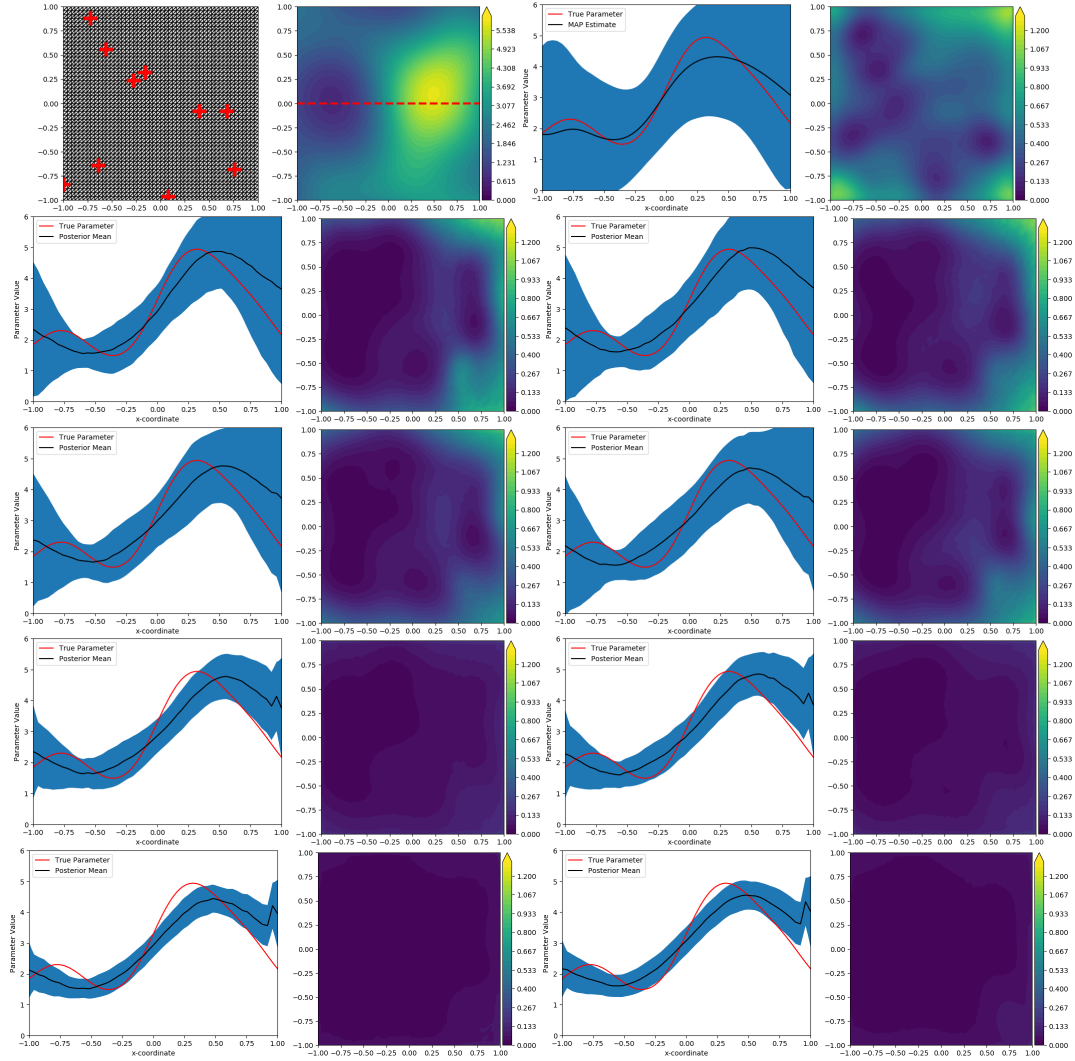


Figure 7: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.5. $\eta = 0.01, M = 50$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	29.19%	30.67%	24.59%
0.001	39.52%	29.25%	31.88%
0.1	28.99%	28.95%	34.97%
0.5	32.92%	30.05%	41.18%

Table 5: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 25.11%.

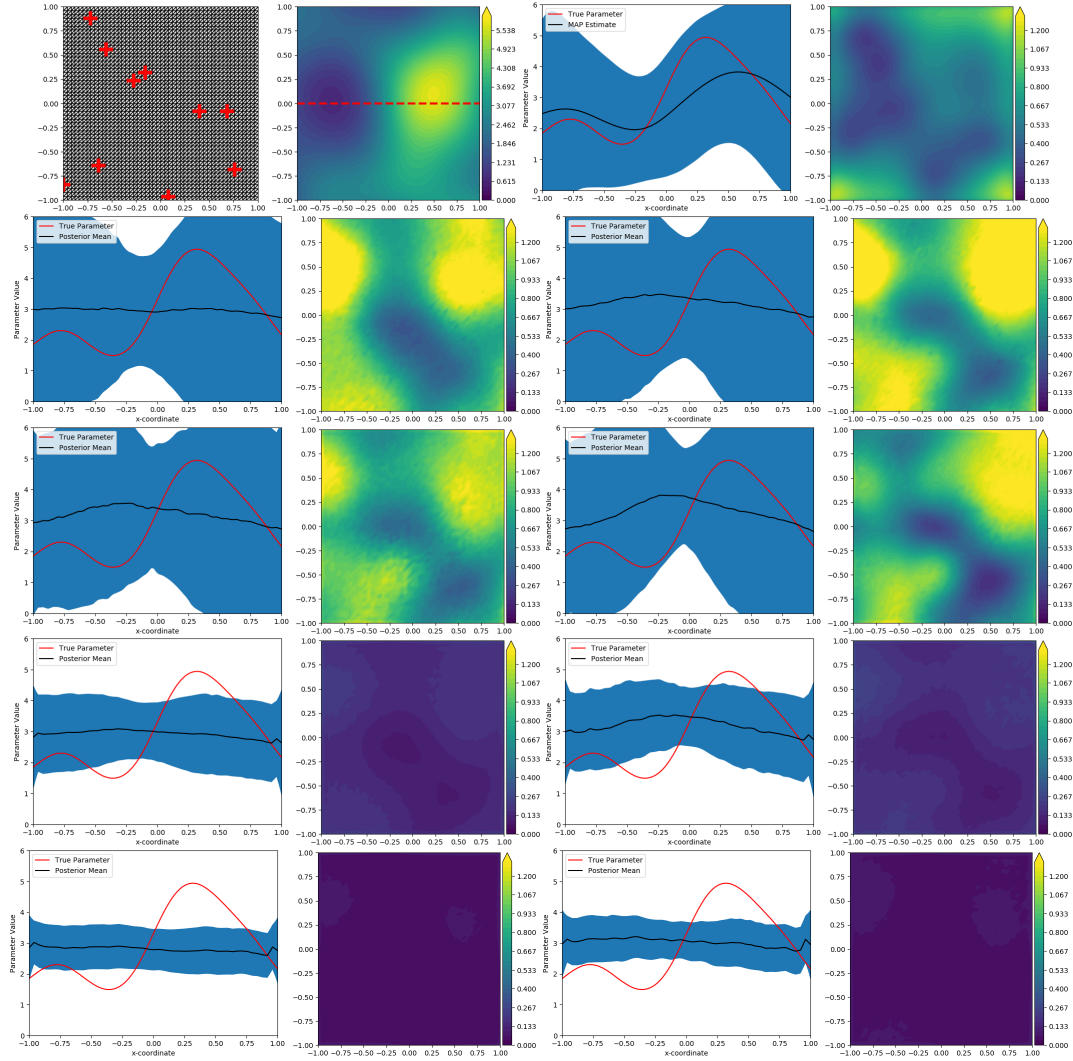


Figure 8: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.6. $\eta = 0.01, M = 500$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	24.17%	23.42%	22.02%
0.001	22.79%	23.23%	17.41%
0.1	25.51%	29.52%	28.16%
0.5	29.93%	29.82%	32.68%

Table 6: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 25.11%.

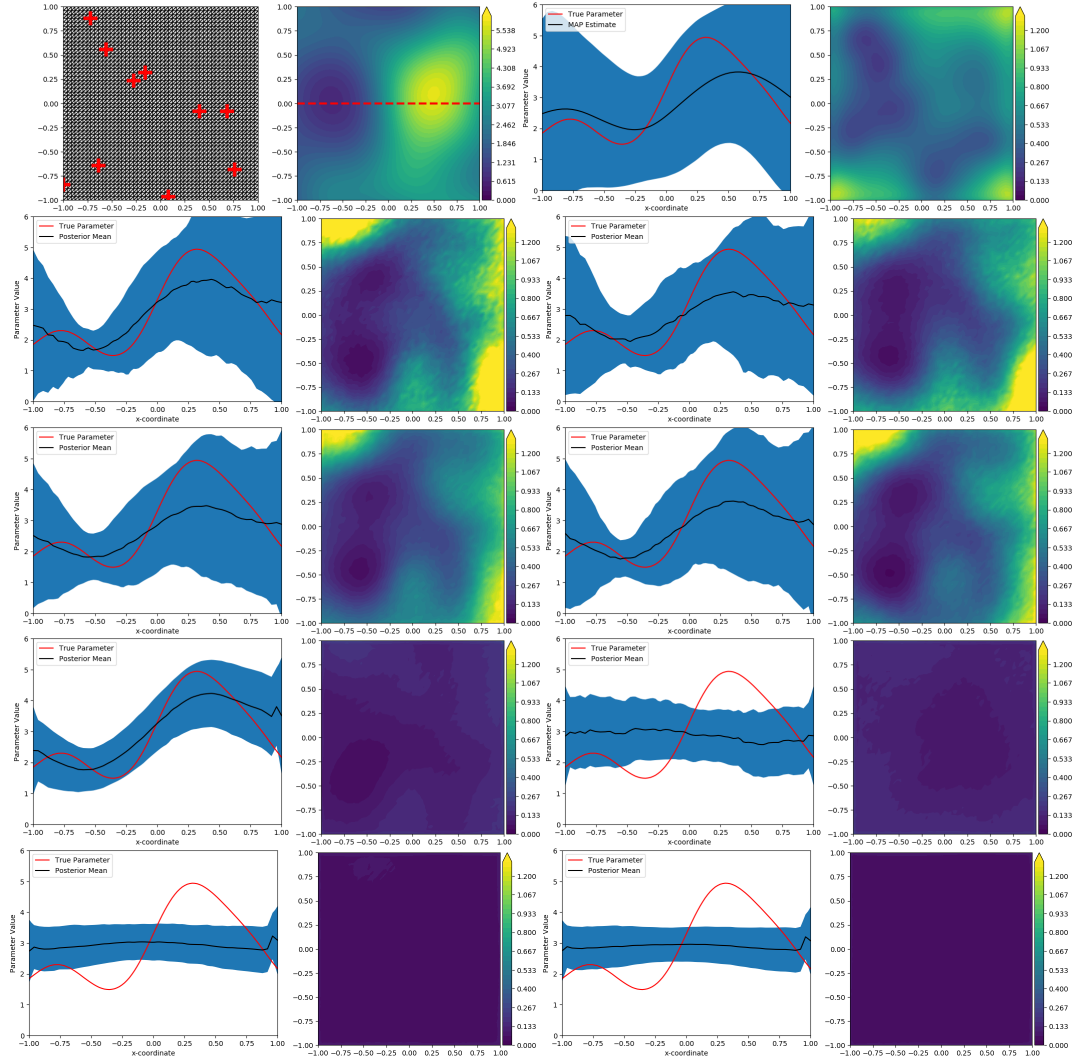


Figure 9: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.7. $\eta = 0.01, M = 1000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	21.99%	23.64%	19.98%
0.001	22.24%	21.78%	14.42%
0.1	22.84%	22.73%	20.28%
0.5	24.16%	24.04%	22.78%

Table 7: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 25.11%.

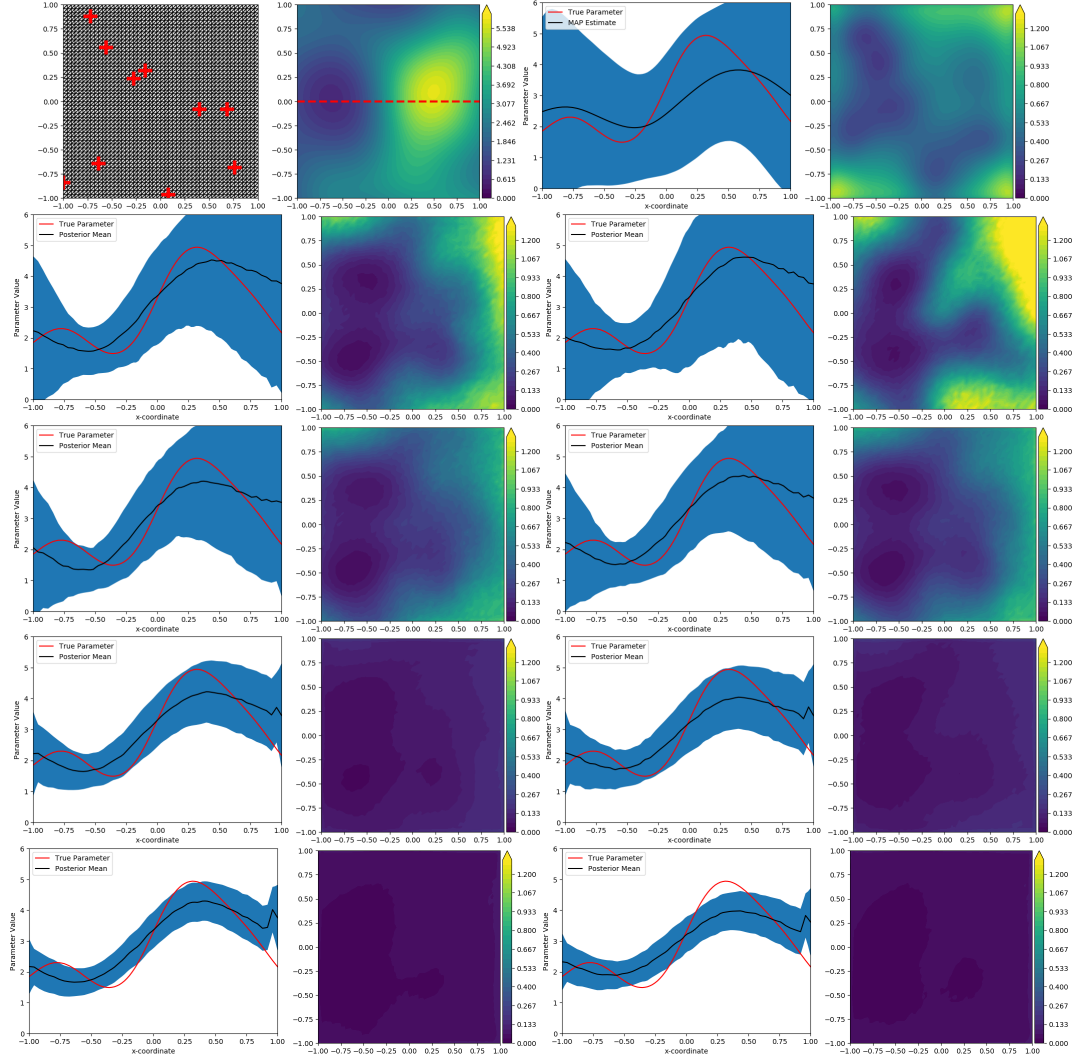


Figure 10: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.8. $\eta = 0.01, M = 5000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	21.28%	21.14%	4.29%
0.001	21.34%	20.92%	3.79%
0.1	21.37%	21.47%	5.27%
0.5	21.72%	21.19%	3.57%

Table 8: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 25.11%.

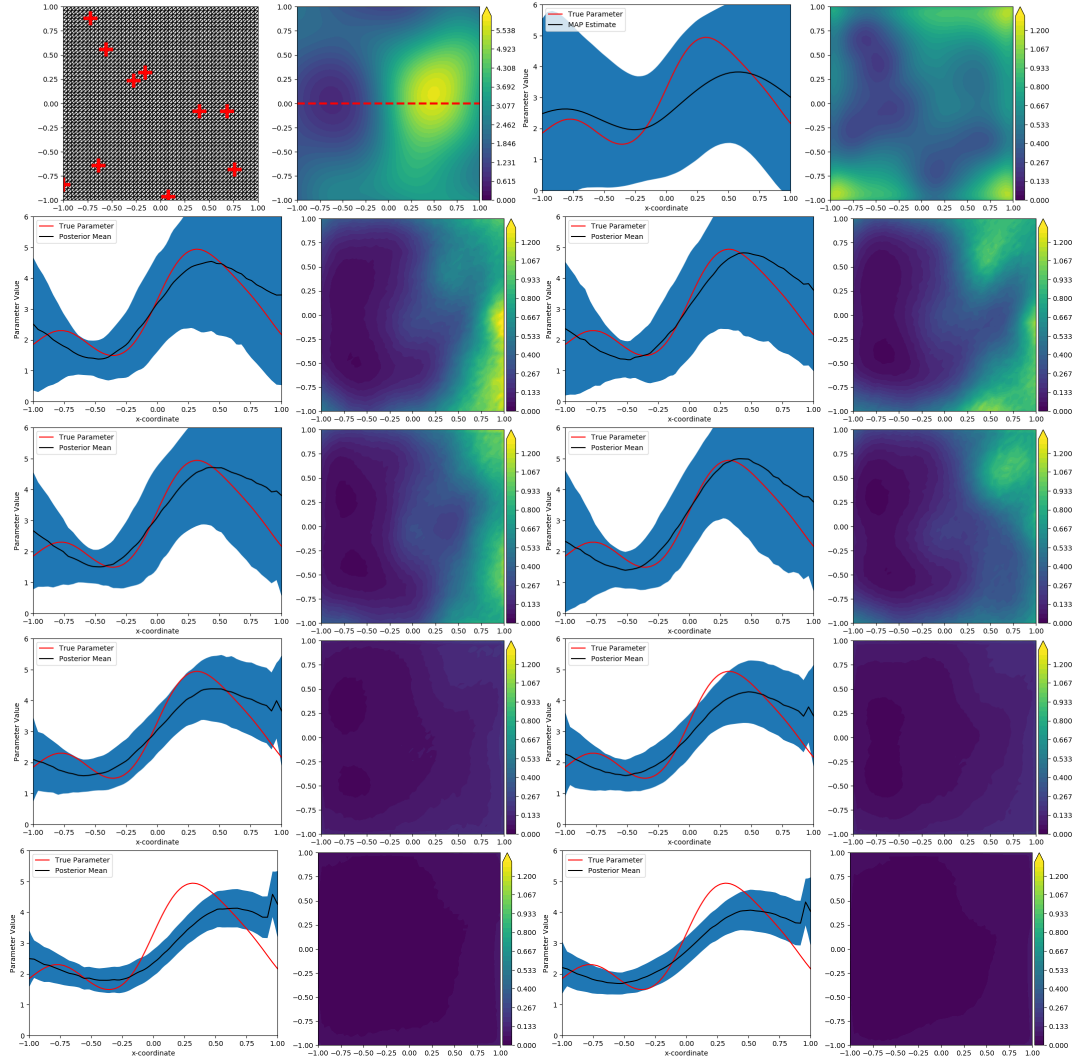


Figure 11: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.9. $\eta = 0.05, M = 50$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	31.76%	33.37%	22.32%
0.001	34.78%	34.52%	22.24%
0.1	33.45%	31.00%	39.73%
0.5	32.87%	33.00%	38.23%

Table 9: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 43.19%.

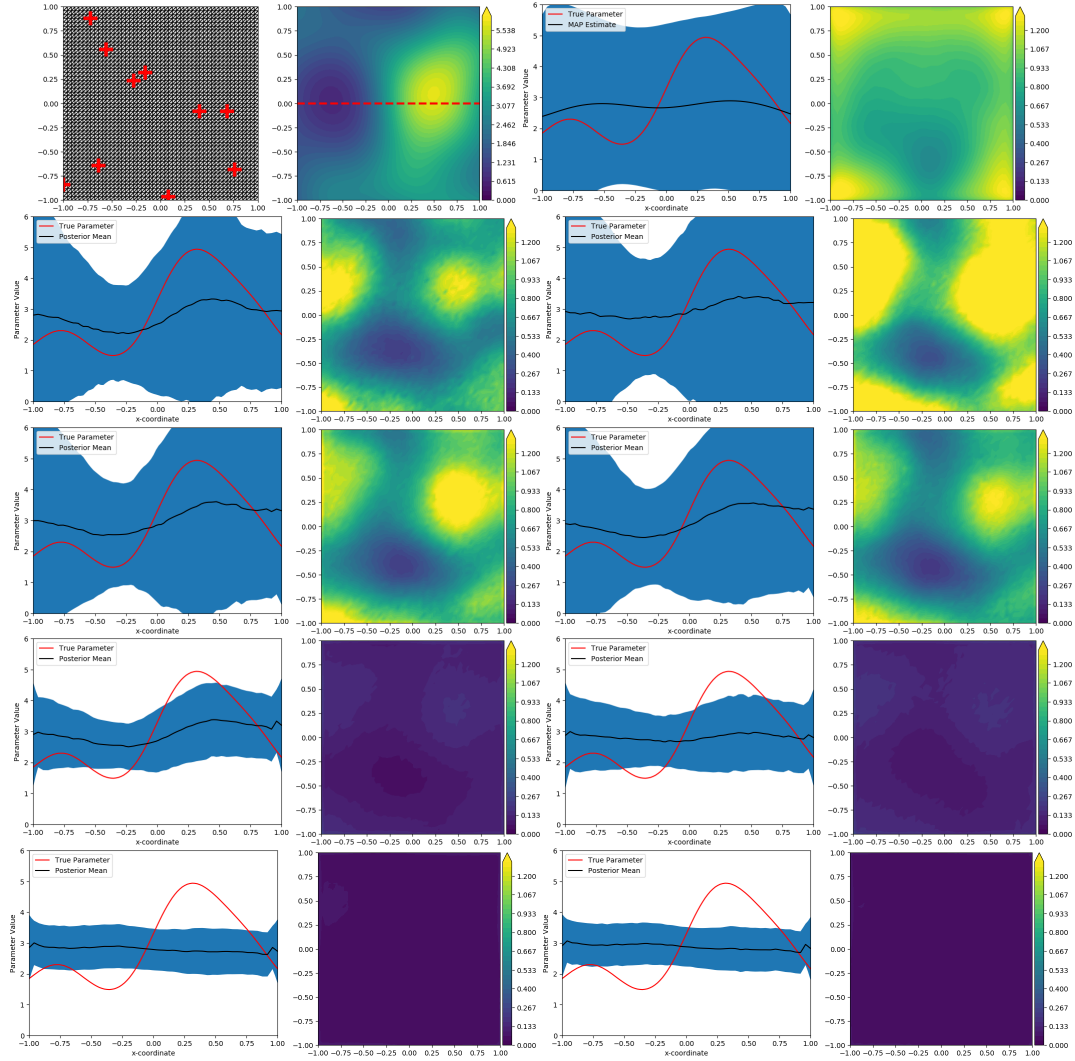


Figure 12: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.10. $\eta = 0.05, M = 500$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	31.72%	30.05%	21.22%
0.001	29.79%	30.69%	25.16%
0.1	30.31%	30.61%	26.02%
0.5	31.56%	31.33%	38.14%

Table 10: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 43.19%.

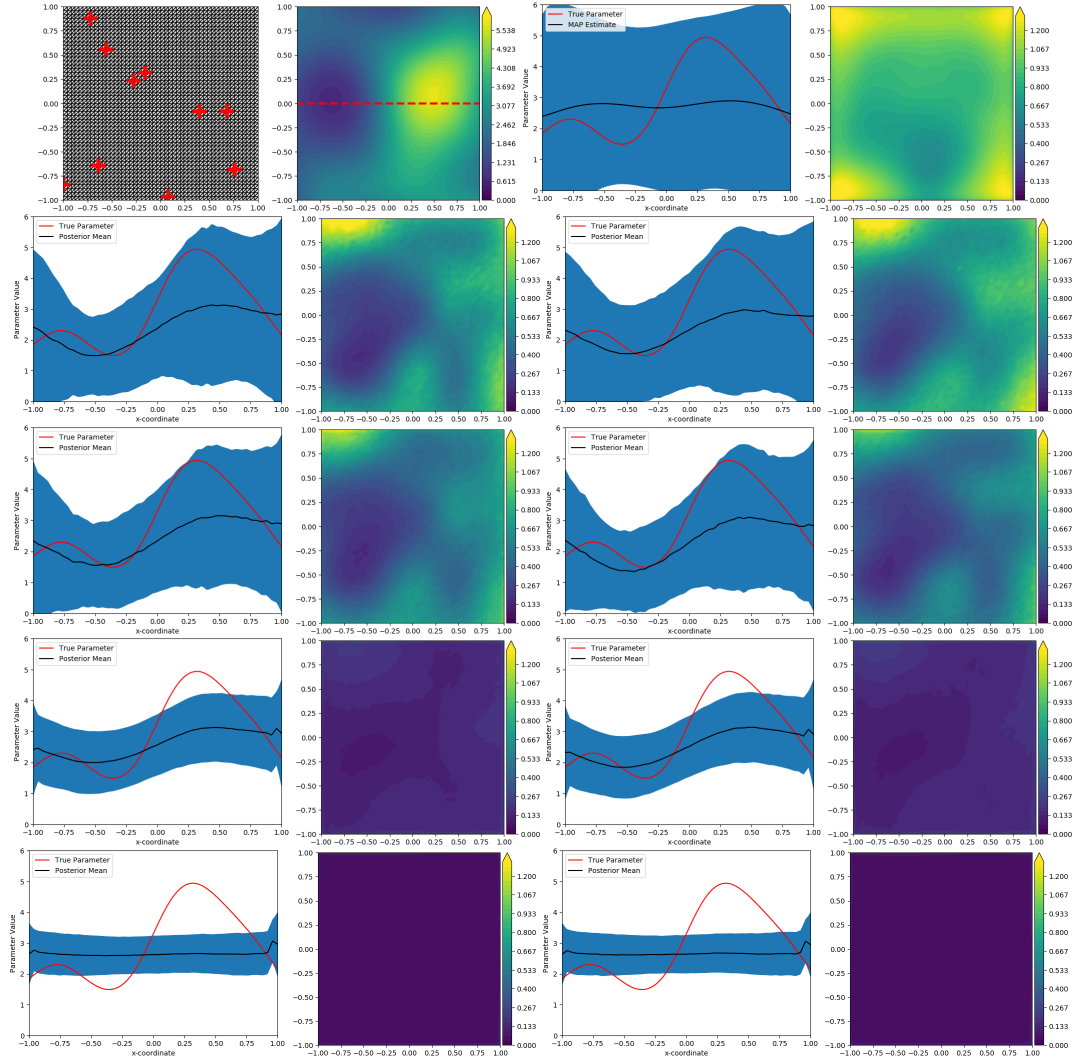


Figure 13: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.11. $\eta = 0.05$, $M = 1000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	31.96%	32.59%	18.36%
0.001	32.14%	31.85%	20.17%
0.1	31.69%	30.87%	25.05%
0.5	33.31%	33.09%	28.05%

Table 11: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 43.19%.

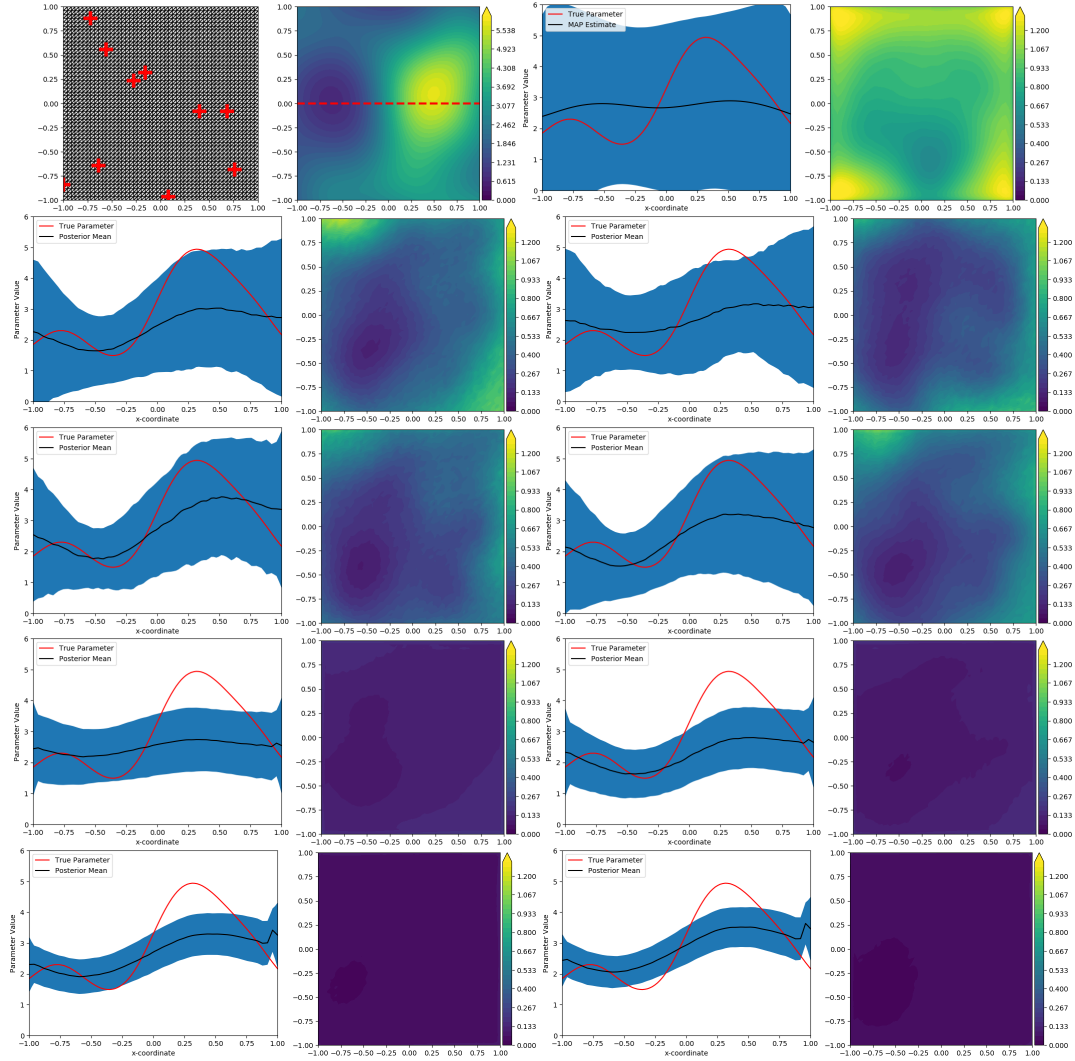


Figure 14: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.12. $\eta = 0.05$, $M = 5000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	36.78%	35.41%	8.11%
0.001	36.90%	35.82%	8.10%
0.1	37.84%	36.48%	7.94%
0.5	36.88%	36.89%	9.41%

Table 12: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 43.19%.

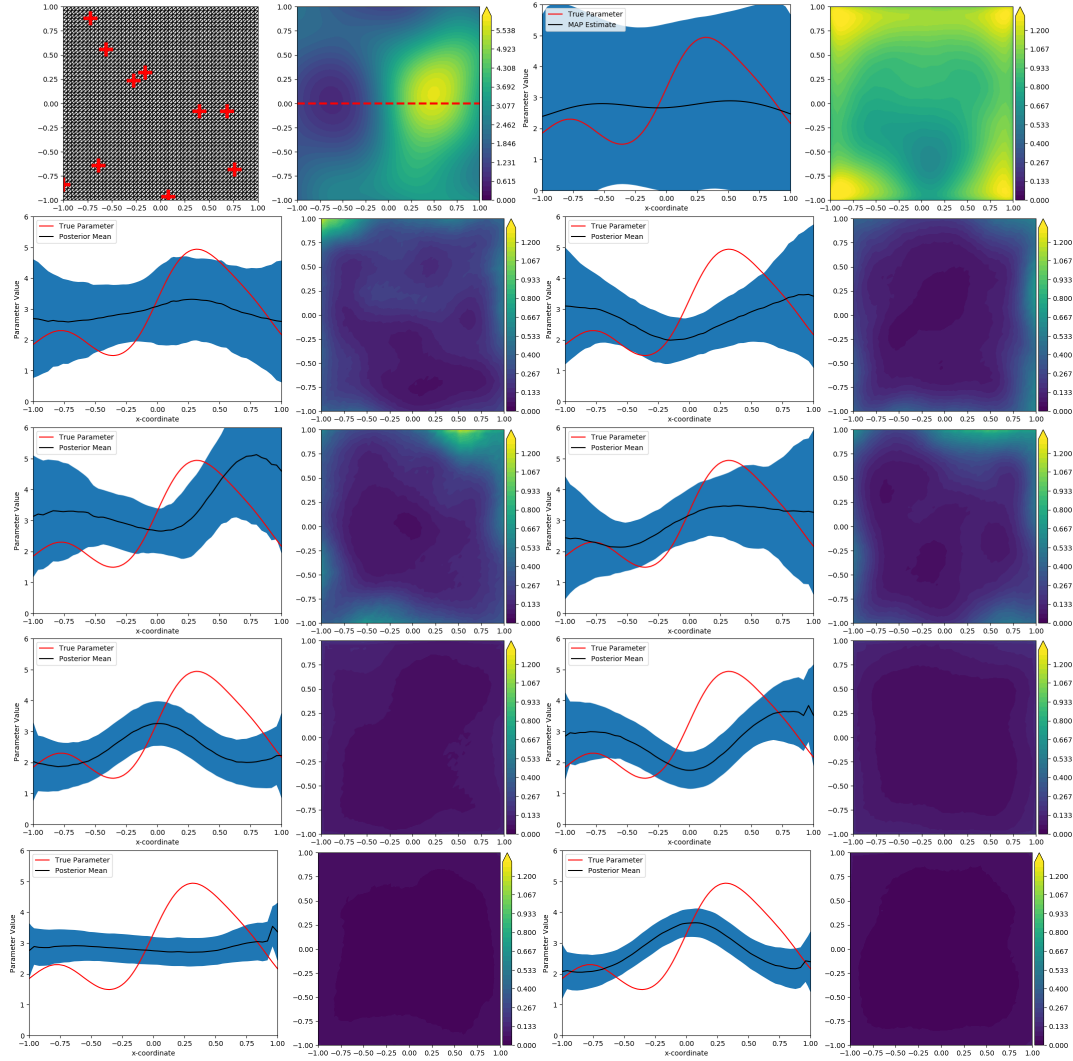


Figure 15: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.13. $\eta = 0.1, M = 50$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	34.95%	36.13%	30.69%
0.001	37.49%	36.15%	27.56%
0.1	35.13%	34.02%	40.31%
0.5	38.44%	34.44%	38.70%

Table 13: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 45.48%.

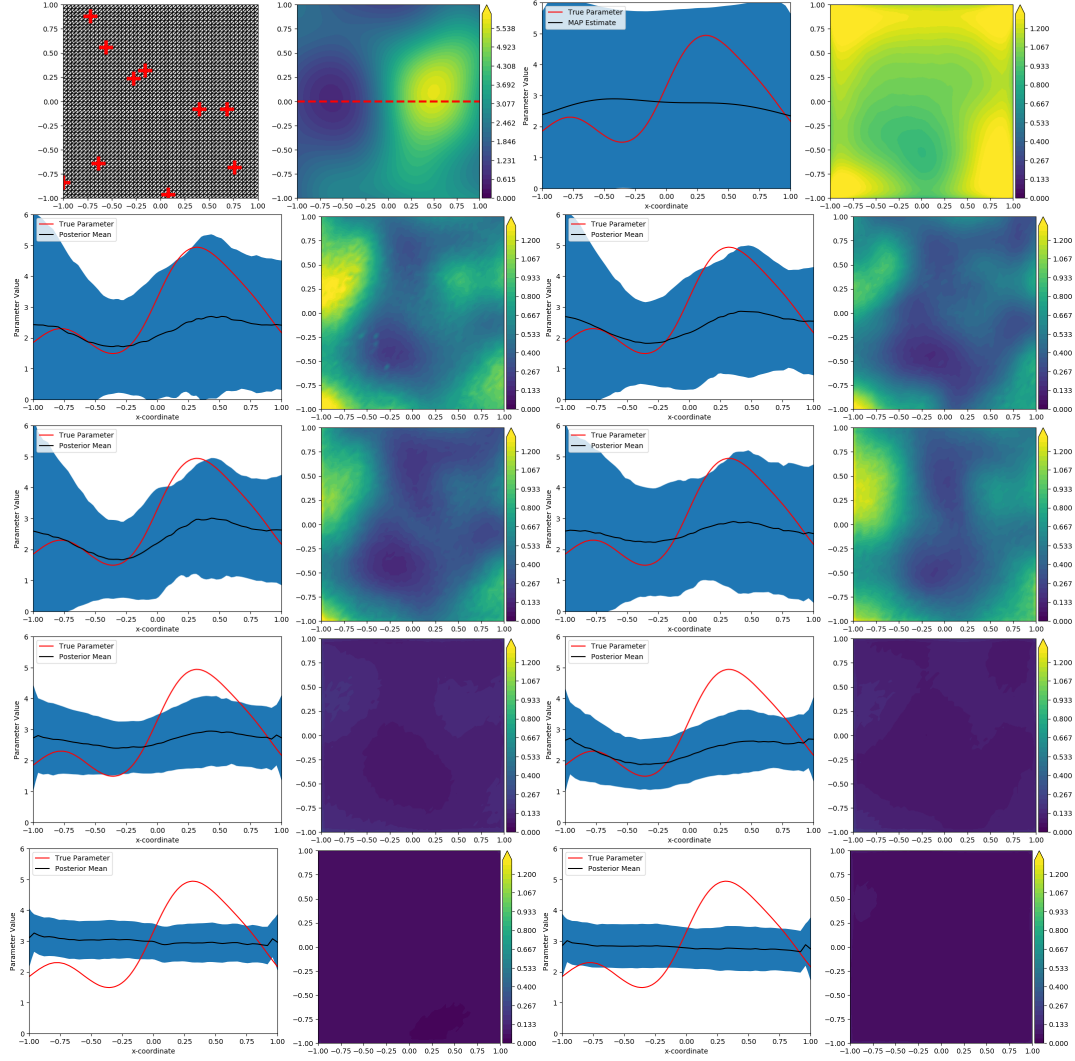


Figure 16: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.14. $\eta = 0.1, M = 500$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	36.61%	38.59%	23.15%
0.001	37.61%	37.82%	23.02%
0.1	35.77%	36.01%	30.22%
0.5	33.78%	35.37%	34.93%

Table 14: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 45.48%.

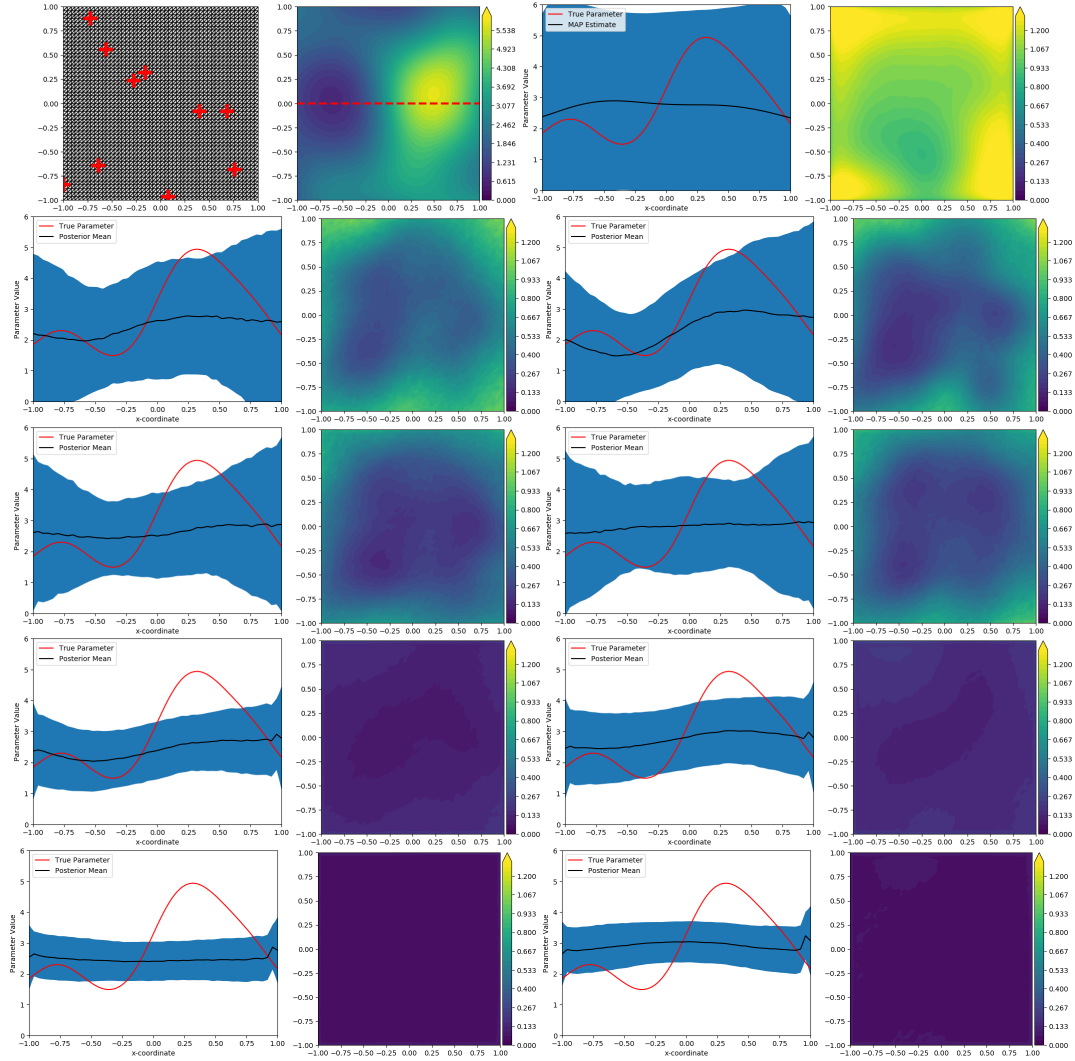


Figure 17: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.15. $\eta = 0.1, M = 1000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	39.96%	32.59%	19.83%
0.001	38.11%	31.85%	20.87%
0.1	37.77%	30.87%	26.37%
0.5	35.51%	33.09%	26.62%

Table 15: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 45.48%.

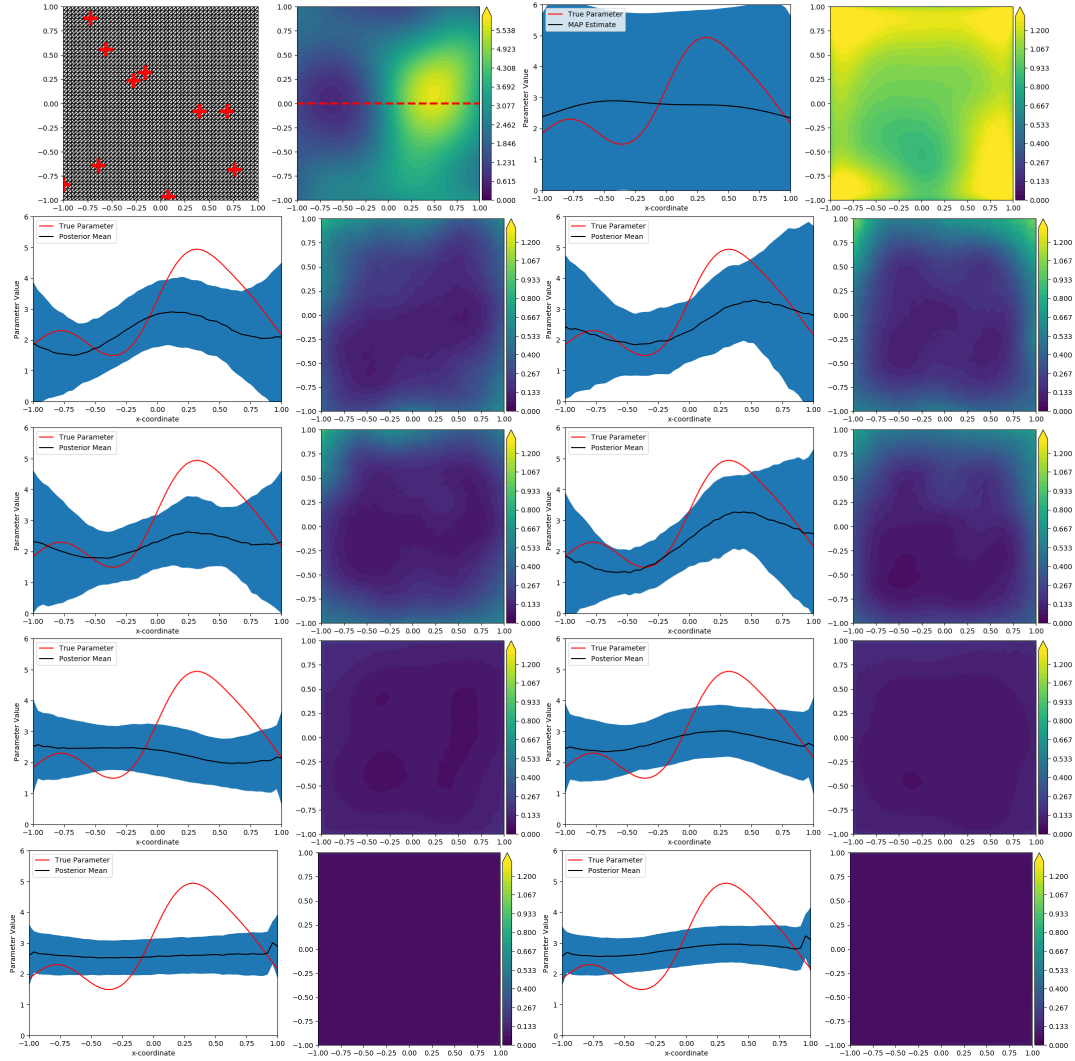


Figure 18: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

D.16. $\eta = 0.1, M = 5000$

α	Relative Error: $\mathbf{u}$		Relative Error: $\mathbf{y}_{\text{obs}}$
	Modelled PtO	Learned PtO	Learned PtO
0.00001	40.45%	35.41%	15.64%
0.001	40.78%	35.82%	15.18%
0.1	41.26%	36.48%	15.39%
0.5	40.85%	36.89%	15.35%

Table 16: Table displaying the relative errors for UQ-VAE. Relative error of MAP estimate: 45.48%.

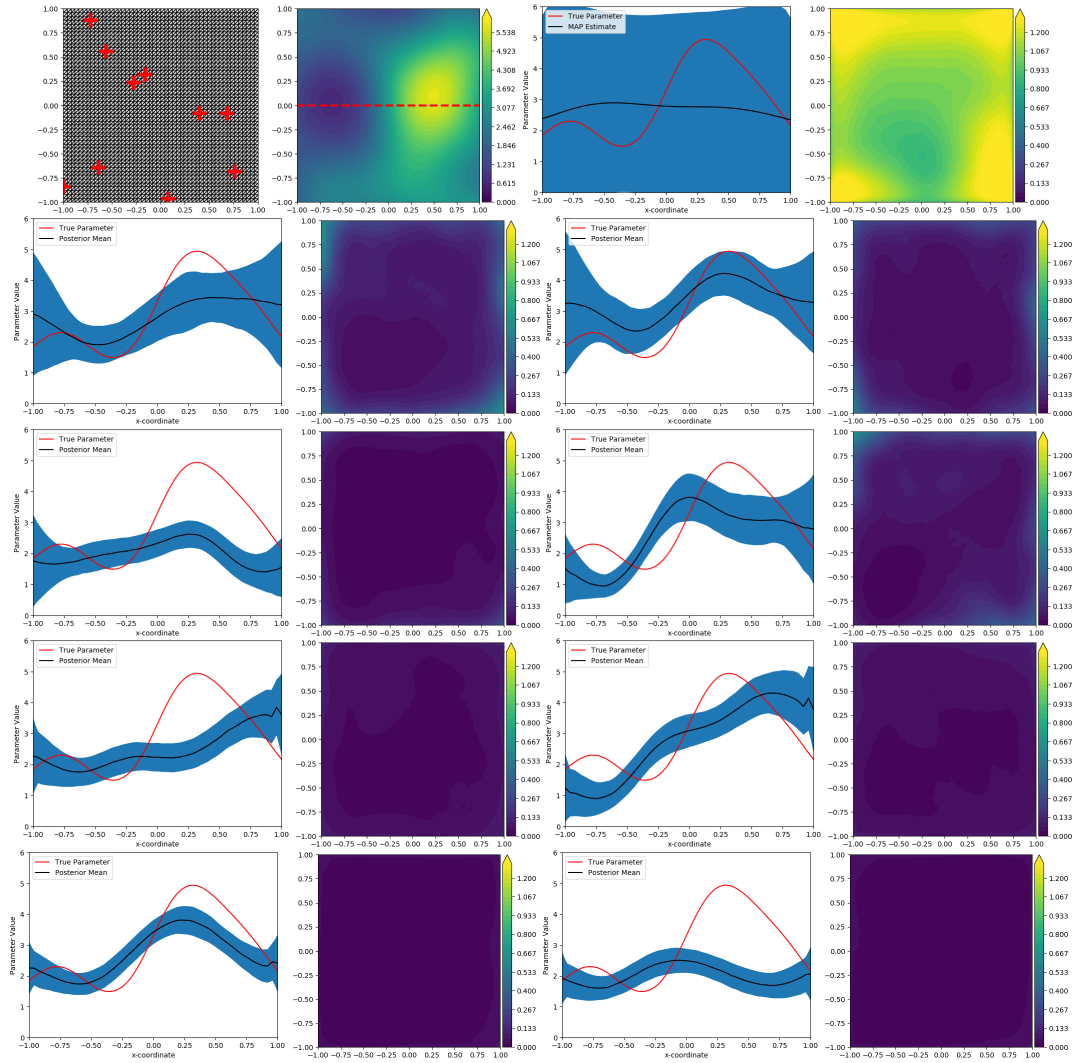


Figure 19: Top row left to right: mesh with sensors denoted with a red cross, true PoI, cross-sectional uncertainty estimate and pointwise posterior variance from Laplace approximation. Second to fourth rows: $\alpha = 0.00001, 0.001, 0.1, 0.5$. First and third columns: cross-sectional uncertainty estimates. Second and fourth columns: approximate pointwise posterior variance. First and second columns: modelled PtO map. Third and fourth columns: learned PtO map.

