



Deep Bayesian Active Learning for Accelerating Stochastic Simulation

Dongxia Wu
University of California, San Diego
La Jolla, CA, USA
dowu@ucsd.edu

Ruijia Niu
University of California, San Diego
La Jolla, CA, USA
rniu@ucsd.edu

Matteo Chinazzi
Northeastern University
Boston, MA, USA
m.chinazzi@northeastern.edu

Alessandro Vespignani
Northeastern University
Boston, MA, USA
a.vespignani@northeastern.edu

Yi-An Ma
University of California, San Diego
La Jolla, CA, USA
yianma@ucsd.edu

Rose Yu
University of California, San Diego
La Jolla, CA, USA
roseyu@ucsd.edu

ABSTRACT

Stochastic simulations such as large-scale, spatiotemporal, age-structured epidemic models are computationally expensive at fine-grained resolution. While deep surrogate models can speed up the simulations, doing so for stochastic simulations and with active learning approaches is an underexplored area. We propose Interactive Neural Process (INP), a deep Bayesian active learning framework for learning deep surrogate models to accelerate stochastic simulations. INP consists of two components, a spatiotemporal surrogate model built upon Neural Process (NP) family and an acquisition function for active learning. For surrogate modeling, we develop Spatiotemporal Neural Process (STNP) to mimic the simulator dynamics. For active learning, we propose a novel acquisition function, Latent Information Gain (LIG), calculated in the latent space of NP based models. We perform a theoretical analysis and demonstrate that LIG reduces sample complexity compared with random sampling in high dimensions. We also conduct empirical studies on three complex spatiotemporal simulators for reaction diffusion, heat flow, and infectious disease. The results demonstrate that STNP outperforms the baselines in the offline learning setting and LIG achieves the state-of-the-art for Bayesian active learning.

CCS CONCEPTS

• Computing methodologies → Active learning settings; Neural networks; Machine learning.

KEYWORDS

Bayesian active learning, neural processes, deep learning

ACM Reference Format:

Dongxia Wu, Ruijia Niu, Matteo Chinazzi, Alessandro Vespignani, Yi-An Ma, and Rose Yu. 2023. Deep Bayesian Active Learning for Accelerating Stochastic Simulation. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '23)*, August 6–10, 2023, Long Beach, CA, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3580305.3599300>



work is licensed under a Creative Commons Attribution International 4.0 License.

KDD '23, August 6–10, 2023, Long Beach, CA, USA.
© 2023 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-0103-0/23/08.
<https://doi.org/10.1145/3580305.3599300>

1 INTRODUCTION

Computational modeling is more than ever at the forefront of infectious disease research due to the COVID-19 pandemic. Stochastic simulations play a critical role in understanding and forecasting infectious disease dynamics, creating what-if scenarios, and informing public health policy making [7]. More broadly, stochastic simulations [2, 42] produce forecasts about complex interactions among people, environment, space, and time given a set of parameters. They provide the numerical tools to simulate stochastic processes in finance [23], chemistry [14] and many other scientific disciplines.

Unfortunately, stochastic simulations at fine-grained spatial and temporal resolution can be extremely computationally expensive. In example, epidemic models for realistic diffusion dynamics simulation via in-silico experiments require a large parameter space (e.g. characteristics of a virus, policy interventions, people's behavior). Similarly, reaction-diffusion systems that play an important role in chemical reaction and bio-molecular processes also involve a large number of simulation conditions. Therefore, hundreds of thousands of simulations are required to explore and calibrate the simulation model with observed experimental data. This process significantly hinders the adaptive capability of existing stochastic simulators, especially in “war time” emergencies, due to the lead time needed to execute new simulations and produce actionable insights that could help guide decision makers.

Learning deep surrogate models to speed up complex simulation has been explored in climate modeling and fluid dynamics for *deterministic* dynamics [3, 16, 41, 43, 50], but not for *stochastic* simulations. These surrogate models can only approximate specific system dynamics and fail to generalize under different parametrization. Especially for pandemic scenario planning, we desire models that can predict futuristic scenarios under different conditions. Furthermore, the majority of the surrogate models are trained *passively* using a simulation data set. This requires a large number of simulations beforehand to cover different parameter regimes of the simulator and ensure generalization.

We propose Interactive Neural Process (INP), a deep Bayesian active learning framework to speed up stochastic simulations. Given parameters such as disease reproduction number, incubation and infectious periods, mechanistic simulators generate future outbreak states with time-consuming numerical integration. INP accelerates

the simulation by guiding a surrogate model to learn the input-output map between parameters and future states, hence bypassing numerical integration.

The deep surrogate model of INP is built upon Neural Process (NP) [13], which lies between Gaussian process (GP) and neural network (NN). NPs can approximate stochastic processes and therefore are well-suitable for surrogate modeling of stochastic simulators. They learn distributions over functions and can generate prediction uncertainty for Bayesian active learning. Compared with GPs, NPs are more flexible and scalable for high-dimensional data with spatiotemporal dependencies. We design a novel Spatiotemporal Neural Process model (STNP) by introducing a time-evolving latent process for temporal dynamics and integrating spatial convolution for spatial modeling.

Instead of learning passively, we design *active learning* algorithms to interact with the simulator and update our model in “real-time”. We derive a new acquisition function, Latent Information Gain (LIG), based on our unique model design. Our algorithm selects the parameters with the highest LIG, queries the simulator to generate new simulation data, and continuously updates our model. We provide theoretical guarantees for the sample efficiency of this procedure over random sampling. We also demonstrate the efficacy of our method on large-scale spatiotemporal epidemic and reaction diffusion models. In summary, our contributions include:

- Interactive Neural Process: a deep Bayesian active learning framework for accelerating large-scale stochastic simulation.
- New surrogate model, Spatiotemporal Neural Process (STNP), for high-dimensional spatiotemporal data that integrates temporal latent process and spatial convolution.
- New acquisition function, Latent Information Gain (LIG), based on the inferred temporal latent process to quantify uncertainty with theoretical guarantees.
- Real-world application to speed up complex stochastic spatiotemporal simulations including reaction-diffusion system, heat flow, and age-structured epidemic dynamics.

2 RELATED WORK

Bayesian Active Learning and Experimental Design. Bayesian active learning, or experimental design is well-studied in statistics and machine learning [4, 6]. Gaussian Processes (GPs) are popular for posterior estimation e.g. [17] and [56], but often struggle in high-dimension. Deep neural networks provide scalable solutions for active learning. Deep active learning has been applied to discrete problems such as image classification [12] and sequence labeling [46] whereas our task is continuous time series. Our problem can also be viewed as sequential experimental design where we design simulation parameters to obtain the desired outcome (imitating the simulator). Foster et al. [9] propose deep design networks for Bayesian experiment design but they require a explicit likelihood model and conditional independence in experiments. Kleinesgesse and Gutmann [22] consider implicit models where the likelihood function is intractable, but computing the Jacobian through sampling path can be expensive and their experiments are mostly limited to low (≤ 10) dimensional design. In contrast, our design space is of much higher-dimension and we do not have access to an explicit likelihood model for the simulator.

Neural Processes. Neural Processes (NP) [13] model distributions over functions and imbue neural networks with the ability of GPs to estimate uncertainty. NP has many extensions such as attentive NP [19] and functional NP [28]. However, NP implicitly assumes permutation invariance in the latent variables and can be limiting in modeling temporal dynamics. Singh et al. [47] proposes sequential NP by incorporating a temporal transition model into NP. Still, sequential NP assumes the latent variables are independent conditioned on the hidden states. We propose STNP with temporal latent process and spatial convolution, which is well-suited for modeling the spatiotemporal dynamics of infectious disease. We apply our model to real-world large-scale Bayesian active learning. Note that even though Garnelo et al. [13] has demonstrated NP for Bayesian optimization, it is only for toy 1-D functions.

Stochastic Simulation and Dynamics Modeling. Stochastic simulations are fundamental to many scientific fields [42] such as epidemic modeling. Data-driven models of infectious diseases are increasingly used to forecast the evolution of an ongoing outbreak [1, 7, 29]. However, very few models can mimic the internal mechanism of a stochastic simulator and answer “what-if questions”. GPs are commonly used as surrogate models for expensive simulators [15, 18, 31, 39], but GPs do not scale well to high-dimensional data. Likelihood-free inference methods [30, 34, 38, 52] learn the posterior of the parameters given the observed data. They do neural density estimation, but require a lot of simulations. For active learning, instead of relying on Monte Carlo sampling, we directly compute the information gain in the latent process. Qian et al. [39] use GPs as a prior for a SEIR model for learning lockdown policy effects, but GPs are computationally expensive and the simple SEIR model cannot capture the real-world large-scale, spatiotemporal dynamics considered in this work. We demonstrate the use of deep sequence model as a prior distribution in Bayesian active learning. Our framework is also compatible with other deep sequence models for time series, e.g. Deep State Space [40], Neural ODE [5].

3 METHODOLOGY

Consider a stochastic process $\{X_1, \dots, X_T\}$, governed by time-varying parameters $\theta_t \in \mathbb{R}^K$, and the initial state $x_0 \in \mathbb{R}^D$. In epidemic modeling, θ_t can represent the effective reproduction number of the virus at a given time, the effective contact rates between individuals belonging to different age groups, the people’s degree of short- or long-range mobility, or the effects of time varying policy interventions (e.g. non-pharmaceutical interventions). The state $x_t \in \mathbb{R}^D$ includes both the daily prevalence and daily incidence for each compartment of the epidemic model (e.g. number of people that are infectious and number of new infected individuals at time t).

Stochastic simulation uses a mechanistic model $F(\theta; \xi)$ to simulate the process where the random variable ξ represents the randomness in the simulator. Let $\theta := (x_0, \theta_1, \dots, \theta_T)$ represent the initial state and all the parameters over time. For each θ , we obtain a different set of simulation data $\{(x_1, \dots, x_T)_m\}_{m=1}^M$. However, realistic large-scale stochastic simulations require the exploration of a large parameter space and are extremely computationally intensive. In the following section, we describe the Interactive Neural

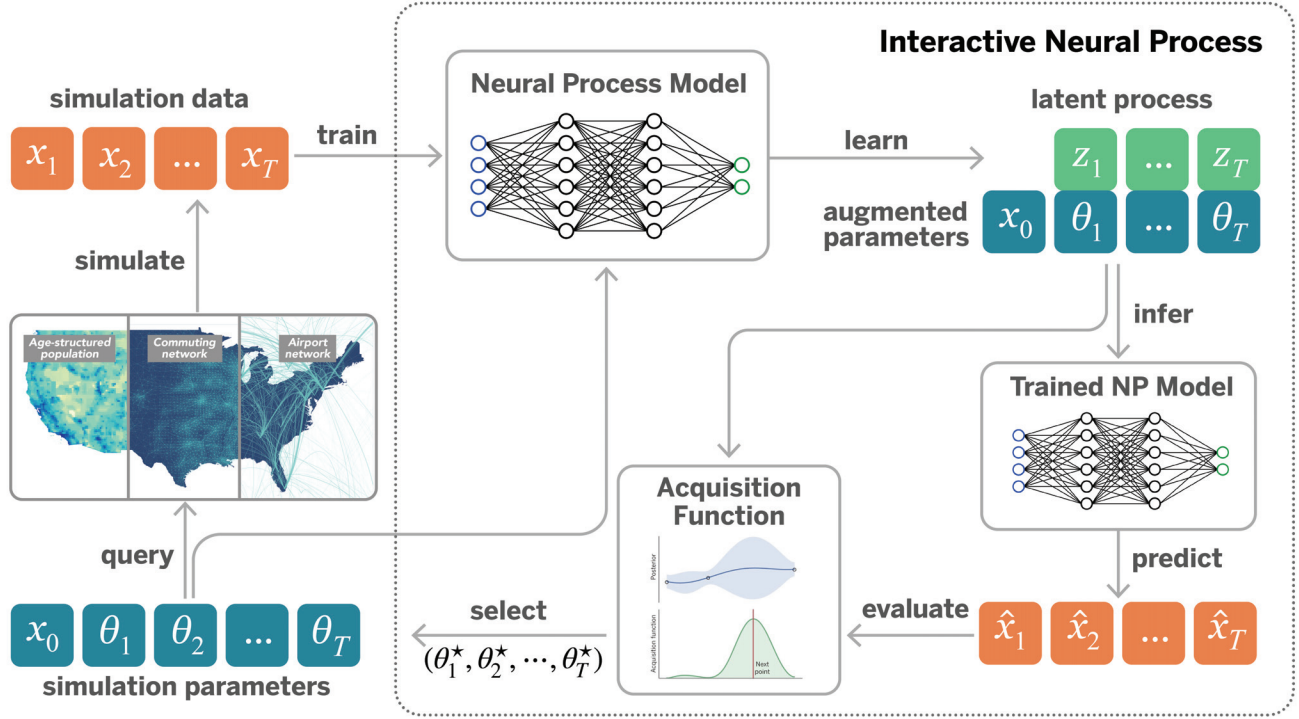


Figure 1: Illustration of the interactive Neural Process (INP). Given simulation parameters and data, INP trains a surrogate model (e.g. STNP) to infer the latent process. The inferred latent process allows prediction and uncertainty quantification. They are used to calculate the acquisition function (e.g. LIG) to select the next set of parameters to query, and simulate more data.

Process (INP) framework to proactively query the stochastic simulator, generate simulation data, in order to learn a fast surrogate model for rapid simulation.

3.1 Interactive Neural Process

INP is used to train a deep surrogate model to mimic the stochastic simulator. As shown in Figure 1, given parameters θ , we query the simulator, i.e., the mechanistic model to obtain a set of simulations $\{(x_1, \dots, x_T)_m\}_{m=1}^M$. We train a NP based model to learn the probabilistic map from parameters to future states. Our NP model can be spatiotemporal to capture complex dynamics such as the disease dynamics of the epidemic simulator. During inference, the model needs to generate predictions $(\hat{x}_1, \dots, \hat{x}_T)$ at the target parameters θ corresponding to different scenarios.

Instead of simulating at a wide range of parameter regimes, we take a Bayesian active learning approach to proactively query the simulator and update the model incrementally. Using NP, we can infer the latent temporal process $(z_1, \dots, z_T)$ that encodes the uncertainty of the current surrogate model. Then we propose a new acquisition function, Latent Information Gain (LIG), to select the θ^* with the highest reward. We use θ^* to query the simulator, and in turn generate new simulation to further improve the model. Next, we describe each of the components in detail.

3.2 Spatiotemporal Neural Process

Neural Process (NP) [13] is a type of deep generative model that represents distributions over functions. It introduces a global latent

variable z to capture the stochasticity and learns the conditional distribution $p(x_{1:T}|\theta)$ by optimizing the evidence lower bound (ELBO):

$$\log p(x_{1:T}|\theta) \geq \mathbb{E}_{q(z|x_{1:T}, \theta)} [\log p(x_{1:T}|z, \theta)] - \text{KL}(q(z|x_{1:T}, \theta) \| p(z)) \quad (1)$$

Here $p(z)$ is the prior distribution for the latent variable. We use $x_{1:T}$ as a shorthand for $(x_1, \dots, x_T)$. The prior distribution $p(z)$ is conditioned on a set of context points $\theta^c, x_{1:T}^c$ as $p(z|x_{1:T}^c, \theta^c)$.

However, the global latent variable z in NP can be limiting for non-stationary, spatiotemporal dynamics in the epidemics. We propose Spatiotemporal Neural Process (STNP) with two extensions. First, we introduce a temporal latent process $(z_1, \dots, z_T)$ to represent the unknown dynamics. The latent process provides an expressive description of the internal mechanism of the stochastic simulator. Each latent variable z_t is sampled conditioning on the past history. Second, we explicitly model the spatial dependency in $x_t \in \mathbb{R}^D$. Rather than treating the dimensions in x_t as independent features, we capture their correlations with regular grids or graphs. For instance, the travel graph between locations can be represented as an adjacency matrix $A \in \mathbb{R}^{D \times D}$.

Given parameters $\{\theta\}$, simulation data $\{x_{1:T}\}$, and the spatial graph A as inputs, STNP models $p(x_{1:T}|\theta, A)$ by optimizing the following ELBO objective:

$$\log p(x_{1:T}|\theta, A) \geq \mathbb{E}_{q(z_{1:T}|x_{1:T}, \theta, A)} \log p(x_{1:T}|z_{1:T}, \theta, A) - \text{KL}(q(z_{1:T}|x_{1:T}, \theta, A) \| p(z_{1:T})) \quad (2)$$

Algorithm 1 Interactive Neural Process

Input: Initial simulation dataset $\mathcal{S}_1$
 Train the model $\text{NP}^{(1)}(\mathcal{S}_1)$;
for $i = 1, 2, \dots$ **do**
 Learn $(z_1, z_2, \dots, z_T) \sim q^{(i)}(z_{1:T}|x_{1:T}, \theta, \mathcal{S}_i)$;
 Predict $(\hat{x}_1, \hat{x}_2, \dots, \hat{x}_T) \sim p^{(i)}(x_{1:T}|z_{1:T}, \theta, \mathcal{S}_i)$;
 Select a batch of data:
 $\{\theta^{(i+1)}\} \leftarrow \arg\max_{\theta} \mathbb{E}_{p(x_{1:T}|z_{1:T}, \theta)} [r(\hat{x}_{1:T}|z_{1:T}, \theta)]$;
 Simulate $\{x_{1:t}^{(i+1)}\} \leftarrow \text{Query the simulator } F(\theta^{(i+1)}; \xi)$;
 Augment training set $\mathcal{S}_{i+1} \leftarrow \mathcal{S}_i \cup \{\theta^{(i+1)}, x_{1:T}^{(i+1)}\}$;
 Update the model $\text{NP}^{(i+1)}(\mathcal{S}_{i+1})$;
end for

where the distributions $q(z_{1:T}|x_{1:T}, \theta, A)$ and $p(x_{1:T}|z_{1:T}, \theta, A)$ are parameterized with neural networks. The prior distribution $p(z_{1:T})$ is conditioned on a set of contextual sequences $p(z_{1:T}|x_{1:T}^c, \theta^c, A)$. Figure 2 visualizes the graphical models of our STNP, the original NP [13] model and Sequential NP [47]. The main difference between STNP and baselines is the encoding procedure to infer the temporal latent process. Compared with STNP which directly embeds the history for z inference at the current timestamp, NP ignores the history and SNP only embeds the partial history information from the previous z .

We implement STNP following an encoder-decoder architecture. The encoder parametrizes the mean and standard deviation of the variational posterior $q(z_{1:T}|x_{1:T}, \theta, A)$ and the decoder approximates the predictive distribution $p(x_{1:T}|z_{1:T}, \theta, A)$. To incorporate the spatial graph information, we use a Diffusion Convolutional Gated Recurrent Unit (DCGRU) layer [25] which integrates graph convolution in a GRU cell. We use multi-layer GRUs to obtain hidden states from the inputs. Using re-parametrization [20], we sample z_t from the encoder and then decode x_t conditioned on z_t in an auto-regressive fashion. To ensure fair comparisons, we adapt NP and SNP to graph-based settings and use the same architecture as STNP to generate the hidden states. Noted if the spatial dependency is regular grid-based, then the DCGRU layer is replaced to Convolutional LSTM layer [27, 45, 51, 53, 54], and there is no adjacency matrix A in Equation 2.

3.3 Bayesian Active Learning

Algorithm 1 details a Bayesian active learning algorithm, based on Bayesian optimization [11, 44]. We train an NP model to interact with the simulator and improve learning. Let the superscript (i) denote the i -th interaction. We start with an initial data set $\mathcal{S}_1 = \{\theta^{(1)}, x_{1:T}^{(1)}\}$ and use it to train our NP model and learn the latent process. During inference, given the augmented parameters θ , we use the trained NP model to predict the future states $(\hat{x}_1, \dots, \hat{x}_T)$. We evaluate the current models' predictions with an acquisition function $r(\hat{x}_{1:T}|z_{1:T}, \theta)$ and select the set of parameters $\{\theta^{(i+1)}\}$ with the highest reward. We query the simulator with $\{\theta^{(i+1)}\}$ to augment the training data set $\mathcal{S}_{i+1}$ and update the NP model for the next iteration.

The choice of the reward (acquisition) function r depends on the goal of the active learning task. For example, to find the model

that best fits the data, the reward function can be the log-likelihood $r = \log p(\hat{x}_{1:T}|\theta, A)$. To collect data and reduce model uncertainty in Bayesian experimental design, the reward function can be the mutual information. In what follows, we discuss different strategies to design the reward/acquisition function. We also propose a novel acquisition function based on information gain in the latent space tailored to our STNP model.

3.4 Reward/Acquisition functions

For regression tasks, standard acquisition functions for active learning include Maximum Mean Standard Deviation (Mean STD), Maximum Entropy, Bayesian Active Learning by Disagreement (BALD) or expected information gain (EIG), and random sampling [12]. We explore various acquisition functions and their approximations in the context of NP. We also introduce a new acquisition function based on our unique NP design called Latent Information Gain (LIG). The details of Mean STD and Maximum Entropy are shown in the Appendix C.

BALD/Expected Information Gain (EIG). BALD [17] quantifies the mutual information between the prediction and model posterior $H(\hat{x}_{1:T}|\theta) - H(\hat{x}_{1:T}|z_{1:T}, \theta)$, which is equivalent to the expected information gain (EIG). Computing the EIG for surrogate modeling is challenging since $p(\hat{x}_{1:T}|z_{1:T}, \theta)$ cannot be found in closed form in general. The integrand is intractable and conventional MC methods are not applicable [10]. One way to get around this is to employ a nested MC estimator with quadratic computational cost for sampling [35, 49], which is computationally infeasible. To reduce the computational cost, we assume $p(\hat{x}_{1:T}|z_{1:T}, \theta)$ follows multivariate Gaussian distribution. Each feature of $\hat{x}_{1:T}$ can be parameterized with mean and standard deviation predicted from the surrogate model, assuming output features are independent with each other. This distribution assumption can be limiting in the high-dimensional spatiotemporal domain, which makes EIG less informative.

Latent Information Gain (LIG). To overcome the limitations mentioned above, we propose a novel acquisition function by computing the expected information gain in the latent space rather than the observational space. To design this acquisition function, we prove the equivalence between the expected information gain in the observational space and the expected KL divergence in the latent processes w.r.t. a candidate parameter θ , as illustrated by the following proposition.

PROPOSITION 1. *The expected information gain (EIG) for Neural Process is equivalent to the KL divergence between the prior and posterior in the latent process, that is*

$$\begin{aligned} \text{EIG}(\hat{x}_{1:T}, \theta) &:= \mathbb{E}[H(\hat{x}_{1:T}) - H(\hat{x}_{1:T}|z_{1:T}, \theta)] \\ &= \mathbb{E}_{p(\hat{x}_{1:T}|\theta)} [\text{KL}(p(z_{1:T}|\hat{x}_{1:T}, \theta) \| p(z_{1:T}))] \quad (3) \end{aligned}$$

See proof in the Appendix C. Inspired by this fact, we propose a novel acquisition function computing the expected KL divergence in the latent processes and name it LIG. Specifically, the trained NP model produces a variational posterior given the current dataset $\mathcal{S}$ as $p(z_{1:T}|\mathcal{S})$. For every parameter θ remained in the search space, we can predict $\hat{x}_{1:T}$ with the decoder. We use $\hat{x}_{1:T}$ and θ as input to the encoder to re-evaluate the posterior $p(z_{1:T}|\hat{x}_{1:T}, \theta, \mathcal{S})$.

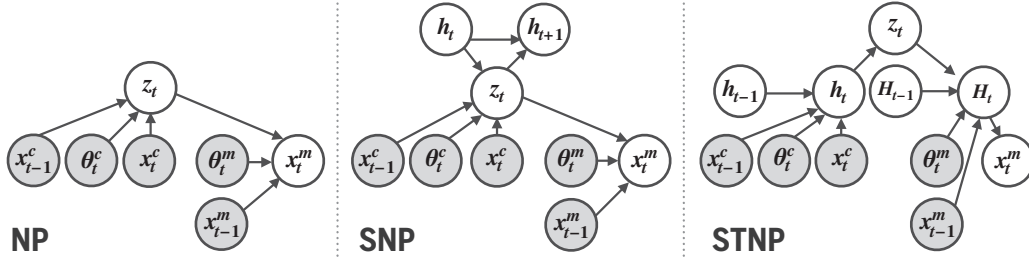


Figure 2: Graphical model comparison: Neural Process, Sequential Neural Process and our Spatiotemporal Neural Process.

LIG computes the distributional difference with respect to the latent process $z_{1:T}$ as $\mathbb{E}_{p(\hat{x}_{1:T}|\theta)} [\text{KL}(p(z_{1:T}|\hat{x}_{1:T}, \theta, S) \| p(z_{1:T}|S))]$, where $\text{KL}(\cdot \| \cdot)$ denotes the KL-divergence between two distributions.

In this way, conventional MC method becomes applicable, which helps reduce the quadratic computational cost to linear. At the same time, although $z_{1:T}$ are assumed to be multivariate Gaussian and are parameterized with mean and standard deviation, they are only in the latent space not the observational space. Moreover, LIG is also more computationally efficient and accurate for batch active learning. Due to the context aggregation mechanism of NP, we can directly calculate LIG with respect to a batch of θ in the candidate set. This is not available for baseline acquisition functions. They all require calculating the scores one by one for all θ in the candidate set and select a batch of θ based on their scores. Such approach is both slow and inaccurate as acquiring points that are informative individually are not necessarily informative jointly [21].

3.5 Theoretical Analysis

We shed light onto the intuition behind choosing adaptive sample selection over random sampling via analyzing a simplifying situation. Assume that at a certain stage we have learned a feature map Ψ which maps the input θ of the neural network to the last layer. Then the output X can be modeled as $X = \langle \Psi(\theta), z^* \rangle + \epsilon$, where z^* is the true hidden variable, ϵ is the random noise.

Our goal is to generate an estimate $\hat{z}$, and use it to make predictions $\langle \Psi(\theta), \hat{z} \rangle$. A good estimate shall achieve small error in terms of $\|\hat{z} - z^*\|_2$ with high probability. In the following theorem, we prove that greedily maximizing the variance of the prediction to choose θ will lead to an error of order $O(d)$ less than that of random exploration in the space of θ , which is significant in high dimension.

THEOREM 1. *For random feature map $\Psi(\cdot)$, greedily optimizing the KL divergence, $\text{KL}(p(z|\hat{x}, \theta) \| p(z))$, or equivalently the variance of the posterior predictive distribution $\mathbb{E}[(\langle \Psi(\theta), \hat{z} \rangle - \mathbb{E}[\langle \Psi(\theta), \hat{z} \rangle])^2]$ in search of θ will lead to an error $\|\hat{z}_t - z^*\|_2$ of order $O(\sigma d / \sqrt{t})$ with high probability. On the other hand, random sampling of θ will lead to an error of order $O(\sigma d^2 / \sqrt{t})$ with high probability.*

See proofs in the Appendix A.1.

4 EXPERIMENTS

We evaluate our proposed STNP for its surrogate modeling performance in the offline learning setting and LIG acquisition function for active learning performance. We aim to verify that (a) LIG outperforms other acquisition functions in the NP and GP model setting for deep Bayesian active learning on non-spatiotemporal surrogate modeling, (b) STNP outperforms other existing baselines for spatiotemporal surrogate modeling in the offline learning setting, and (c) LIG outperforms other acquisition functions in the STNP model setting for deep Bayesian active learning on spatiotemporal surrogate modeling. The implementation code is available at <https://github.com/Rose-STL-Lab/Interactive-Neural-Process>.

4.1 Experimental Setup

We experiment with the following four stochastic simulators.

SEIR Compartmental Model. To highlight the difference between NP and GP, we begin with a simple stochastic, discrete, chain-binomial SEIR compartmental model as our stochastic simulator. In this model, susceptible individuals (S) become exposed (E) through interactions with infectious individuals (I) and are eventually removed (R), details are deferred to the Appendix C.

We set the total population $N = S + E + I + R$ as 100,000, the initial number of exposed individuals as $E_0 = 2,000$, and the initial number of infectious individuals as $I_0 = 2,000$. We assume latent individuals move to the infectious stage at a rate $\epsilon \in [0.25, 0.65]$ (step 0.05), the infectious period μ^{-1} is set to be equal to 1 day, and we let the basic reproduction number R_0 (which in this case coincides with the transmissibility rate β) vary between 1.1 and 4.0 (step 0.1). Here, each (β, ϵ) pair corresponds to a specific scenario, which determines the parameters θ . We simulate the first 100 days of the epidemic with a total of 300 scenarios and generate 30 samples for each scenario.

We predict the number of individuals in the infectious compartment. The input is (β, ϵ) pair and the output is the 100 days' infection prediction. As the simulator is not spatiotemporal, we use the vanilla NP model with the global latent variable z . For each epoch, we randomly select 10% of the samples as context. Implementation details are deferred to Appendix C.

Reaction Diffusion Model. The reaction-diffusion (RD) system [48] is a spatiotemporal model that simulates how two chemicals might react to each other as they diffuse through a medium together. The simulation is based on initial pattern, feed rate (θ_0),

removal rate (θ_1) and reaction between two substances. We use an RD simulator to generate sequences from 0 to 500 timestamps, sampled every 100 timestamps, resulting into 5 timestamps for each simulated sequence. Every timestamp is a 3D tensor ($2 \times 32 \times 32$) with dimension 0 corresponds to the two substances in the reaction and dimension 1, 2 are the image representation of the reaction diffusion processes. Each sequence is simulated with a unique feed rate $\theta_0 \in [0.029, 0.045]$ and kill rate $\theta_1 \in [0.055, 0.062]$ combination. There are 200 uniformly sampled scenarios, corresponding to (θ_0, θ_1) combinations.

We implement STNP to mimic the reaction diffusion simulator with feed rate (θ_0) and kill rate (θ_1) as input. The initial state of the reaction is fixed. We use multiple convolutional layers with a linear layer to encode the spatial data into latent space. We use an LSTM layer to encode the latent spatial data with θ_0, θ_1 to map the input-output pairs to hidden features $z_{1:5}$. With (θ_0, θ_1) , and $z_{1:5}$ sampled from the posterior distribution, we use an LSTM layer and deconvolutional layers to simulate reaction diffusion sequence. For each epoch, we randomly select 20% samples as context sequence.

Heat Model. The model is to predict the spatial solution fields of the Heat equation [36]. The ground-truth data is generated from the standard numerical solver used in [24]. The experiment setting also follows [24]. The examples are generated by solvers running with 32×32 meshes. The corresponding output dimension is 1024. The input consists of three parameters that control the thermal conductivity and the flux rate.

Local Epidemic and Mobility model. The Local Epidemic and Mobility model (LEAM-US) is a stochastic, spatial, age-structured epidemic model based on a metapopulation approach which divides the US in more than 3,100 subpopulations, each one corresponding to a each US county or statistically equivalent entity. Population size and county-specific age distributions reflect Census' annual resident population estimates for year 2019. We consider individuals divided into 10 age groups. Contact mixing patterns are age-dependent and state specific and modeled considering contact matrices that describe the interaction of individuals in different social settings [33]. LEAM-US integrates a human mobility layer, represented as a network, using both short-range (i.e., commuting) and long-range (i.e., flights) mobility data, see more details in Appendix C.

We separate data in California monthly to predict the 28 days' sequence from the 2nd to the 29th day of each month from March to December. Each θ includes the county-level parameters of LEAM-US and state level incidence and prevalence compartments. The total number of dimension in θ is 16, 912, see details in Appendix C. Overall, there are 315 scenarios in the search space, corresponding to 315 different θ with total 16,254 samples. We split 78% of the data as the candidate set, and 11% for validation and test. For active learning, we use the candidate set as the search space.

We instantiate an STNP model to mimic an epidemic simulator that has θ at both county and state level and x_t at the state level. We use county-level parameter θ together with a county-to-county mobility graph A in California as input. We use the DCGRU layer [25] to encode the mobility graph in a GRU. We use a linear layer to map the county-level output to hidden features at the state level. For both the state-level encoder and decoder, we use multi-layer

GRUs. For each epoch, we randomly select 20% samples as context sequence.

4.2 Offline Learning Performance

We compared our proposed STNP with vanilla NP [13], SNP [47], Masked Autoregressive Flow (MAF) [37], the RNN baseline with variational dropout (RNN) [55], and VAE-based deep surrogate model for multi-fidelity active learning (DMFAL) [24]. The implementation details can be seen in Appendix C. The key innovation of STNP is the introduced temporal latent process. To ensure fair comparison, we modified baselines for the RD and Heat model by adding convolutional layers for data encoding and deconvolutional layers for sequence generation. For the LEAM-US model, we modified NP by adding the convolutional layers with diffusion convolution [26] to embed the graphs. Similarly, we modified SNP by replacing the convolutional layers with diffusion convolution. The rest baselines do not support LEAM-US surrogate modeling. Table 1 shows the testing MAE of different NP models trained in an offline fashion. Our STNP significantly improves the performance and can accurately learn the simulator dynamics for all three experiments.

Figure 3 left compares the NP and GP performance on one scenario in the held-out test set. It shows the ground truth and the predicted number of infectious population for the first 50 days. We also include the confidence intervals (CI) with 5 standard deviations for ground truth and NP predictions and 1 standard deviation for GP predictions. We observe that NP fits the simulation dynamics better than GP for mean prediction. Moreover, NP has closer CIs to the truth, reflecting the simulator's intrinsic uncertainty. GP shows larger CIs which represent the model's own uncertainty. Note that NP is much more flexible than GP and can scale easily to high-dimensional data. Figure 3 middle indicates STNP can accurately predict various patterns corresponding to different (θ_0, θ_1) . This confirms that our STNP is able to capture the high-dimensional spatiotemporal dependencies in RD simulations. Figure 3 right visualize the STNP predictions in four key compartments of a typical scenario with $R_0 = 3.1$ from March 2nd to March 29th. The confidence interval is plotted with 2 standard deviations. We can see that both the mean and confidence interval of STNP predictions match the truth well. These two results demonstrate the promise that the generative STNP model can serve as a deep surrogate model for RD and LEAM-US simulator.

4.3 Active Learning

Implementation Details. We compare 6 different acquisition functions with NP for SEIR model and STNP for RD, Heat, and LEAM-US model. For SEIR, the initial training dataset has 2 scenarios and we continue adding 1 scenario per iteration to the training set until the test loss converges to the offline modeling performance. We also include GP with 3 different acquisition functions and Sequential Neural Likelihood (SNL) (see Table 3). For the RD and Heat model, all acquisition functions start with the same 5 scenarios randomly picked from the training dataset. Then we continue adding 5 scenarios per iteration to the training set until the test loss converges. Similarly, the LEAM-US model begins with 27 training data and we continue adding 8 scenarios per iteration to the training set until

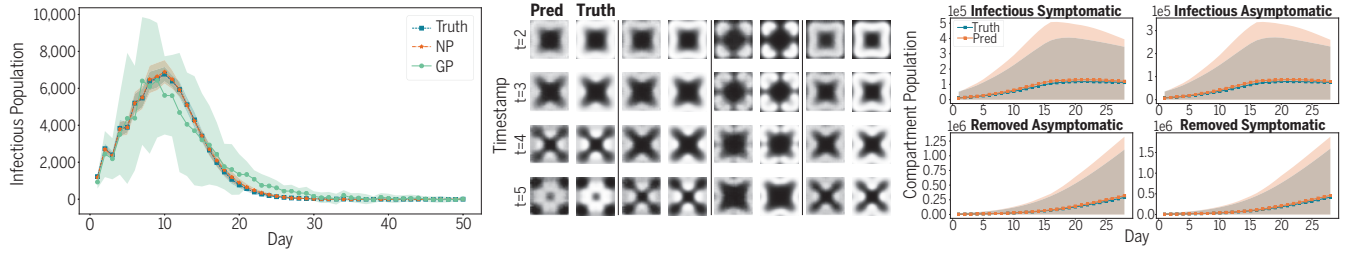


Figure 3: Prediction visualizations, Left: Accuracy and uncertainty quantification comparison between Neural Process (NP) and Gaussian process (GP) in SEIR simulator. Middle: STNP predictions for spatiotemporal patterns of substances in Reaction-Diffusion simulator. Right: STNP predictions for the number of individuals in Infectious and Removed compartments in LEAM-US simulator.

Table 1: Surrogate model performance comparison using MAE in Reaction-Diffusion, Heat, and LEAM simulator (population divided by 1000).

Model	NP	SNP	MAF	RNN	DMFAL	STNP
RD	3.37 ± 0.18	3.11 ± 0.07	7.22 ± 0.66	3.44 ± 0.07	4.1 ± 0.02	2.84 ± 0.17
HEAT	$1.05e-2 \pm 1.1e-3$	$9.62e-3 \pm 1e-3$	$2.1e-2 \pm 4.8e-3$	$3.2e-2 \pm 1.7e-3$	$1.35e-2 \pm 3e-4$	$7.36e-3 \pm 6.7e-4$
LEAM-US	24.2 ± 5.9	21.8 ± 0.8	–	–	–	6.3 ± 0.8

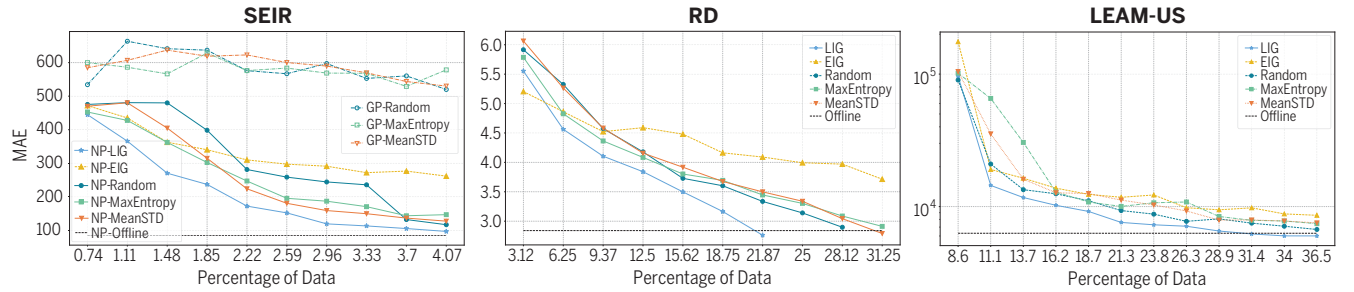


Figure 4: MAE loss versus the percentage of samples for Bayesian active learning. The Black dash line shows the offline learning performance with the entire data set available for training. Left: GP and NP for SEIR. Middle: STNP for RD. Right: STNP for LEAM-US.

the validation loss converges. We measure the average performance over three random runs and report the MAE for the test set.

Active Learning Performance. Figure 4 shows the testing MAE versus the percentage of samples included for training. The percentage of data is linearly proportional to the overall running time. This figure shows our proposed LIG always has the best MAE performance until the convergence for SEIR, RD, and LEAM-US tasks. As shown in figure 4 left, we compare different acquisition functions on both NP and GP for SEIR model. It shows none of the GP methods converge after selecting 4.07% of the data for training while NP methods converge much faster. Our proposed acquisition function LIG is the most sample efficient in acquisition functions used for NP. It takes only 4.07% of the data to converge and reach the NP offline performance, which uses the entire training set for training. Moreover, there is an enormous gap between LIG and EIG with respect to the active learning performance. This validates our theory that the uncertainty of the deep surrogate model is better

measured on the latent space instead of the predictions. Similarly in figure 4 middle and right, we compared LIG with other acquisition functions on STNP for RD and LEAM-US model. It shows LIG converges to the offline performance using only 21.87% of data for RD experiment and 31.4% of data for LEAM-US experiment. Therefore, it is consistent among all three experiments that our proposed LIG always has the best MAE performance until convergence. Notice that for figure 4 right, it shows the log scale MAE versus the percentage of samples included for training. Detailed performance comparison including mean and standard deviation for all four tasks including Heat can be seen in Appendix B.2. Our proposed LIG also has the best MAE performance for the Heat task.

Exploration Exploitation Trade-off. To understand the large performance gap for LIG vs. baselines, we visualize the values of test MAE and the acquisition function score for each Bayesian active learning iteration for SEIR model, shown in Figure 5. For EIG, Mean STD, and Maximum Entropy, they all tend to exploit the region

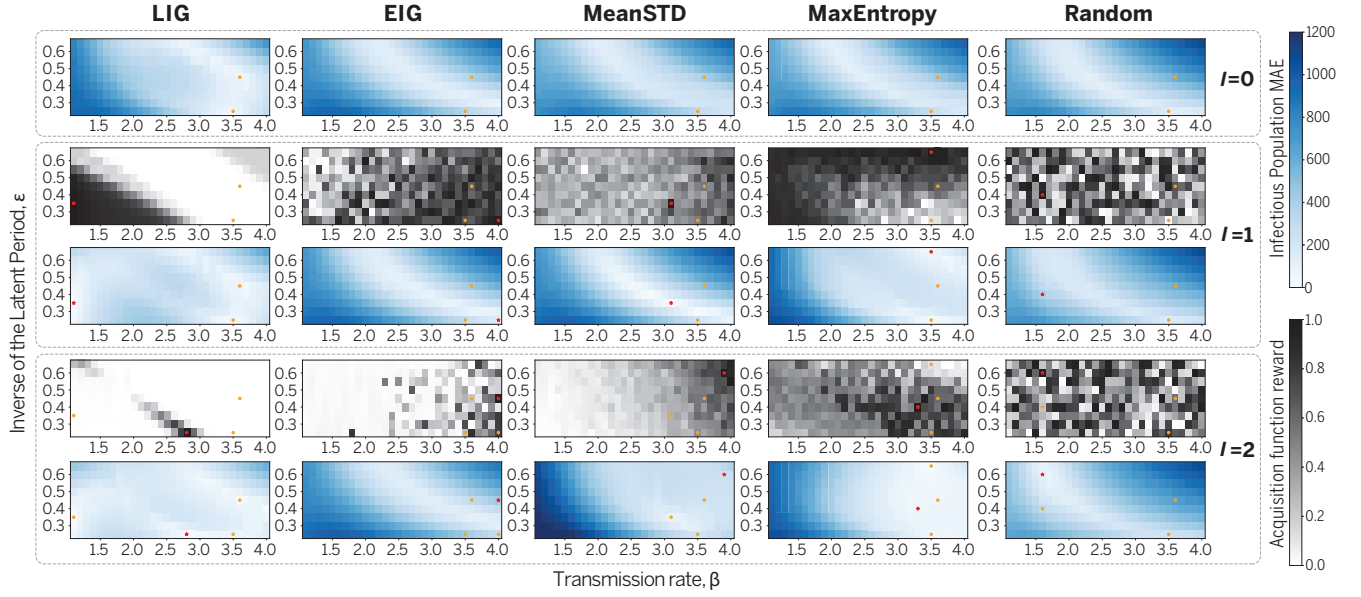


Figure 5: Acquisition function behavior visualization in SEIR model. For each iteration, top row is the current MAE mesh in infectious population for all (β, ϵ) candidates. Bottom row is the acquisition function score. Yellow dots are existing parameters. Red stars are the newly selected parameters.

with large transmission rate for the first 2 iterations. Including these scenarios makes the training set unbalanced. The MAE in the region with small transmission rate become worse after 2 iterations. Meanwhile, Random is doing pure exploration. The improvement of MAE performance is not apparent after 2 iterations. Our proposed LIG is able to reach a balance by exploiting the uncertainty in the latent process and encouraging exploration. Hence, with a small number of iterations ($I = 2$), it has already selected “informative scenarios” in the search space.

5 CONCLUSION

We propose a unified framework Interactive Neural Processes (INP) for deep Bayesian active learning, that can seamlessly interact with existing stochastic simulators and accelerate simulation. Specifically, we design STNP to approximate the underlying simulation dynamics. It infers the latent process which describes the intrinsic uncertainty of the simulator. We exploit this uncertainty and propose LIG as a powerful acquisition function in deep Bayesian active learning. We perform a theoretical analysis and demonstrate that our approach reduces sample complexity compared with random sampling in high dimension. We also did extensive empirical evaluations on several complex real-world spatiotemporal simulators to demonstrate the superior performance of our proposed STNP and LIG. For the future work, we plan to leverage Bayesian optimization techniques to directly optimize for the target parameters with auto-differentiation.

ACKNOWLEDGMENTS

This work was supported in part by U.S. Department Of Energy, Office of Science, Facebook Data Science Research Awards, U. S. Army

Research Office under Grant W911NF-20-1-0334, and NSF Grants #2134274 and #2146343, as well as NSF-SCALE MoDL (2134209) and NSF-CCF-2112665 (TILOS). M.C. and A.V. acknowledge support from grant HHS/CDC 5U01IP0001137.

REFERENCES

- [1] Sercan Arik, Chun-Liang Li, Jinsung Yoon, Rajarishi Sinha, Arkady Epshteyn, Long Le, Vikas Menon, Shashank Singh, Leyou Zhang, Martin Nikoltchev, et al. 2020. Interpretable Sequence Learning for Covid-19 Forecasting. *Advances in Neural Information Processing Systems* 33 (2020).
- [2] Søren Asmussen and Peter W Glynn. 2007. *Stochastic simulation: algorithms and analysis*. Vol. 57. Springer Science & Business Media.
- [3] Salva Rühling Cachay, Venkatesh Ramesh, Jason N. S. Cole, Howard Barker, and David Rolnick. 2021. ClimART: A Benchmark Dataset for Emulating Atmospheric Radiative Transfer in Weather and Climate Models. In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track*. <https://arxiv.org/abs/2111.14671>
- [4] Kathryn Chaloner and Isabella Verdinelli. 1995. Bayesian experimental design: A review. *Statist. Sci.* (1995), 273–304.
- [5] Ricky TQ Chen, Yulia Rubanova, Jesse Bettencourt, and David Duvenaud. 2018. Neural ordinary differential equations. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*. 6572–6583.
- [6] David A Cohn, Zoubin Ghahramani, and Michael I Jordan. 1996. Active learning with statistical models. *Journal of artificial intelligence research* 4 (1996), 129–145.
- [7] Estee Y Cramer, Velma K Lopez, Jarad Niemi, Glover E George, Jeffrey C Cegan, Ian D Dettwiller, William P England, Matthew W Farthing, Robert H Hunter, Brandon Lafferty, et al. 2021. Evaluation of individual and ensemble probabilistic forecasts of COVID-19 mortality in the US. *medRxiv* (2021).
- [8] Simon Du, Jason Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. 2019. Gradient Descent Finds Global Minima of Deep Neural Networks. In *Proceedings of the 36th International Conference on Machine Learning (ICML)*. 1675–1685.
- [9] Adam Foster, Desi R Ivanova, Ilyas Malik, and Tom Rainforth. 2021. Deep Adaptive Design: Amortizing Sequential Bayesian Experimental Design. *Proceedings of the 38th International Conference on Machine Learning (ICML)* (2021).
- [10] A Foster, M Jankowiak, E Bingham, P Horsfall, YW Tee, T Rainforth, and N Goodman. 2019. Variational Bayesian optimal experimental design. *Conference on Neural Information Processing Systems*.
- [11] Peter I Frazier. 2018. A tutorial on Bayesian optimization. *arXiv preprint arXiv:1807.02811* (2018).

- [12] Yarin Gal, Riashat Islam, and Zoubin Ghahramani. 2017. Deep bayesian active learning with image data. In *International Conference on Machine Learning*. PMLR, 1183–1192.
- [13] Marta Garnelo, Jonathan Schwarz, Dan Rosenbaum, Fabio Viola, Danilo J Rezende, SM Eslami, and Yee Whye Teh. 2018. Neural processes. *arXiv preprint arXiv:1807.01622* (2018).
- [14] Daniel T Gillespie. 2007. Stochastic simulation of chemical kinetics. *Annu. Rev. Phys. Chem.* 58 (2007), 35–55.
- [15] Michael U Gutmann, Jukka Corander, et al. 2016. Bayesian optimization for likelihood-free inference of simulator-based statistical models. *Journal of Machine Learning Research* (2016).
- [16] Philipp Holl, Nils Thuerey, and Vladlen Koltun. 2019. Learning to Control PDEs with Differentiable Physics. In *International Conference on Learning Representations*.
- [17] Neil Houlsby, Ferenc Huszár, Zoubin Ghahramani, and Máté Lengyel. 2011. Bayesian active learning for classification and preference learning. *arXiv preprint arXiv:1112.5745* (2011).
- [18] Marko Järvenpää, Michael U Gutmann, Arijus Pleska, Aki Vehtari, and Pekka Marttinen. 2019. Efficient acquisition rules for model-based approximate Bayesian computation. *Bayesian Analysis* 14, 2 (2019), 595–622.
- [19] Hyunjik Kim, Andriy Mnih, Jonathan Schwarz, Marta Garnelo, Ali Eslami, Dan Rosenbaum, Oriol Vinyals, and Yee Whye Teh. 2019. Attentive neural processes. *International Conference on Learning Representation* (2019).
- [20] Diederik P Kingma and Max Welling. 2013. Auto-Encoding Variational Bayes. *arXiv preprint arXiv:1312.6114* (2013).
- [21] Andreas Kirsch, Joost Van Amersfoort, and Yarin Gal. 2019. Batchbald: Efficient and diverse batch acquisition for deep bayesian active learning. *Advances in neural information processing systems* 32 (2019).
- [22] Steven Kleinegesse and Michael U Gutmann. 2020. Bayesian experimental design for implicit models by mutual information neural estimation. In *International Conference on Machine Learning*. PMLR, 5316–5326.
- [23] Damien Lamberton and Bernard Lapeyre. 2007. *Introduction to stochastic calculus applied to finance*. CRC press.
- [24] Shibo Li, Robert M Kirby, and Shandian Zhe. 2020. Deep Multi-Fidelity Active Learning of High-dimensional Outputs. *arXiv preprint arXiv:2012.00901* (2020).
- [25] Yaguang Li, Rose Yu, Cyrus Shahabi, and Yan Liu. 2017. Diffusion convolutional recurrent neural network: Data-driven traffic forecasting. *arXiv preprint arXiv:1707.01926* (2017).
- [26] Yaguang Li, Rose Yu, Cyrus Shahabi, and Yan Liu. 2018. Diffusion Convolutional Recurrent Neural Network: Data-Driven Traffic Forecasting. In *International Conference on Learning Representations (ICLR)*.
- [27] Haoxing Lin, Rufan Bai, Weijia Jia, Xinyu Yang, and Yongjian You. 2020. Preserving dynamic attention for long-term spatial-temporal prediction. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. 36–46.
- [28] Christos Louizos, Xiahan Shi, Klammer Schutte, and Max Welling. 2019. The Functional Neural Process. *Advances in Neural Information Processing Systems* (2019).
- [29] Jose Lourenco, Robert Paton, Mahan Ghafari, Moritz Kraemer, Craig Thompson, Peter Simmonds, Paul Klennerman, and Sunetra Gupta. 2020. Fundamental principles of epidemic spread highlight the immediate need for large-scale serological surveys to assess the stage of the SARS-CoV-2 epidemic. *MedRxiv* (2020).
- [30] Jan-Matthis Lueckmann, Giacomo Bassetto, Theofanis Karaletsos, and Jakob H Macke. 2019. Likelihood-free inference with emulator networks. In *Symposium on Advances in Approximate Bayesian Inference*. PMLR, 32–53.
- [31] Edward Meeds and Max Welling. 2014. GPS-ABC: Gaussian process surrogate approximate Bayesian computation. In *Proceedings of the Thirtieth Conference on Uncertainty in Artificial Intelligence*. 593–602.
- [32] Song Mei and Andrea Montanari. 2019. The generalization error of random features regression: Precise asymptotics and double descent curve. (2019). *arXiv:1908.05355*.
- [33] Dina Mistry, Maria Litvinova, Ana Pastore y Piontti, Matteo Chinazzi, Laura Fumanelli, Marcelo FC Gomes, Syed A Haque, Quan-Hui Liu, Kunpeng Mu, Xinyue Xiong, et al. 2021. Inferring high-resolution human mixing patterns for disease modeling. *Nature communications* 12, 1 (2021), 1–12.
- [34] Andreas Munk, Adam Scibior, Atılım Güneş Baydin, Andrew Stewart, Goran Fernlund, Anoush Poursartip, and Frank Wood. 2019. Deep probabilistic surrogate networks for universal simulator approximation. *arXiv preprint arXiv:1910.11950* (2019).
- [35] Jay I Myung, Daniel R Cavagnaro, and Mark A Pitt. 2013. A tutorial on adaptive design optimization. *Journal of mathematical psychology* 57, 3-4 (2013), 53–67.
- [36] Louise Olsen-Kettle. 2011. Numerical solution of partial differential equations. *Lecture notes at University of Queensland, Australia* (2011).
- [37] George Papamakarios, Theo Pavlakou, and Iain Murray. 2017. Masked autoregressive flow for density estimation. *Advances in neural information processing systems* 30 (2017).
- [38] George Papamakarios, David Sterratt, and Iain Murray. 2019. Sequential neural likelihood: Fast likelihood-free inference with autoregressive flows. In *The 22nd International Conference on Artificial Intelligence and Statistics*. PMLR, 837–848.
- [39] Zhaozhi Qian, Ahmed M Alaa, and Mihaela van der Schaar. 2020. When and How to Lift the Lockdown? Global COVID-19 Scenario Analysis and Policy Assessment using Compartmental Gaussian Processes. *Advances in Neural Information Processing Systems* 33 (2020).
- [40] Syama Sundar Rangapuram, Matthias W Seeger, Jan Gasthaus, Lorenzo Stella, Yuyang Wang, and Tim Januschowski. 2018. Deep state space models for time series forecasting. *Advances in neural information processing systems* 31 (2018), 7785–7794.
- [41] Stephan Rasp, Michael S Pritchard, and Pierre Gentine. 2018. Deep learning to represent subgrid processes in climate models. *Proceedings of the National Academy of Sciences* 115, 39 (2018), 9684–9689.
- [42] Brian D Ripley. 2009. *Stochastic simulation*. Vol. 316. John Wiley & Sons.
- [43] Alvaro Sanchez-Gonzalez, Jonathan Godwin, Tobias Pfaff, Rex Ying, Jure Leskovec, and Peter Battaglia. 2020. Learning to simulate complex physics with graph networks. In *International Conference on Machine Learning*. PMLR, 8459–8468.
- [44] Bobak Shahriari, Kevin Swersky, Ziyu Wang, Ryan P Adams, and Nando De Freitas. 2015. Taking the human out of the loop: A review of Bayesian optimization. *Proc. IEEE* 104, 1 (2015), 148–175.
- [45] Xingjian Shi, Zhourong Chen, Hao Wang, Dit-Yan Yeung, Wai-Kin Wong, and Wang-chun Woo. 2015. Convolutional LSTM network: A machine learning approach for precipitation nowcasting. *Advances in neural information processing systems* 28 (2015).
- [46] Aditya Siddhant and Zachary C Lipton. 2018. Deep bayesian active learning for natural language processing: Results of a large-scale empirical study. *arXiv preprint arXiv:1808.05697* (2018).
- [47] Gautam Singh, Jaesik Yoon, Youngsung Son, and Sungjin Ahn. 2019. Sequential Neural Processes. *Advances in Neural Information Processing Systems* 32 (2019), 10254–10264.
- [48] Alan Mathison Turing. 1990. The chemical basis of morphogenesis. *Bulletin of mathematical biology* 52, 1 (1990), 153–197.
- [49] Benjamin T Vincent and Tom Rainforth. 2017. The DARC Toolbox: automated, flexible, and efficient delayed and risky choice experiments using Bayesian adaptive design. *PsyArXiv*. October 20 (2017).
- [50] Rui Wang, Karthik Kashinath, Mustafa Mustafa, Adrian Albert, and Rose Yu. 2020. Towards physics-informed deep learning for turbulent flow prediction. In *Proceedings of the 26th ACM SIGKDD international conference on Knowledge discovery and data mining*. ACM, 2020.
- [51] Yunbo Wang, Mingsheng Long, Jianmin Wang, Zhifeng Gao, and Philip S Yu. 2017. Predrnn: Recurrent neural networks for predictive learning using spatiotemporal lstrms. *Advances in neural information processing systems* 30 (2017).
- [52] Frank Wood, Andrew Warrington, Saeid Naderiparizi, Christian Weillbach, Vaden Masrani, William Harvey, Adam Scibior, Boyan Beronov, John Grefenstette, Duncan Campbell, et al. 2020. Planning as inference in epidemiological models. *arXiv preprint arXiv:2003.13221* (2020).
- [53] Huaxiu Yao, Xianfeng Tang, Hua Wei, Guanjie Zheng, and Zhenhui Li. 2019. Revisiting spatial-temporal similarity: A deep learning framework for traffic prediction. In *Proceedings of the AAAI conference on artificial intelligence*, Vol. 33. 5668–5675.
- [54] Huaxiu Yao, Fei Wu, Jintao Ke, Xianfeng Tang, Yitian Jia, Siyu Lu, Pinghua Gong, Jieping Ye, and Zhenhui Li. 2018. Deep multi-view spatial-temporal network for taxi demand prediction. In *Proceedings of the AAAI conference on artificial intelligence*, Vol. 32.
- [55] Lingxue Zhu and Nikolay Laptev. 2017. Deep and confident prediction for time series at uber. In *2017 IEEE International Conference on Data Mining Workshops (ICDMW)*. IEEE, 103–110.
- [56] Christoph Zimmer, Mona Meister, and Duy Nguyen-Tuong. 2018. Safe active learning for time-series modeling with gaussian processes. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*. 2735–2744.

A THEORETICAL ANALYSIS

A.1 Sample Efficiency of Active Learning

From the main text we know that in each round, the output random variable

$$X = \langle \Psi(\theta), z^* \rangle + \epsilon. \quad (4)$$

We further assume that the random noise ϵ is mean zero and σ -subGaussian.

Using this information, we treat z as an unknown parameter and define a likelihood function so that $p(X|z; \theta)$ has good coverage over the observations:

$$p(X_k|z; \theta_k) \propto \exp\left(-\frac{1}{2\sigma^2} (X_k - \langle \Psi(\theta_k), z \rangle)^2\right).$$

Let the prior distribution over z be:

$$p(z|\theta_k) = p(z) \propto \exp\left(-\frac{m}{2\sigma^2} \|z\|^2\right).$$

Here we use k instead of (i) in the Algorithm 1 to represent the number of iterations. We can form a posterior over z in the k -th round:

$$p(z|X_1, \theta_1, \dots, X_k, \theta_k) \propto \exp\left(-\frac{m}{2\sigma^2} \|z\|^2 - \frac{1}{2\sigma^2} \sum_{s=1}^k (X_s - \langle \Psi(\theta_s), z \rangle)^2\right).$$

Focusing on the random variable $z \sim p(\cdot|X_1, \theta_1, \dots, X_k, \theta_k)$, the estimate of the hidden variable, we can express it at k -th round as:

$$z_k = \hat{z}_k + \sigma V_k^{-1} \eta_k, \quad (5)$$

where $\hat{z}_k = V_k^{-1} \sum_{s=1}^k X_s \Psi(\theta_s)$, $V_k = mI + \sum_{s=1}^k \Psi(\theta_s) \Psi(\theta_s)^T$, and η_k is a standard normal random variable.

We can either choose action θ randomly or greedily. A random choice of θ corresponds to taking

$$\theta_k \sim \mathcal{N}(0, I), \quad (6)$$

A greedy procedure is to choose action θ_k in the k -th round to optimize $\text{KL}(p(z|\hat{x}, \theta) \| p(z)) = \mathbb{E}_{p(z|\hat{x}, \theta)} \left(\log \frac{p(z|\hat{x}, \theta)}{p(z)} \right)$, where we denote the estimated output variable $\hat{x}$ given θ and z as $\hat{x} = \langle \Psi(\theta), z \rangle$. This optimization procedure is equivalent to maximizing the variance of the prediction:

$$\theta_k = \arg \max_{\theta \in \mathbb{R}^d} \mathbb{E}_{z \sim p(\cdot|X_1, \theta_1, \dots, X_{k-1}, \theta_{k-1})} \left(\langle \Psi(\theta), z \rangle - \mathbb{E}_{z \sim p(\cdot|X_1, \theta_1, \dots, X_{k-1}, \theta_{k-1})} \langle \Psi(\theta), z \rangle \right)^2. \quad (7)$$

For both approaches, we assume that the features $\Psi(\theta)$ are normalized.

We compare the statistical risk of this approach with the random sampling approach.

Assume that the features are normalized, so that for all $\theta \in \mathbb{R}^d$, $\Psi(\theta) \in \mathbb{S}^{d-1}$. Define a matrix $A_k \in \mathbb{R}^{d \times k}$ containing all the column vectors $\{\Psi(\theta_1), \dots, \Psi(\theta_k)\}$. We can then express the estimation error in the following lemma.

LEMMA 1. *The estimation error $\|\hat{z}_k - z^*\|_2$ can be bounded as follow.*

$$\|\hat{z}_k - z^*\|_2 \leq m \left(m + \sigma_{\min}(A_k A_k^T) \right)^{-1} \cdot \|z^*\|_2 + \min \left\{ 1/(2\sqrt{m}), 1/\left(\sqrt{\sigma_{\min}(A_k A_k^T)} + \frac{m}{\sqrt{\sigma_{\min}(A_k A_k^T)}} \right) \right\} \cdot \sigma\sqrt{d}.$$

We analyze random sampling of θ versus greedy search for θ .

If the feature map $\Psi(\cdot) = \text{id}$, then from random matrix theory, we know that for θ randomly sampled from a normal distribution and normalized to $\|\theta\| = 1$, $\sigma_{\min}\left(\frac{1}{k} A_k A_k^T\right)$ will converge to $\left(\sqrt{1/k} - \sqrt{1/d}\right)^2$ for large k , which is of order $\Omega(1/d)$. This will lead to an appealing risk bound for $\|\hat{z}_k - z^*\|_2$ on the order of $O(d/\sqrt{k})$.

However, in high dimension, this feature map is often far from identity. In the proof of Theorem 1 below, we demonstrate that even when $\Psi(\cdot)$ is simply a linear random feature map, with i.i.d. normal entries, random exploration in θ can lead to deteriorated error bound. This setting is motivated by the analyses in wide neural networks, where the features learned from gradient descent are close to those generated from random initialization [8, 32].

THEOREM 1 (FORMAL STATEMENT). *Assume that the noise ϵ in equation 4 is σ -subGaussian.*

For a normalized linear random feature map $\Psi(\cdot)$, greedily optimizing the KL divergence, $\text{KL}(p(z|\hat{x}, \theta) \| p(z))$ (or equivalently the variance of the posterior predictive distribution defined in equation 7) in search of θ will lead to an error $\|\hat{z}_k - z^\|_2 = O(\sigma d/\sqrt{k})$ with high probability.*

On the other hand, random sampling of θ following equation 6 will lead to $\|\hat{z}_k - z^\|_2 = O(\sigma d^2/\sqrt{k})$ with high probability.*

See proofs in the Appendix C.

B ADDITIONAL RESULTS

B.1 INP, GP, and SNL Model

Table 2 and Table 3 show the average results together with the standard deviation of INP, GP, and SNL model for SEIR simulator after running experiments three times. The performance of INP with the proposed LIG is much better than GP and SNL baselines at each iteration.

B.2 Active learning performance comparison.

Table 4, Table 5, and Table 6 show the active learning performance comparison results on Heat, RD, and LEAM-US simulation task. The performance of INP with the proposed LIG always outperforms the baselines at each iteration among all 3 tasks.

C OTHER INFORMATION

The proof of Proposition 1, the proof of Theorem A.1 and experiment details can be found at <https://arxiv.org/pdf/2106.02770.pdf>.

Table 2: Active learning Performance comparison using MAE for different acquisition functions in NP model on SEIR simulator

Percentage of samples	LIG	EIG	Random	MeanSTD	MaxEntropy
1.11%	365.87 $\pm$ 142.87	435.08 $\pm$ 32.38	480.68 $\pm$ 5.24	480.22 $\pm$ 12.63	427.73 $\pm$ 61.36
1.85%	236.9 $\pm$ 50.6	340.27 $\pm$ 30.84	398.33 $\pm$ 131.05	314.75 $\pm$ 111.42	302.24 $\pm$ 119.84
2.96%	119.26 $\pm$ 14.22	291.15 $\pm$ 10.60	244.27 $\pm$ 148.89	158.94 $\pm$ 36.6	186.88 $\pm$ 57.48
4.07%	96.73 $\pm$ 17.07	261.60 $\pm$ 7.78	116.8 $\pm$ 9.1	127.36 $\pm$ 27.97	146.72 $\pm$ 26.06

Table 3: Active learning performance comparison using MAE for different acquisition functions in GP model and SNL model on SEIR simulator

Percentage of samples	Random (GP)	MeanSTD (GP)	MaxEntropy (GP)	SNL
1.11%	663.76 $\pm$ 46.36	606.81 $\pm$ 6.89	586.25 $\pm$ 58.44	707.61 $\pm$ 44.42
1.85%	637.12 $\pm$ 13.45	619.15 $\pm$ 36.42	628.54 $\pm$ 71.34	669.03 $\pm$ 73.19
2.96%	597.3 $\pm$ 19.59	589.72 $\pm$ 24.9	568.84 $\pm$ 19.05	668.67 $\pm$ 72.42
4.07%	519.98 $\pm$ 17.86	530.07 $\pm$ 32.95	578.34 $\pm$ 68.7	685.28 $\pm$ 53.00

Table 4: Active learning performance comparison using MAE for different acquisition functions in STNP model on Heat simulator.

Percentage of samples	LIG	EIG	Random	MeanSTD	MaxEntropy
5.21%	1.55e-2 $\pm$ 1.9e-3	1.64e-2 $\pm$ 1.9e-3	1.74e-2 $\pm$ 2.2e-3	1.77e-2 $\pm$ 2e-3	1.83e-2 $\pm$ 1.1e-3
7.81%	1.33e-2 $\pm$ 1.7e-3	1.56e-2 $\pm$ 1.1e-3	1.49e-2 $\pm$ 3.3e-3	1.60e-2 $\pm$ 3.7e-3	1.58e-2 $\pm$ 4e-4
10.42%	1.02e-2 $\pm$ 3.4e-3	1.38e-2 $\pm$ 1.6e-3	1.23e-2 $\pm$ 1.5e-3	1.30e-2 $\pm$ 2.4e-3	1.61e-2 $\pm$ 8e-4
13.02%	9.3e-3 $\pm$ 3.1e-3	1.05e-2 $\pm$ 2.2e-3	1.14e-2 $\pm$ 1.3e-3	1.27e-2 $\pm$ 2.5e-3	1.46e-2 $\pm$ 6e-4
15.62%	6.7e-3 $\pm$ 5e-4	1.08e-2 $\pm$ 1.8e-3	1.05e-2 $\pm$ 9e-4	1.17e-2 $\pm$ 1.2e-3	1.43e-2 $\pm$ 2e-4

Table 5: Active learning performance comparison using MAE for different acquisition functions in STNP model on RD simulator

Percentage of samples	LIG	EIG	Random	MeanSTD	MaxEntropy
6.25%	4.562 $\pm$ 0.114	4.861 $\pm$ 0.433	5.325 $\pm$ 0.361	5.264 $\pm$ 0.298	4.826 $\pm$ 0.336
12.50%	3.841 $\pm$ 0.253	4.590 $\pm$ 0.529	4.179 $\pm$ 0.045	4.157 $\pm$ 0.252	4.084 $\pm$ 0.042
18.75%	3.165 $\pm$ 0.142	4.162 $\pm$ 0.696	3.602 $\pm$ 0.182	3.675 $\pm$ 0.229	3.694 $\pm$ 0.140
25.00%	2.415 $\pm$ 0.083	3.993 $\pm$ 0.847	3.140 $\pm$ 0.165	3.339 $\pm$ 0.111	3.302 $\pm$ 0.284
31.25%	2.302 $\pm$ 0.007	3.714 $\pm$ 0.861	2.561 $\pm$ 0.243	2.791 $\pm$ 0.072	2.912 $\pm$ 0.473

Table 6: Active learning performance comparison using MAE for different acquisition functions in STNP model on LEAM-US simulator, population divided by 1000.

Percentage of samples	LIG	EIG	Random	MeanSTD	MaxEntropy
11.1%	14.447 $\pm$ 1.087	19.067 $\pm$ 3.981	20.961 $\pm$ 5.548	35.356 $\pm$ 28.706	65.498 $\pm$ 13.324
13.7%	11.704 $\pm$ 0.216	16.372 $\pm$ 3.663	13.418 $\pm$ 0.815	16.092 $\pm$ 3.11	30.496 $\pm$ 24.333
21.3%	7.593 $\pm$ 0.822	11.754 $\pm$ 1.713	9.332 $\pm$ 0.601	11.191 $\pm$ 0.184	10.028 $\pm$ 2.065
28.9%	6.539 $\pm$ 0.618	9.455 $\pm$ 0.595	8.077 $\pm$ 0.657	7.908 $\pm$ 0.536	8.417 $\pm$ 0.616
36.5%	6.008 $\pm$ 1.079	8.596 $\pm$ 0.741	6.719 $\pm$ 0.383	7.533 $\pm$ 0.861	7.431 $\pm$ 0.776