

Full Length Article

A new local time-decoupled squared Wasserstein-2 method for training stochastic neural networks to reconstruct uncertain parameters in dynamical systems

Mingtao Xia^{a,*}, Qijing Shen^b, Philip K. Maini^c, Eamonn A. Gaffney^c, Alex Mogilner^d

^a Department of Mathematics, University of Houston, Philip Guthrie Hoffman Hall, 3551 Cullen Blvd, Houston, TX, 77204, USA

^b Nuffield Department of Medicine, University of Oxford, Henry Wellcome Building for Molecular Physiology, Old Road, Oxford, OX3 7BN, United Kingdom

^c Mathematical Institute, University of Oxford, Radcliffe Observatory, Andrew Wiles Building, Oxford, OX2 6GG, United Kingdom

^d Courant Institute of Mathematical Sciences, New York University, 251 Mercer Street, New York, NY, 10012, USA

ARTICLE INFO

Keywords:

Uncertainty quantification
Stochastic neural network
Wasserstein distance
Dynamical systems

ABSTRACT

In this work, we propose and analyze a new local time-decoupled squared Wasserstein-2 method for reconstructing the distribution of unknown parameters in dynamical systems from a finite number of observed temporal trajectories. Specifically, we show that a stochastic neural network model, which can be effectively trained by minimizing our proposed local time-decoupled squared Wasserstein-2 loss function, is an effective model for approximating the distribution of uncertain model parameters in dynamical systems. Through several numerical examples, we showcase the effectiveness of our proposed method in reconstructing the distribution of parameters in different dynamical systems.

1. Introduction

The inverse problem of reconstructing a noisy dynamical system from time-series data finds wide applications across different disciplines. Efficient algorithms to solve such inverse-type problems advance different fields including inferring neural circuit dynamics from spiking data (Pillow et al., 2008) in neuroscience, modeling and predicting complex weather patterns from historical data (Carrassi et al., 2018) in climate science, uncovering disease transmission dynamics from infection case counts over time (Roda et al., 2020) in epidemiology, and deducing reaction rates from experimental concentration-time profiles in reaction kinetics in biochemistry (Loskot et al., 2019). However, such inverse-type problems pose substantial mathematical and computational challenges, particularly when data are limited and noisy, motivating ongoing research into novel algorithms and theoretical frameworks to improve model reconstruction accuracy and efficiency.

In this paper, we study the inverse problem of inferring the distribution of model parameters for several dynamical systems including ordinary differential equations (ODEs), partial differential equations (PDEs), and stochastic differential equations (SDEs) from time-series data or spatiotemporal data. Existing methods for such problems can be broadly categorized into traditional statistical approaches and modern data-driven techniques. Traditional statistical methods often in-

volve parameter estimation frameworks. For example, linear and non-linear regression methods play a role in simpler systems where the functional form of the model is partially known (Chib & Greenberg, 2002). Furthermore, maximum likelihood estimation and Bayesian inference methods (Cresswell, 2015; Marin et al., 2021) are often adopted. Maximum likelihood estimation optimizes the likelihood of model parameter values in a proposed model from observed data, while Bayesian methods incorporate prior information and compute posterior distributions. These approaches are widely used in applications such as reaction network reconstruction and epidemiological modeling. On the other hand, data-driven methods, leveraging advances in machine learning, offer a complementary toolkit for inverse problems. For example, neural networks and reservoir computing frameworks have been successful in reconstructing chaotic systems and inferring governing equations directly from data (Brunton et al., 2016; Pathak et al., 2018). Sparse identification of nonlinear dynamics (SINDy) has emerged as a powerful tool for discovering interpretable dynamical systems by identifying a parsimonious set of governing equations from time-series data (Brunton et al., 2016). Gaussian process regression has proven effective for nonparametric inference, especially in uncertainty quantification (Raissi et al., 2017). Hybrid approaches, which integrate data-driven and traditional statistical methods such as the physics-informed neural network method (Raissi et al., 2019), are also gathering increasing attention as they take

* Corresponding author.

E-mail address: mxia4@uh.edu (M. Xia).

<https://doi.org/10.1016/j.neunet.2025.107893>

Received 6 March 2025; Received in revised form 9 July 2025; Accepted 20 July 2025

Available online 5 August 2025

0893-6080/© 2025 Elsevier Ltd. All rights are reserved, including those for text and data mining, AI training, and similar technologies.

advantage of both statistical methods and data-driven methods and provide more interpretable machine-learning tools.

The Wasserstein distance, which can effectively measure the discrepancy between two probability distributions (Balasubramanian et al., 2024; Panaretos & Zemel, 2019; Villani, 2009), was utilized for various uncertainty quantification (UQ) tasks. In (Bernton et al., 2019), the Wasserstein distance was proposed to estimate model parameters that govern a probabilistic model. Compared to empirical maximum likelihood estimates, using the Wasserstein distance to estimate model parameters in uncertainty models can be more efficient and accurate (Blanchet & Kang, 2021). However, previous methods mainly focused on point estimates of model parameters (inferring the exact values of the parameters) and could not quantify the uncertainty in the model parameters, i.e., they cannot reconstruct a distribution of the model parameters from data. Recently, a time-decoupled Wasserstein-2 (W_2) distance method, which compares the distributions of two stochastic processes across different time points, has been revealed to be an efficient loss function for reconstructing intrinsically noisy stochastic processes such as pure-diffusion processes and jump-diffusion processes using parameterized neural networks (Xia et al., 2024a,b). Furthermore, in Xia and Shen (2024), a local squared W_2 method, which adopts a “neighborhood technique” to enlarge the amount of available data, was proposed to infer the distribution of y given observed data x in the uncertainty model $y = f(x, \omega)$ where ω are uncertain latent unobserved variables. However, for the specific problem of inferring unknown parameters in deterministic or stochastic dynamical models, these two methods are not suitable: the time-decoupled W_2 distance method assumes that there is no uncertainty in the underlying model (i.e., the ODE or SDE to be reconstructed is fixed), and the local squared W_2 method is not directly applicable to parameter inference problems. Two major difficulties arise: first, an efficient model is required to approximate the distribution of parameters in a dynamical system; second, it is necessary to distinguish uncertainty in model parameters from uncertainty in the initial state of the dynamical systems and intrinsic stochasticity in the dynamical system, such as the Wiener process in SDEs.

In this work, we propose and analyze a local time-decoupled squared W_2 method, which builds upon the time-decoupled squared W_2 method (Xia et al., 2024b) for efficiently reconstructing the underlying dynamics from noisy time-series data and the local squared W_2 method (Xia & Shen, 2024) for handling the uncertainty in the initial state. This in turn is implemented to infer the distribution of model parameters in deterministic or stochastic dynamical systems given a finite number of observations. As an illustration, consider the following ODE:

$$dX(t; \theta) = f(X(t; \theta), t; \theta)dt, \quad X \in \mathbb{R}^d, \quad t \in [0, T], \quad X(0) \sim \nu_0, \quad \theta \sim \mu, \quad (1.1)$$

where $\theta \in \mathbb{R}^\ell$ is a continuous random variable representing uncertain parameters in the ODE and $f: \mathbb{R}^{d+1+\ell} \rightarrow \mathbb{R}^d$. ν_0 is a probability measure defined on the Borel σ -algebra $\mathcal{B}(\mathbb{R}^d)$. For each realization of the ODE (1.1), θ is sampled independently. Thus, we obtain different trajectories by solving the ODE (1.1) multiple times due to uncertainty in model parameters. We use another ODE model as an approximation to Eq. (1.1):

$$d\hat{X}(t; \hat{\theta}) = f(\hat{X}(t; \hat{\theta}), t; \hat{\theta})dt, \quad \hat{X} \in \mathbb{R}^d, \quad t \in [0, T], \quad \hat{X}(0) = X_0, \quad \hat{\theta} \sim \hat{\mu}. \quad (1.2)$$

Here, $\hat{\theta} \in \mathbb{R}^\ell$ is another continuous random variable and denotes uncertain parameters in the approximate ODE. We aim to construct the probability density function $\hat{\mu}$, i.e. the distribution of $\hat{\theta} \in \mathbb{R}^\ell$, such that $\hat{\mu}$ can match μ well in Eq. (1.1). Specifically, we show that minimizing our proposed loss function can lead to efficient training of a stochastic neural network (SNN) model with weight uncertainty.

SNNs are effective in uncertainty quantification, generative modeling, and time-series analysis (Gal & Ghahramani, 2016; Rezende et al., 2014; Senan et al., 2017; Vadivel et al., 2020). Instead of deterministic outputs, SNNs introduce randomness in the weights and/or biases (Blundell et al., 2015; Yu et al., 2021) or utilize stochastic neurons whose outputs are binary (Tang & Salakhutdinov, 2013). Theoretical results exist

for the approximation errors for certain types of SNNs, such as the approximation error of a single-layer SNN (Gonon et al., 2023) as well as the universal approximation ability of dropout neural networks (Manita et al., 2022). In this work, we prove that the SNN model we use can serve as an effective model to approximate a general random field model in the W_2 metric under moderate assumptions. Our result generalizes the universal approximation ability property of deterministic multilayer feedforward neural networks for approximating deterministic functions (Hornik et al., 1989; Leshno et al., 1993) to SNNs for approximating random fields. As a special case, when the input of the SNN is fixed, we show that the output of the SNN can approximate any continuous random variable under moderate assumptions, making it an ideal surrogate model for reconstructing the distribution of unknown parameters in dynamical systems. Compared with traditional Bayesian methods, our proposed approach has the advantage of not requiring any knowledge or prior distributions of model parameters and directly reconstructing model parameter distributions from time-series data. Compared to other data-driven methods, such as Bayesian neural networks (Neal, 2012), generative modeling methods (Böhm et al., 2019), and Wasserstein generative adversarial networks (WGANs) that train a generator and a discriminator (Arjovsky et al., 2017; Boukraichi et al., 2022), our method is more physics-informed and provides more insights and interpretability of the dynamical system. Our method directly outputs the distribution of unknown parameters governing the dynamical system. Furthermore, our method does not require deep neural networks and we shall show that shallow SNNs with hundreds of neurons are capable of reconstructing the joint distribution of several model parameters when inputting only a few hundred trajectories as training data.

The main contributions of our work are as follows:

- We propose and analyze a local time-decoupled squared W_2 method for the direct reconstruction of model parameters in specific dynamical systems including ODEs, SDEs, and PDEs from time-series or spatiotemporal data. Our method takes into account both uncertainties in the initial state as well as intrinsic fluctuations, e.g. Wiener processes, of the dynamical system when reconstructing model parameters.
- We analyze an SNN model whose weights are sampled from independent normal distributions. We prove that this SNN model can approximate any multidimensional continuous random variable under moderate assumptions. Furthermore, this SNN model can be trained by direct minimization of our local time-decoupled squared W_2 loss function.
- Through numerical experiments, we showcase the effectiveness of our proposed method for reconstructing the distribution of model parameters in several deterministic and stochastic dynamical systems.

The structure of this paper is as follows. In Section 2, we analyze a local time-decoupled squared W_2 loss function for reconstructing model parameters in dynamical systems. In Section 3, we prove that an SNN model can approximate a continuous random variable in the squared W_2 sense, which makes this SNN model an ideal approximate model for reconstructing the distribution of uncertain parameters in dynamical systems. In Section 4, we carry out numerical experiments and showcase the effectiveness of training the SNN model by minimizing our proposed local function on the reconstruction of uncertain model parameters in different dynamical systems. In Section 5, we summarize our results and propose potential future research directions.

2. A local time-decoupled squared W_2 loss function

In this section, we propose and analyze a local time-decoupled squared W_2 method for reconstructing the distribution of uncertain parameters in specific dynamical systems. Our local time-decoupled squared W_2 method integrates the time-decoupled squared W_2 -distance method proposed in Xia et al. (2024a,b) and the local squared W_2 -distance method in Xia and Shen (2024). Therefore, both intrinsic

fluctuations in the dynamical systems and uncertainties in the initial condition can be taken into account. First, we introduce the W_2 distance between distributions associated with two multidimensional random variables.

Definition 2.1. For two d -dimensional random variables

$$X = (X_1, \dots, X_d), \quad \hat{X} = (\hat{X}_1, \dots, \hat{X}_d) \in \mathbb{R}^d \quad (2.1)$$

with associated probability measures $\nu, \hat{\nu}$, respectively, the W_2 -distance $W_2(\nu, \hat{\nu})$ between their probability measures is defined as

$$W_2(\nu, \hat{\nu}) := \inf_{\pi(\nu, \hat{\nu})} \mathbb{E}_{(X, \hat{X}) \sim \pi(\nu, \hat{\nu})} [\|X - \hat{X}\|^2]^{\frac{1}{2}}. \quad (2.2)$$

In Eq. (2.2) and throughout this paper, the norm $\|\cdot\|$ denotes the l^2 norm of a vector: $\|X\| := \left(\sum_{i=1}^d (X_i)^2\right)^{\frac{1}{2}}$. $\pi(\nu, \hat{\nu})$ is a coupling probability measure which iterates over all coupled distributions of $X(t), \hat{X}(t)$, defined by the condition:

$$\begin{cases} P_{\pi(\nu, \hat{\nu})}(A \times \mathbb{R}^d) = P_\nu(A), \\ P_{\pi(\nu, \hat{\nu})}(\mathbb{R}^d \times A) = P_{\hat{\nu}}(A), \end{cases} \quad \forall A \in \mathcal{B}(\mathbb{R}^d), \quad (2.3)$$

where $\mathcal{B}(\mathbb{R}^d)$ denotes the Borel σ -algebra associated with the space of d -dimensional functions in $\mathbb{R}^d$.

Next, we define the local squared W_2 distance for the probability measures associated with the trajectories of two dynamical systems, which builds upon the local squared W_2 distance introduced in Xia and Shen (2024).

Definition 2.2. The local squared W_2 distance between the probability measures associated with two dynamical systems $\{X(t)\}_{t \in [0, T]}, \{\hat{X}(t)\}_{t \in [0, T]}$ at a specific time t is defined by:

$$W_{2, \delta}^{2, e}(X(t), \hat{X}(t)) := \int_{\mathbb{R}^d} W_2^2\left(\nu_{X_0, \delta}^e(t), \hat{\nu}_{\hat{X}_0, \delta}^e(t)\right) \nu_0^e(dX_0). \quad (2.4)$$

Here, $\nu_0^e(\cdot)$ is the empirical distribution of the initial condition for $X(t)$ and $\hat{X}(t)$. $\nu_{X_0, \delta}^e(t)$ and $\hat{\nu}_{\hat{X}_0, \delta}^e(t)$ are the empirical conditional probability distributions of $X(t)$ and $\hat{X}(t)$ at time t conditioned on $\|X(0) - X_0\| \leq \delta$ and $\|\hat{X}(0) - \hat{X}_0\| \leq \delta$ for a given initial state X_0 , respectively.

Now, we define the local time-decoupled squared W_2 distance between the probability measures associated with trajectories of two dynamical systems.

Definition 2.3. Let $0 = t_0 < t_1 < \dots < t_n = T$ be a time discretization mesh in $[0, T]$. The local time-decoupled squared W_2 distance between the distributions associated with two dynamical systems $\{X\}_{t \in [0, T]}, \{\hat{X}\}_{t \in [0, T]}$ at time t is defined by:

$$\begin{aligned} \bar{W}_{2, \delta}^{2, e}(X, \hat{X}) &:= \int_0^T W_{2, \delta}^{2, e}(X(t), \hat{X}(t)) dt \\ &= \lim_{n \rightarrow \infty, \max(t_{i+1} - t_i) \rightarrow 0} \sum_{i=0}^{n-1} W_{2, \delta}^{2, e}(X(t_i), \hat{X}(t_i))(t_{i+1} - t_i), \end{aligned} \quad (2.5)$$

where $W_{2, \delta}^{2, e}(X(t), \hat{X}(t))$ is the local squared W_2 distance in Definition 2.2. Empirically, we use

$$\bar{W}_{2, \delta}^{2, e}(X, \hat{X}) \approx \sum_{i=0}^{n-1} W_{2, \delta}^{2, e}(X(t_i), \hat{X}(t_i))(t_{i+1} - t_i), \quad (2.6)$$

as an approximation to Eq. (2.5) for numerically calculating the local time-decoupled squared W_2 distance loss function. Under some technical conditions, the approximation error of using the time-discretized RHS of Eq. (2.6) to approximate the local time-decoupled squared W_2 distance loss function in Eq. (2.5) is $O(\max_{i=1}^{n-1} (t_{i+1} - t_i))$ for ODEs and $O(\sqrt{\max_{i=0}^{n-1} (t_{i+1} - t_i)})$ for jump-diffusion processes, which will be analyzed in the proof of Theorem 2.1 for ODEs and in the proof of Corollary 2.1 for jump-diffusion processes, respectively.

The loss function Eq. (2.6) was also used to reconstruct the dynamics of an ODE using a parameterized neural network in Xia and Shen (2024). Yet, there is no understanding of why minimizing the loss function Eq. (2.6) leads to the successful reconstruction of the distribution of parameters underlying a dynamical system. Additionally, intrinsic stochasticity, such as Wiener processes in stochastic dynamical systems, was not considered in Xia and Shen (2024). In this section, we shall show that: i) the local time-decoupled squared W_2 distance in Eq. (2.5) is well-defined in several deterministic or stochastic dynamical systems and ii) minimizing the loss function in Eq. (2.6) is a necessary condition for the reconstruction of the distribution of parameters in dynamical systems.

First, we prove that the local time-decoupled squared W_2 distance is well-defined in some typical dynamical systems including ODEs and certain SDEs. We can prove that the local squared W_2 distance between the probability measures associated with $\{X\}_{t \in [0, T]}$ and $\{\hat{X}\}_{t \in [0, T]}$, which are trajectories generated by solving the two ODEs (1.1) and (1.2), is well-defined. In Eqs. (1.1), (1.2), and the models we study below, we assume that the model parameters are independent of the initial condition, e.g., in Eqs. (1.1), (1.2) θ is independent of $X(0)$ and $\hat{\theta}$ is independent of $\hat{X}(0)$.

Theorem 2.1. Suppose

$$\sup_{X(0), \theta, t} \|X(t)\|^2 \leq X, \quad \sup_{\hat{X}(0), \hat{\theta}, t} \|\hat{X}(t)\|^2 \leq \hat{X} \quad (2.7)$$

are uniformly bounded, where $X(t)$ and $\hat{X}(t)$ are solutions to the ODEs (1.1) and (1.2), respectively. Furthermore, we assume that f is continuous and uniformly bounded. Then, the limit

$$\lim_{\max(t_{i+1} - t_i) \rightarrow 0} \sum_{i=0}^{n-1} W_{2, \delta}^{2, e}(X(t_i), \hat{X}(t_i))(t_{i+1} - t_i) \quad (2.8)$$

on the RHS of Eq. (2.5) exists.

We provide a proof of Theorem 2.1 in Appendix A, which is similar to the proof of Theorem 3.1 in Xia et al. (2024a). Theorem 2.1 can be extended to reveal that the local time-decoupled squared W_2 distance between probability measures associated with two noisy jump-diffusion processes is also well-defined. We can prove the following corollary.

Corollary 2.1. Consider the following two d -dimensional jump-diffusion processes:

$$\begin{aligned} dX(t) &= f(X(t), t; \theta) dt + \sigma(X(t), t; \theta) dB_t \\ &\quad + \int_U \beta(X(t), \xi, t; \theta) \tilde{N}(dt, \gamma(d\xi)), \quad X(0) \sim \nu_0 \end{aligned} \quad (2.9)$$

and

$$\begin{aligned} d\hat{X}(t) &= f(\hat{X}(t), t; \hat{\theta}) dt + \sigma(\hat{X}(t), t; \hat{\theta}) d\hat{B}_t \\ &\quad + \int_U \beta(\hat{X}(t), \xi, t; \hat{\theta}) \hat{N}(dt, \gamma(d\xi)), \quad \hat{X}(0) = X(0). \end{aligned} \quad (2.10)$$

In Eq. (2.9), $f : \mathbb{R}^{d+\ell+1} \rightarrow \mathbb{R}^d$ and $\sigma : \mathbb{R}^{d+\ell+1} \rightarrow \mathbb{R}^{d \times m}$ denote the drift and diffusion functions of the SDE, respectively; ν_0 is a probability measure defined on the Borel σ -algebra $\mathcal{B}(\mathbb{R}^d)$; $B(t)$ represents an m -dimensional standard Brownian motion; $\tilde{N}(dt, \gamma(d\xi))$ is a compensated Poisson process independent of B_t , defined as follows:

$$\tilde{N}(dt, \gamma(d\xi)) := N(dt, \gamma(d\xi)) - \gamma(d\xi) dt, \quad (2.11)$$

where $N(dt, \gamma(d\xi))$ is a Poisson process with intensity $\gamma(d\xi) dt$, and $\gamma(d\xi)$ is a measure defined on $U \subseteq \mathbb{R}$, the measure space of the Poisson process. $\hat{N}(dt, \gamma(d\xi))$ is another compensated Poisson process of intensity $\gamma(d\xi) dt$ and independent of $B_t, \tilde{N}_t$ in Eq. (2.9) and $\hat{B}_t$ in Eq. (2.10). $\theta, \hat{\theta} \in \mathbb{R}^\ell$ are uncertain model parameters. We assume that f, σ, β are continuous and uniformly bounded. Then

$$\bar{W}_{2, \delta}^{2, e}(X, \hat{X}) := \lim_{\max(t_{i+1} - t_i) \rightarrow 0} \sum_{i=0}^{n-1} W_{2, \delta}^{2, e}(X(t_i), \hat{X}(t_i))(t_{i+1} - t_i) \quad (2.12)$$

exists. Furthermore,

$$\left| \sum_{i=0}^{n-1} W_{2,\delta}^{2,e}(X(t_i), \hat{X}(t_i))(t_{i+1} - t_i) - \tilde{W}_{2,\delta}^{2,e}(X, \hat{X}) \right| \leq 2(X + \hat{X})T \max_i \left(\sqrt{F_i \Delta t + \Sigma_i + B_i} + \sqrt{\hat{F}_i \Delta t + \hat{\Sigma}_i + \hat{B}_i} \right), \quad (2.13)$$

where $\Delta t := \max_{i=0}^{n-1} (t_{i+1} - t_i)$, and $v(t_i)$ and $\hat{v}(t_i)$ are the probability measures of $X(t_i)$ and $\hat{X}(t_i)$, respectively. In Eq. (2.13),

$$\begin{aligned} F_i &:= \sup_{X_0, \theta} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d f_{\ell}^2(X(t^-), t^-; \theta) dt \right], \\ \hat{F}_i &:= \sup_{X_0, \hat{\theta}} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d f_{\ell}^2(\hat{X}(t^-), t^-; \hat{\theta}) dt \right], \\ \Sigma_i &:= \sup_{X_0, \theta} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d \sum_{j=1}^m \sigma_{\ell,j}^2(X(t^-), t^-; \theta) dt \right], \\ \hat{\Sigma}_i &:= \sup_{X_0, \hat{\theta}} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d \sum_{j=1}^m \hat{\sigma}_{\ell,j}^2(\hat{X}(t^-), t^-; \hat{\theta}) dt \right], \\ B_i &:= \sup_{X_0, \theta} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d \int_U \beta_{\ell}^2(X(t^-), \xi, t^-; \theta) \gamma(d\xi) dt \right], \\ \hat{B}_i &:= \sup_{X_0, \hat{\theta}} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d \int_U \hat{\beta}_{\ell}^2(\hat{X}(t^-), \xi, t^-; \hat{\theta}) \gamma(d\xi) dt \right], \\ X &:= \sup_{X_0, \theta, t} \mathbb{E}[\|X(t)\|^2]^{\frac{1}{2}}, \quad \hat{X} := \sup_{X_0, \hat{\theta}, t} \mathbb{E}[\|\hat{X}(t)\|^2]^{\frac{1}{2}}, \end{aligned} \quad (2.14)$$

where $f_{\ell}(X(t^-), t^-; \theta)$, $\sigma_{\ell,j}(X(t^-), t^-; \theta)$, and $\beta_{\ell}(X(t^-), \xi, t^-; \theta)$ refer to the left-hand limits:

$$\begin{aligned} f_{\ell}(X(t^-), t^-; \theta) &= \lim_{s \rightarrow t, s < t} f_{\ell}(X(s), s; \theta), \\ \sigma_{\ell,j}(X(t^-), t^-; \theta) &= \lim_{s \rightarrow t, s < t} \sigma_{\ell,j}(X(s), s; \theta), \\ \beta_{\ell}(X(t^-), t^-, \xi; \theta) &= \lim_{s \rightarrow t, s < t} \beta_{\ell}(X(s), s, \xi; \theta). \end{aligned} \quad (2.15)$$

The proof of Corollary 2.1 is in Appendix B. Next, we show that the local time-decoupled squared W_2 distance $\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})$ in Definition 2.3 can be bounded by the squared W_2 distance between the two probability measures associated with the two sets of parameters θ and $\hat{\theta}$ in Eqs. (1.1) and (1.2), which implies the necessity of minimizing the local time-decoupled squared W_2 distance $\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})$ if we wish to match the distribution of θ using the distribution of the reconstructed $\hat{\theta}$.

Theorem 2.2. Suppose the drift, diffusion, and jump functions in the two jump-diffusion processes Eqs. (2.9) and (2.10) satisfy the following Lipschitz condition:

$$\begin{aligned} \sum_{i=1}^d |f_i(X, t; \theta) - f_i(\hat{X}, t; \hat{\theta})| &\leq C(\|X - \hat{X}\| + \|\theta - \hat{\theta}\|), \\ \sum_{i=1}^d |\sigma_{i,j}(X, t; \theta) - \sigma_{i,j}(\hat{X}, t; \hat{\theta})| &\leq C(\|X - \hat{X}\| + \|\theta - \hat{\theta}\|), \quad j = 1, \dots, m, \\ \sum_{i=1}^d |\beta_i(X, \xi, t; \theta) - \beta_i(\hat{X}, \xi, t; \hat{\theta})| &\leq C(\|X - \hat{X}\| + \|\theta - \hat{\theta}\|), \\ \forall X, \hat{X} \in \mathbb{R}^d, \forall \theta, \hat{\theta} \in \mathbb{R}^{\ell}, \quad C < \infty. \end{aligned} \quad (2.16)$$

In Eq. (2.16), f_i and β_i denote the i th component of f and β , respectively, and $\sigma_{i,j}$ denotes the (i, j) element of the matrix σ in Eq. (2.9). Furthermore, we assume that the assumptions in Corollary 2.1 hold and the sixth-order moments

$$\mathbb{E}[\|\theta\|^6] \leq \Theta_6, \quad \mathbb{E}[\|\hat{\theta}\|^6] \leq \hat{\Theta}_6 \quad (2.17)$$

are uniformly bounded. Then, $\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})$ can be bounded by the squared W_2 distance $W_2^2(\mu, \hat{\mu})$:

$$\begin{aligned} \mathbb{E}[\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})] &\leq 8C_0 T \delta^2 \exp(C_0 T) + \frac{6C_1}{C_0} T \exp(C_0 T) \\ &\times \left(W_2^2(\mu, \hat{\mu}) + 2C_3 \mathbb{E} \left[h(N^{\#}(X_0; \delta), \ell) \cdot \left(\Theta_6^{\frac{1}{3}} + \hat{\Theta}_6^{\frac{1}{3}} \right) \right] \right), \end{aligned} \quad (2.18)$$

where μ and $\hat{\mu}$ are the probability measures associated with θ and $\hat{\theta}$, respectively. In Eq. (2.18), C_0, C_1, C_2 are three constants, and

$$h(N, \ell) := \begin{cases} 2N^{-\frac{1}{2}}(\log(1+N)+1), & \ell \leq 4, \\ 2N^{-\frac{2}{\ell}}, & \ell > 4. \end{cases} \quad (2.19)$$

In Eq. (2.18), $N^{\#}(X_0; \delta)$ refers to the number of trajectories in the data set such that their initial conditions satisfy $\|X(0) - X_0\| \leq \delta$.

We prove Theorem 2.2 in Appendix C. Theorem 2.2 implies that minimizing the local time-decoupled squared W_2 distance $\tilde{W}_{2,\delta}^{2,e}$ is a necessary condition if we wish to match the distribution of θ with the distribution of $\hat{\theta}$. In Eq. (2.18), there is a trade-off between the first term and the third term on the RHS: if we increase δ , then the first term on the RHS of Eq. (2.18) will increase but the factor $h(N^{\#}(X_0; \delta), \ell)$ in the last term will decrease. In (Xia & Shen, 2024), it is found that a δ that is too large may lead to systematic errors, while a δ that is too small is insufficient to quantify the uncertainty. The optimal choice of δ is problem-specific and also depends on the amount of available training data, which may require some fine-tuning. When an appropriate δ is chosen, both the first and last terms on the RHS of Eq. (2.18) can be kept small so that the expected local time-decoupled squared W_2 distance $\mathbb{E}[\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})]$ can be well controlled by $W_2^2(\mu, \hat{\mu})$. In this case, minimizing $\mathbb{E}[\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})]$ is necessary to minimize $W_2^2(\mu, \hat{\mu})$, i.e., match the distribution of θ with the distribution of $\hat{\theta}$. On the other hand, in Eq. (2.19), the rate of convergence when using the empirical probability measure is $2N^{-\frac{1}{2}}(\log(1+N)+1), \ell \leq 4$ or $N^{-\frac{2}{\ell}}, \ell > 4$ as the number of observed trajectories increases, which depends only on the dimensionality of the parameter θ instead of on the dimensionality of $X(t)$. Therefore, it is harder to reconstruct the joint distribution of higher-dimensional uncertain parameters due to a slower convergence rate when only finite training trajectories are available.

ODEs can be regarded as special jump-diffusion processes whose diffusion function and jump function are both 0. Thus, Theorem 2.2 can be applied to bound the time-decoupled local squared W_2 distance between the distributions of trajectories of the two ODEs Eqs. (1.1) and (1.2) using the squared W_2 distance between the probability distributions of θ and $\hat{\theta}$. Finally, Theorem 2.2 can also be generalized to the cases of reconstructing parameters in spatiotemporal stochastic partial differential equations (SPDEs). We provide an example of bounding the local time-decoupled squared W_2 distance between solutions to a parabolic SPDE associated with two different sets of model parameters in Appendix D.

Finally, one can also take into account time discretization errors using a time discretization scheme such as the Runge-Kutta scheme for solving ODEs and the strong Itô-Taylor approximation (Kloeden & Platen, 1992) for numerically solving SDEs. Specifically, for stiff problems, a detailed analysis of the numerical implementation is worth carrying out. Additionally, it is also possible to extend Theorem 2.2 to more complicated spatiotemporal dynamics by replacing B_t with more complicated spatiotemporal cylindrical Brownian noise (Liu & Röckner, 2015) or considering spatiotemporal integrodifferential equations (Deng et al., 2025). Such discussions require a more detailed analysis of SPDEs and are thus beyond the scope of this paper.

3. An SNN model for approximating the distribution of continuous random variables

In this section, we analyze an SNN model used in [Xia and Shen \(2024\)](#). We apply this SNN model to reconstruct the distribution of unknown parameters for general dynamical systems. Specifically, we will prove that this SNN model can approximate the probability distribution of any continuous multidimensional random variable under certain technical assumptions in the W_2 distance sense, and the training of this SNN can be done by direct minimization of the local time-decoupled squared W_2 loss function [Eq. \(2.6\)](#) analyzed in [Section 2](#). We sketch the structure of the SNN with weight uncertainty in [Fig. 1](#). All weights in the neural networks $\{w_{i,j,k}\}$ are sampled from independent normal distributions with $w_{i,j,k} \sim \mathcal{N}(a_{i,j,k}, \sigma_{i,j,k}^2)$. The biases $\{b_{i,k}\}$ for all i, k are deterministic. The means and variances of the weights $\{a_{i,j,k}\}, \{\sigma_{i,j,k}^2\}$ and the biases $\{b_{i,k}\}$ are optimized through training.

First, we prove that the SNN model in [Fig. 1](#) can approximate any continuous random variable whose probability distribution follows a parameterized multivariate normal distribution in the squared W_2 distance sense. In the following, $\mathcal{N}(\mathbf{y} - \mathbf{b}, \Sigma)$ denotes the probability density function of a d' -dimensional multivariate normal distribution with mean $\mathbf{b}$ and the covariance matrix Σ and takes the form:

$$\mathcal{N}(\mathbf{y} - \mathbf{b}, \Sigma) := \frac{1}{\sqrt{2\pi}^{d'}} \cdot \frac{1}{|\Sigma|^{\frac{1}{2}}} \cdot \exp\left(-\frac{1}{2}(\mathbf{y} - \mathbf{b})^T \Sigma (\mathbf{y} - \mathbf{b})\right). \quad (3.1)$$

Theorem 3.1. *Let $\mathbf{y} \in \mathbb{R}^{d'}$ be a continuous random variable whose probability density is $f_{\mathbf{x}}(\mathbf{y})$, $\mathbf{x} \in D \subseteq \mathbb{R}^d$ such that $f_{\mathbf{x}}(\mathbf{y}) = \mathcal{N}(\mathbf{y} - \mathbf{b}(\mathbf{x}), A(\mathbf{x})^T A(\mathbf{x}))$. D is a bounded set and $\mathbf{x}$ has a probability measure $\gamma(\cdot)$ on D . We make the following assumptions:*

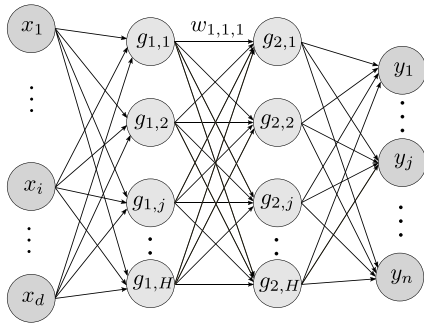
1. For any sequence of sets $\{D_i\}_{i=1}^{\infty}$, if $D_i \rightarrow D$ as $i \rightarrow \infty$, then $\lim_{i \rightarrow \infty} \gamma(D_i) = \gamma(D) = 1$. Furthermore, for any $\Delta x > 0$, we can find a set of equidistance grids $\{\mathbf{x}_i\}_{i=1}^K \subseteq D$ such that the distance between two adjacent grids is Δx , $D \subseteq \bigcup_{i=1}^K \bigotimes_{j=1}^d [x_i^j, x_i^j + \Delta x)$, and $\bigotimes_{j=1}^d [x_{i_1}^j, x_{i_1}^j + \Delta x) \cap \bigotimes_{j=1}^d [x_{i_2}^j, x_{i_2}^j + \Delta x) = \emptyset$ if $i_1 \neq i_2$.
2. $f_{\mathbf{x}}(\mathbf{y})$ is uniformly continuous in $\mathbf{x}$ such that for any $\epsilon > 0$, there exists $\Delta x > 0$ satisfying:

$$W_2^2(f_{\mathbf{x}}, f_{\tilde{\mathbf{x}}}) < \epsilon, \quad \forall \|\mathbf{x} - \tilde{\mathbf{x}}\| \leq \Delta x, \quad \mathbf{x}, \tilde{\mathbf{x}} \in D. \quad (3.2)$$

3. $Y := \sup_{\mathbf{x} \in D} \|\mathbf{b}(\mathbf{x})\|^2 + \|A(\mathbf{x})^T A(\mathbf{x})\|_F^2 < \infty$, where $\|\cdot\|_F$ is the Frobenius norm of a matrix.

Then, for any $\epsilon_0 > 0$, there exists an SNN model as described in [Fig. 1](#) which uses the ReLU activation and the linear forward propagation such that if we denote the probability density function of the output by $\hat{f}_{\mathbf{x}}$ given the input $\mathbf{x}$, the following inequality holds:

$$\int_D W_2^2(f_{\mathbf{x}}, \hat{f}_{\mathbf{x}}) \gamma(d\mathbf{x}) \leq \epsilon_0. \quad (3.3)$$



Linear: $g_{i,k} = \sum_{j=1}^H w_{i-1,j,k} g_{i-1,j} + b_{i,k}$
 ReLU activation: $g_{i,k} = \text{ReLU}(\sum_{j=1}^H w_{i-1,j,k} g_{i-1,j} + b_{i,k})$
 ResNet: $g_{i,k} = \text{ReLU}(\sum_{j=1}^H w_{i-1,j,k} g_{i-1,j} + b_{i,k}) + g_{i-1,k}$

$$w_{i,j,k} \sim \mathcal{N}(a_{i,j,k}, \sigma_{i,j,k}^2)$$

H : the number of neurons per hidden layer
 ReLU: the ReLU activation function

Fig. 1. A sketch of the structure of the neural network model with weight uncertainty used in [Xia and Shen \(2024\)](#) and in this paper. The weights $w_{i,j,k} \sim \mathcal{N}(a_{i,j,k}, \sigma_{i,j,k}^2)$ are independently sampled, i.e., w_{i_1,j_1,k_1} is independent of w_{i_2,j_2,k_2} when $(i_1, j_1, k_1) \neq (i_2, j_2, k_2)$. When using this neural network model to make predictions, for each input $\mathbf{x} = (x_1, \dots, x_d) \in D \subseteq \mathbb{R}^d$, we resample all weights $\{w_{i,j,k}\}$ again. For each neuron in the hidden layer, one of the following four forward propagation methods is considered: the linear operation, the ReLU activation, the ELU activation, and the ResNet technique.

We prove [Theorem 3.1](#) in [Appendix E](#). [Theorem 3.1](#) can be further generalized and we can prove that the probability density distribution of the output of the SNN model in [Fig. 1](#) can approximate a multivariate Gaussian mixture model (defined in [Corollary 3.1](#) below) in the squared W_2 distance sense. First, we prove the following lemma.

Lemma 3.1. *Let $\mathbf{y} \in \mathbb{R}^{d'}$ be a continuous random variable that has the following probability density function of a Gaussian mixture model:*

$$f(\mathbf{y}) = \sum_{i=1}^s p_i \mathcal{N}(\mathbf{y} - \mathbf{b}_i, A_i^T A_i), \quad p_i > 0, \quad \sum_{i=1}^s p_i = 1. \quad (3.4)$$

Suppose another continuous random variable $\hat{\mathbf{y}} \in \mathbb{R}^{d'}$ has the following probability density function:

$$\hat{f}(\hat{\mathbf{y}}) = \sum_{i=1}^s \hat{p}_i \mathcal{N}(\hat{\mathbf{y}} - \mathbf{b}_i, A_i^T A_i) + p(\hat{\mathbf{y}}), \quad \hat{p}_i > 0 \quad (3.5)$$

where $\hat{p}_i \leq p_i$, $p(\cdot) : \mathbb{R}^{d'} \rightarrow \mathbb{R}^+ \cup \{0\}$ is non-negative, and $\int_{\mathbb{R}^{d'}} \|\hat{\mathbf{y}}\|^2 p(\hat{\mathbf{y}}) d\hat{\mathbf{y}} < \infty$. Then, the following bound of the squared W_2 distance between the probability measures of $\mathbf{y}$ and $\hat{\mathbf{y}}$ holds:

$$W_2^2(f, \hat{f}) \leq 2 \left(\sum_{i=1}^s (p_i - \hat{p}_i) (\|\mathbf{b}_i\|^2 + \|A_i^T A_i\|_F^2) + \int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 p(\mathbf{y}) d\mathbf{y} \right). \quad (3.6)$$

Proof. First, if $p_i = \hat{p}_i, i = 1, \dots, s$, then $f(\mathbf{y}) = \hat{f}(\mathbf{y})$ and $W_2^2(f, \hat{f}) = 0$, indicating that [Eq. \(3.6\)](#) holds. Next, we assume that $\sum_{i=1}^s \hat{p}_i < 1$. Without loss of generality, we assume that $\mathbf{y}$ and $\hat{\mathbf{y}}$ are independent of each other. We define a special coupling probability measure of the random variable $(\mathbf{y}, \hat{\mathbf{y}}) \in \mathbb{R}^{2d'}$:

$$\pi(f, \hat{f})(\mathbf{y}, \hat{\mathbf{y}}) = \left[\sum_{i=1}^s \hat{p}_i \mathcal{N}(\mathbf{y} - \mathbf{b}_i, A_i^T A_i) \right] \delta(\mathbf{y} - \hat{\mathbf{y}}) + \frac{1}{p} \left[\sum_{i=1}^s (p_i - \hat{p}_i) \mathcal{N}(\mathbf{y} - \mathbf{b}_i, A_i^T A_i) \right] \cdot p(\hat{\mathbf{y}}), \quad (3.7)$$

where $p := \sum_{i=1}^s (p_i - \hat{p}_i) = \int_{\mathbb{R}^{d'}} p(\hat{\mathbf{y}}) d\hat{\mathbf{y}}$, and δ is the Dirac delta measure. We can check that the marginal distributions of $\pi(f, \hat{f})$ coincide with $f(\mathbf{y})$ and $\hat{f}(\hat{\mathbf{y}})$, respectively. Furthermore, we have

$$\begin{aligned} W_2^2(f, \hat{f}) &\leq \mathbb{E}_{(\mathbf{y}, \hat{\mathbf{y}}) \sim \pi(f, \hat{f})} [\|\mathbf{y} - \hat{\mathbf{y}}\|^2] \\ &\leq \int_{\mathbb{R}^{2d'}} \frac{1}{p} \left[\sum_{i=1}^s (p_i - \hat{p}_i) \mathcal{N}(\mathbf{y} - \mathbf{b}_i, A_i^T A_i) \right] \cdot p(\hat{\mathbf{y}}) \|\mathbf{y} - \hat{\mathbf{y}}\|^2 d\mathbf{y} d\hat{\mathbf{y}} \\ &\leq \int_{\mathbb{R}^{2d'}} \frac{1}{p} \left[\sum_{i=1}^s (p_i - \hat{p}_i) \mathcal{N}(\mathbf{y} - \mathbf{b}_i, A_i^T A_i) \right] \cdot p(\hat{\mathbf{y}}) \\ &\quad \cdot 2(\|\mathbf{y}\|^2 + \|\hat{\mathbf{y}}\|^2) d\mathbf{y} d\hat{\mathbf{y}} \end{aligned}$$

$$\leq 2 \sum_{i=1}^s (p_i - \hat{p}_i) (\|b_i\|^2 + \|A_i^T A_i\|_F^2) + 2 \int_{\mathbb{R}^{d'}} \|\hat{y}\|^2 p(\hat{y}) d\hat{y}, \quad (3.8)$$

which proves the inequality (3.6). $\square$

Next, we show that the SNN model in Fig. 1 can approximate a multivariate Gaussian mixture model in the squared W_2 distance sense.

Corollary 3.1. *Let $y \in \mathbb{R}^{d'}$ be a continuous random variable with a probability density function f_x , where $x \in D \subseteq \mathbb{R}^d$ is continuous and has a probability density $\gamma(\cdot)$. At each x , f_x is the probability density function of a Gaussian mixture model:*

$$f_x(y) = \sum_{r=1}^s p_r(x) \mathcal{N}(y - b_r(x), A_r^T(x) A_r(x)), \quad \sum_{r=1}^s p_r(x) = 1, \quad p_r(x) > 0. \quad (3.9)$$

We make the following three additional assumptions:

1. For any sequence of sets $\{D_i\}_{i=1}^\infty$, if $D_i \rightarrow D$ as $i \rightarrow \infty$, then $\lim_{i \rightarrow \infty} \gamma(D_i) = \gamma(D) = 1$. Additionally, D is a bounded set in $\mathbb{R}^d$ and for any $\Delta x > 0$, we can find a set of equidistance grids $\{x_i\}_{i=1}^K \subseteq D$ such that $D \subseteq \bigcup_{i=1}^K \bigotimes_{j=1}^d [x_{i_j}^j, x_{i_j}^j + \Delta x]$ and $\bigotimes_{j=1}^d [x_{i_1}^j, x_{i_1}^j + \Delta x] \cap \bigotimes_{j=1}^d [x_{i_2}^j, x_{i_2}^j + \Delta x] = \emptyset$ if $i_1 \neq i_2$.
2. $f_x(y)$ is uniformly continuous in x such that for any $\epsilon > 0$, there exists a $\delta > 0$:

$$W_2^2(f_x(y), f_{\tilde{x}}(y)) < \epsilon, \quad \forall \|x - \tilde{x}\| \leq \delta, \quad \forall x, \tilde{x} \in D. \quad (3.10)$$

3. The quantity

$$\max_{1 \leq r \leq s} \|A_r^T(x) A_r^T(x)\|_F^2 + \|b_r(x)\|^2 \quad (3.11)$$

is uniformly bounded for all $x \in D$.

Then, for any positive number $c_0 > 0$, there exists an SNN with weight uncertainty described in Fig. 1 such that:

$$\int_D W_2^2(f_x, \hat{f}_x) \gamma(dx) \leq c_0. \quad (3.12)$$

Here, $\hat{f}_x$ is the distribution of the output of the SNN when the input is x .

We prove Corollary 3.1 in Appendix F. Finally, we prove that for each continuous random variable $y \in \mathbb{R}^{d'}$ with a probability distribution function $f(y)$, under certain conditions, we can find a random variable $\hat{y} \in \mathbb{R}^{d'}$ obeying a Gaussian mixture distribution whose probability distribution function is denoted by $\hat{f}$ such that $\hat{f}$ can approximate f in the W_2 sense. We have the following result.

Theorem 3.2. *Suppose $y = (y_1, \dots, y_{d'}) \in \mathbb{R}^{d'}$ is a continuous random variable with a smooth probability density function $f(y) \in L^2(\mathbb{R}^{d'}) \cap L^\infty(\mathbb{R}^{d'})$. Furthermore, we assume the following conditions hold: 1. $f(y)$ is uniformly continuous in $\mathbb{R}^n$ 2.*

$$\|f\|_{\text{mix}} := \sum_{|n|_0 \leq d'} \left\| \partial_n^{n|_0} f \right\|_{L^2} < \infty, \quad \left| \sqrt{f} \right|_{\text{mix}} < \infty \quad (3.13)$$

where $|n|_0$ is the number of non-zero components in n , $n = (n_1, \dots, n_j)$ satisfying $1 \leq n_1 < \dots < n_j \leq d'$, and $\partial_n f := \partial_{y_{n_1}} \dots \partial_{y_{n_j}} f$.

3. $\|f y_i^2\|_{\text{mix}} < \infty$ and $\|f y_i^2 y_j^2\|_{\text{mix}} < \infty$.

Then, for every $\sigma > 0$, $n_0(\sigma) > 0$, there exists a probability density function of a Gaussian mixture model:

$$\tilde{f}_{\sigma^2, n_0(\sigma)}(y) := \sum_{i=1}^{(n_0(\sigma)+1)^{d'}} p_i \mathcal{N}(y - y_i; \sigma^2 I_{d' \times d'}), \quad y_i \in \mathbb{R}^{d'} \quad (3.14)$$

such that as $\sigma \rightarrow 0^+$ and $n(\sigma) \rightarrow \infty$:

$$\tilde{f}_{\sigma^2, n_0(\sigma)}(y) \rightarrow f(y) \quad (3.15)$$

uniformly in $\mathbb{R}^{d'}$. Furthermore, $\forall \epsilon > 0$, there exist $\sigma > 0$, $n_0(\sigma)$ and $\tilde{f}_{\sigma^2, n_0(\sigma)}(y)$ in Eq. (3.14) such that

$$W_2^2(f, \tilde{f}_{\sigma^2, n_0(\sigma)}) \leq 24\epsilon. \quad (3.16)$$

In Eq. (3.14), $I_{d' \times d'}$ is a $d' \times d'$ identity matrix.

We prove Theorem 3.2 in Appendix G. For any continuous random variable $y \in \mathbb{R}^{d'}$ with a probability density function f satisfying the assumptions in Theorem 3.2, there exists a random variable $\hat{y}$ whose probability density function is the probability density function of a multivariate Gaussian mixture model denoted by $\tilde{f}_{\sigma^2, n_0}$ such that $W_2^2(f, \tilde{f}_{\sigma^2, n_0}) \leq \epsilon$. As a special case of the proof of Corollary 3.1 in Appendix F, there exists an SNN (Fig. 1) which takes the input $n_i^3 = x \equiv 1$, i.e. this SNN consists of the fourth, fifth, sixth, and seventh layers of the SNN model we constructed in Appendix F, such that $W_2^2(\tilde{f}_{\sigma^2, n_0}, \hat{f}_{x \equiv 1}) \leq \epsilon$, where $\hat{f}_{x \equiv 1}$ is the probability density function of the output of the SNN. Thus, we have $W_2^2(f, \hat{f}_{x \equiv 1}) \leq 4\epsilon$, i.e., the SNN model in Fig. 1 can approximate the distribution of any continuous random variable in the squared W_2 sense under the assumptions specified in Theorem 3.2. From Theorem 2.2, minimizing the local time-decoupled squared W_2 loss function Eq. (2.6) is necessary to minimize $W_2^2(\mu, \hat{\mu})$, i.e., to match the distribution of model parameters θ by the distribution of $\hat{\theta}$ in dynamical systems like Eqs. (2.9) and (2.10). Thus, we can train the SNN model in Fig. 1 by direct minimization of the local time-decoupled squared W_2 loss function Eq. (2.6).

Remark: Consider the more general model $y = f(x, \theta)$, $x \in \mathbb{R}^d$, $y \in \mathbb{R}^{d'}$ as in Xia and Shen (2024), where θ is the latent random model parameters (e.g. measurement noise). y is a continuous random variable whose distribution is determined by the observed continuous variable $x \in D \subseteq \mathbb{R}^d$. D is a bounded set in $\mathbb{R}^d$, and for any sequence of sets $\{D_i\}_{i=1}^\infty$, if $D_i \rightarrow D$ as $i \rightarrow \infty$, then $\lim_{i \rightarrow \infty} \gamma(D_i) = \gamma(D) = 1$. Additionally, we assume that for any $\Delta x > 0$, we can find a set of equidistance grids $\{x_i\}_{i=1}^K \subseteq D$ such that $D \subseteq \bigcup_{i=1}^K \bigotimes_{j=1}^d [x_{i_j}^j, x_{i_j}^j + \Delta x]$ and $\bigotimes_{j=1}^d [x_{i_1}^j, x_{i_1}^j + \Delta x] \cap \bigotimes_{j=1}^d [x_{i_2}^j, x_{i_2}^j + \Delta x] = \emptyset$ if $i_1 \neq i_2$. ω are continuous latent parameters in the model sampled from an unknown distribution. We denote the distribution of y given x by f_x . Under certain assumptions, for any $\epsilon > 0$, there exists an SNN in Fig. 1 such that:

$$\int_D W_2^2(f_x, \hat{f}_x) \gamma(dx) \leq \epsilon. \quad (3.17)$$

where $\hat{f}_x$ is the probability density function of the output of the SNN given the input x . We give a brief discussion on this “universal approximation” property of the SNN model to approximate a family of continuous random variables $y = f(x, \theta)$ for all $x \in D \subseteq \mathbb{R}^d$ in Appendix H. Our SNN can approximate a family of probability density functions $f(x, \theta)$, $x \in D$ simultaneously, while in Lu and Lu (2020) only a single probability density function $f(\theta)$ depending on θ is to be approximated. Furthermore, our SNN utilizes only seven hidden layers and the number of neurons in each layer scales linearly with the dimensionality of either the input or the output variable. This implies that even with a shallow neural network, we might be able to reconstruct a family of uncertainty models $y = f(x, \theta)$ characterized by different $x \in D$.

4. Numerical results

In this section, we conduct numerical experiments to test our proposed local squared W_2 method, which involves training the SNN model in Fig. 1 by minimizing the loss function Eq. (J.1), a scaled numerical approximation to the local time-decoupled squared W_2 distance loss function Eq. (2.6). The W_2 distance between two empirical probability measures is then numerically evaluated using the PoT package of Python in Flamary et al. (2021). A pseudocode of our method is given in Algorithm 1, and a schematic diagram on the training of the SNN using our proposed local time-decoupled squared W_2 loss function Eq. (2.4) is presented in Fig. 2.

Specifically, when applying the SNN model in Fig. 1 for the reconstruction of the distribution for uncertain model parameters, we assume that uncertain model parameters are sampled from the same underlying distribution across all samples and thus always input a scalar 1 as the input into the SNN. Default hyperparameters and training settings are given in Appendix I. We use the ELU activation function to replace the ReLU activation function in Examples 4.1–4.3 to tackle the issue

Algorithm 1 The pseudocode of our local time-decoupled squared W_2 method for training the SNN (the local time-decoupled squared W_2 loss Eq. (2.4) can be replaced with other loss functions).

Given N observed time-series data $\{X_i(t_j), t_j = j\Delta t, j = 1, \dots, N_T\}_{i=1}^N$, the underlying dynamical system with unknown model parameters θ (such as the jump-diffusion process Eq. (2.9)), and the maximal epochs $i_{\max}$.

Initialize the SNN model in Fig. 1.

Input a scalar 1 into the SNN and evaluate the SNN N times independently (each time, the weights in the SNN are resampled), which outputs N sets of approximate parameters denoted by $\{\hat{\theta}\}_{i=1}^N$.

Generate N trajectories $\{\hat{X}\}_{i=1}^N$ from the approximate dynamical system (such as Eq. (2.10)) with the approximate parameters $\{\hat{\theta}\}_{i=1}^N$.

while $i < i_{\max}$ **do**

 Perform gradient descent to minimize the loss function $\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})$ (Eq. (2.4)) and update the parameters (biases & means and variances of weights) in the SNN model.

 Input a scalar 1 into the updated SNN and evaluate the updated SNN N times independently, which outputs N sets of approximate parameters denoted by $\{\hat{\theta}\}_{i=1}^N$.

 Generate N trajectories from the approximate dynamical system with the approximate parameters $\{\hat{\theta}\}_{i=1}^N$.

end while

return the trained SNN model

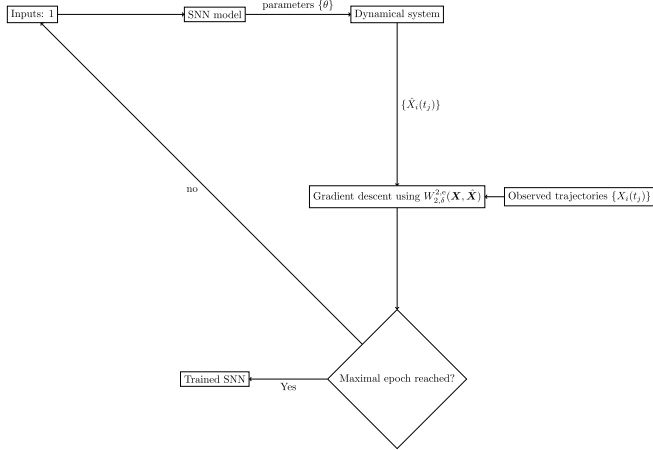


Fig. 2. A schematic diagram on the training of the SNN using our proposed local time-decoupled squared W_2 loss function.

of vanishing gradients and improve the representational power of the SNN. In the following, errors in the distribution of reconstructed model parameters denote the scaled squared W_2 distance:

$$\text{error} := \frac{W_2^2(\mu_\theta, \hat{\mu}_\theta)}{\|\theta\|^2}, \quad (4.1)$$

where θ and $\hat{\theta}$ are unknown ground truth model parameters and reconstructed model parameters, e.g. in Eqs. (1.1) and (1.2), and μ_θ and $\hat{\mu}_\theta$ are the distributions of θ and $\hat{\theta}$, respectively. For numerically solving ODEs in all examples, we use the odeint function with default settings in the torchdiffeq package (Chen et al., 2018). In Example 4.4, we use the package developed in Xia et al. (2024a) for numerically solving a jump-diffusion process using the Euler scheme (Platen & Bruti-Liberati, 2010), which allows for back-propagation and gradient descent for hyperparameter optimization in the neural networks. The Itô integral is adopted for evaluating the stochastic integrals. The numerical experiment in Example 4.1 is conducted using Python 3.11 on a desktop with a 32-core Intel® i9-13900KF CPU (when comparing runtimes and RAM usage, we train each model on just one core). Numerical experiments in

Examples 4.2–4.4 are carried out using Python 3.11 on NYU HPC with a GPU (NYU HPC, 2025).

Example 4.1. First, we consider reconstructing the following 2D ODE characterizing the Lotka–Volterra predator–prey dynamics with one uncertain predation rate parameter:

$$\begin{aligned} \frac{dx}{dt} &= 2x - cxy, \quad \frac{dy}{dt} = \frac{1}{4}cxy - 2y, \\ (x(0), y(0)) &= (\xi_1, \xi_2), \quad c \sim \mathcal{U}(2, 4), \quad \xi_1, \xi_2 \sim \mathcal{U}(1, 2), \quad t \in [0, 8]. \end{aligned} \quad (4.2)$$

In Eq. (4.2), c is the uncertain predation rate parameter, and ξ_1, ξ_2 are two independent random variables. We train the SNN model in Fig. 1 by minimizing the loss function Eq. (J.1) (the neighborhood size $\delta = 0.4$) to reconstruct the distribution of the parameter c . For comparison, we minimize other loss functions (definitions given in Appendix J) commonly used in statistical inference tasks to train the SNN model in Fig. 1. We also compare our method with a Bayesian neural network (BNN) method used in Bayesian Neural Network Derivation (2020), Mullachery et al. (2018) which minimizes the Kullback–Leibler divergence as well as a WGAN method presented in Arjovsky et al. (2017), Gulrajani et al. (2017). For implementing the WGAN method, the generator is the same as the SNN in Fig. 1 when utilizing other loss functions, while the discriminator is a feedforward neural network with one hidden layer equipped with 32 neurons and the ReLU activation function. Within each epoch for training the generator, 256 randomly selected samples are provided to train the discriminator, repeated 4 times. After training, the generator (SNN) is used to generate the distribution of the reconstructed model parameter.

From Fig. 3(a), (b), the distribution of ground truth trajectories of the prey and predator population can be matched well by the distribution of trajectories generated by numerically solving Eq. (4.2) with $\hat{c}$ sampled from the reconstructed distribution. Furthermore, the distribution of the reconstructed $\hat{c}$ generated by the SNN model trained using our loss function Eq. (J.2) aligns well with the distribution of the ground truth predation rate c (shown in Fig. 3(c)). Using the previous time-decoupled squared W_2 loss function in Xia et al. (2024b), the MMD loss function, or the Mean² + Var loss function all lead to an inaccurate calculated value of the mean in the reconstructed $\hat{c}$. The WGAN method and the BNN method both yield qualitatively incorrect results. We find that training the discriminator and generator in the WGAN approach requires fine-tuning of hyperparameters by setting the learning rate to be a relatively small value of 10^{-5} , otherwise the training of the generator and discriminator will terminate prematurely due to blowup, as reported in the literature (Hoffer et al., 2017). Nonetheless, the WGAN could not yield a qualitatively correct mean of the reconstructed $\hat{c}$. On the other hand, the BNN method also does not perform well, and one possible reason could be that randomly initialized means and variances of weights as well as the biases in the neural network do not provide a good prior distribution for the BNN. Using the MSE as the loss function to train the SNN yields an almost degenerate distribution of $\hat{c}$ and the calculated value of the variance in $\hat{c}$ is inaccurate, indicating that it is not suitable to train the SNN model for reconstructing the distribution of unknown parameters. Using our proposed local time-decoupled squared W_2 loss function yields more accurate values of the mean and variance of the reconstructed $\hat{c}$ compared to other methods (shown in Fig. 3(d)). Finally, from Table 1, the runtime and RAM usage when using our proposed local time-decoupled squared W_2 loss function is generally comparable to those when using other benchmark loss functions, and the runtime and RAM usage of the WGAN method and BNN method are significantly larger than training the SNN in Fig. 1 by directly minimizing any one of the loss functions we use. Thus, using our proposed local time-decoupled squared W_2 method is more efficient than using other benchmark statistical loss functions for training the SNN in Fig. 1 to reconstruct the distribution of the unknown predation rate in Eq. (4.2), and it also outperforms the WGAN method and the BNN method.

Next, we apply our local time-decoupled squared W_2 method for the reconstruction of model parameters in a spatiotemporal PDE.

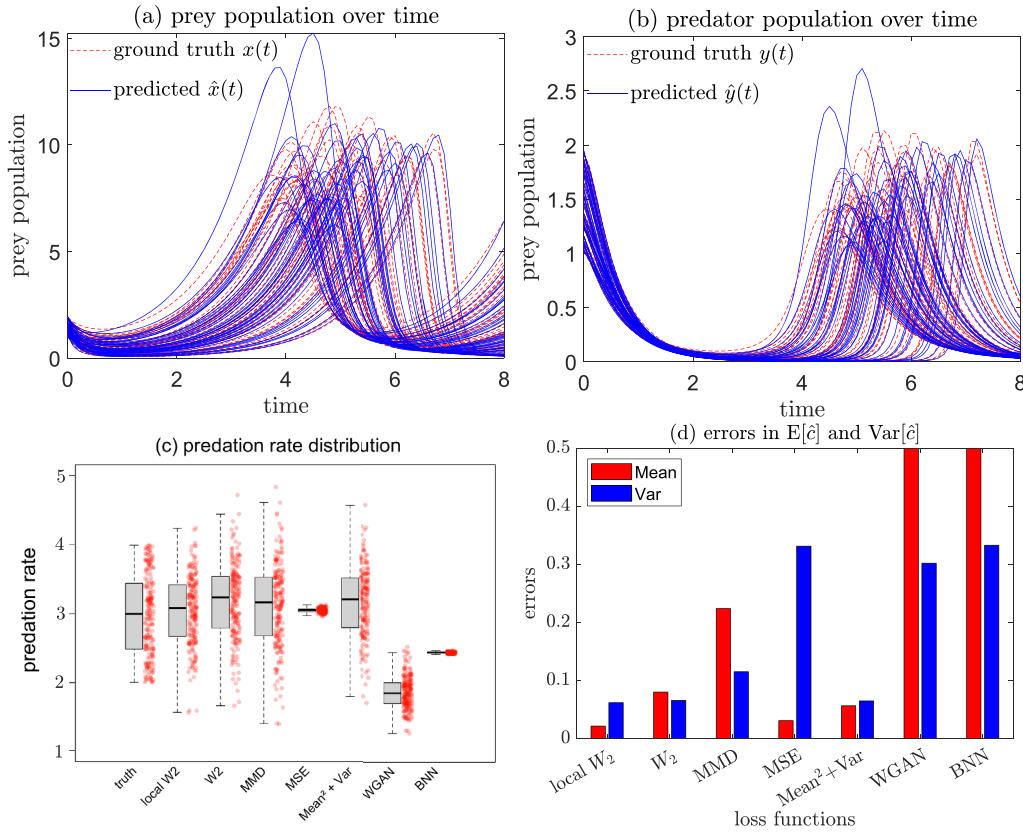


Fig. 3. (a) Ground truth (red dashed lines) prey population dynamics versus predicted prey population dynamics (blue solid lines) obtained with reconstructed predation rate $\hat{c}$. (b) Ground truth (red) predator population dynamics versus predicted predator population dynamics (blue) obtained with reconstructed predation rate $\hat{c}$. In (a) and (b), for clarity, we plot the first 50 groups of prey and predatory trajectories. Since the predation rate c in Eq. (4.2) is sampled independently for each realization of the model Eq. (4.2), the ground truth trajectories also form a distribution. (c) Ground truth $c \sim \mathcal{U}(2, 4)$ versus the distribution of the approximate $\hat{c}$ when minimizing different loss functions or using the BNN or WGAN method. The black horizontal line and the box indicate the median and the interquartile range of the ground truth or predicted predation rate. (d) Errors in the predicted mean $|\mathbb{E}[\hat{c}] - \mathbb{E}[c]|$ and predicted variance $|\text{Var}[\hat{c}] - \text{Var}[c]|$ when minimizing different loss functions. The errors are their averaged values over 5 independent experiments. In (c) and (d), “local W_2 ” refers to our scaled local time-decoupled squared W_2 loss function Eq. (J.1) while “ W_2 ” refers to previous time-decoupled squared W_2 loss function in Xia et al. (2024b).

Table 1

Computational time and memory usage when utilizing different loss functions. Mean and the standard deviation in the runtime (unit: hour) and RAM (unit: Mb) usage over five repeated experiments are recorded.

	Ours (Eq. (J.1))	Time-decoupled W_2^2	MMD	MSE	Mean ² + Var	WGAN	BNN
Time	1.53 ± 0.35	1.57 ± 0.36	1.91 ± 0.47	1.37 ± 0.46	1.73 ± 0.57	7.88 ± 1.17	21.31 ± 5.40
RAM	729.8 ± 56.3	648.4 ± 108.7	834.3 ± 19.3	616.5 ± 110.3	614.3 ± 108.2	1341.4 ± 293.5	3406.2 ± 300.0

Example 4.2. We consider reconstructing the distributions of parameters in the following parabolic PDE:

$$\begin{aligned} \partial_t u(x, t; c_1, c_2) &= \frac{c_1}{c_2^2} \partial_{xx} u(x, t; c_1, c_2) + \frac{c_1}{c_1 t + 1} u(x, t; c_1, c_2), \quad (x, t) \in \mathbb{R} \times [0, 2], \\ u(x, 0) &= (1 + \xi) \frac{x}{\sqrt{4}} \cdot \exp\left(-\frac{x^2}{2}\right), \quad c_1 \sim \mathcal{N}(0.5, \sigma_1^2), \\ c_2 &= \tilde{\xi} + \beta(0.5 - c_1), \quad \xi \sim \mathcal{N}(0, \sigma_3^2), \quad \tilde{\xi} \sim \mathcal{N}(1.5, \sigma_2^2). \end{aligned} \quad (4.3)$$

We use another parabolic PDE model to approximate Eq. (4.3):

$$\begin{aligned} \partial_t \hat{u}(x, t; \hat{c}_1, \hat{c}_2) &= \frac{\hat{c}_1}{\hat{c}_2^2} \partial_{xx} \hat{u}(x, t; \hat{c}_1, \hat{c}_2) + \frac{\hat{c}_1}{\hat{c}_1 t + 1} \hat{u}(x, t; \hat{c}_1, \hat{c}_2), \quad (x, t) \in \mathbb{R} \times [0, 2], \\ \hat{u}(x, 0) &= u(x, 0). \end{aligned} \quad (4.4)$$

We wish to reconstruct the distribution of (c_1, c_2) in Eq. (4.3) using the distribution of $(\hat{c}_1, \hat{c}_2)$ in the approximate Eq. (4.4). We use a pseudo-spectral method with a spectral expansion in space to solve Eqs. (4.3) and (4.4)

numerically:

$$\begin{aligned} u(x, t; \theta) &\approx u_{n-1}(x, t; \theta) = \sum_{i=0}^{n-1} u_i(t; \theta) \hat{H}_i(x), \quad \theta := (c_1, c_2) \\ \hat{u}(x, t; \hat{\theta}) &\approx \hat{u}_{n-1}(x, t; \hat{\theta}) = \sum_{i=0}^{n-1} \hat{u}_i(t; \hat{\theta}) \hat{H}_i(x), \quad \hat{\theta} := (\hat{c}_1, \hat{c}_2), \end{aligned} \quad (4.5)$$

where $\hat{H}_i$ is the generalized Hermite function described in Shen and Wang (2010).

We carry out the following sensitivity tests:

1. Vary (σ_1, σ_2) which determines the variance in (c_1, c_2) to investigate how variances in model parameters would affect the reconstruction accuracy of the joint distribution of (c_1, c_2) . Other parameters are set as $\beta = 1$, $\sigma_3 = 0.2$, $n = 12$ and $\delta = 0.1$.
2. Change the value of σ_3 characterizing the uncertainty in the initial condition and the size of the neighborhood δ in our loss function Eq. (J.1) to investigate how uncertainty in the initial condition and the hyperparam-

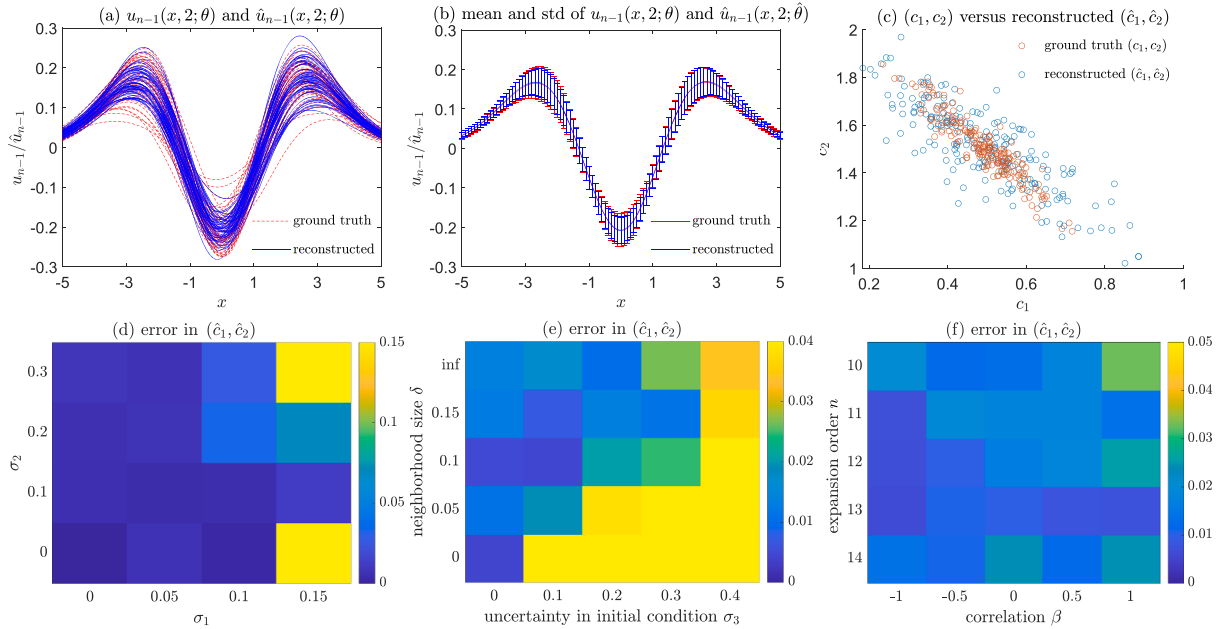


Fig. 4. (a) Ground truth $u_{n-1}(x, 2; \theta)$ (red dashed lines) versus reconstructed $\hat{u}_{n-1}(x, 2; \hat{\theta})$ (blue solid lines). For clarity, we only plot 50 ground truth $u_{n-1}(x, 2; \theta)$ numerical solutions versus 50 approximate numerical solutions $\hat{u}_{n-1}(x, 2; \hat{\theta})$ in Eq. (4.5). (b) Mean and standard deviations of the ground truth $u_{N-1}(x, 2)$ versus reconstructed $\hat{u}_{n-1}(x, 2)$. (c) The ground truth (c_1, c_2) versus reconstructed $(\hat{c}_1, \hat{c}_2)$ when $\beta = 1, \sigma_1 = 0.15, \sigma_2 = 0.1, n = 12$. In (a), (b) and (c), the parameters are $n = 12$ and $\beta = 1, \sigma_1 = 0.15, \sigma_2 = 0.1, \sigma_3 = 0.2, N = 12$ and $\delta = 0.1$ in the loss function Eq. (J.1). (d) Errors in $(\hat{c}_1, \hat{c}_2)$ w.r.t. different variances σ_1, σ_2 for (c_1, c_2) in Eq. (4.3) (Case 1 on Page 23). (e) Errors in $(\hat{c}_1, \hat{c}_2)$ w.r.t. different values of the variance σ_3 in the initial condition ϵ and different δ . $\delta = \text{inf}$ indicates that we set $\delta = \infty$, which corresponds to the time-decoupled squared W_2 loss function in Xia et al. (2024a) (Case 2 on Page 23). (f) Errors in $(\hat{c}_1, \hat{c}_2)$ w.r.t. different values N and c_0 (Case 3 on Page 23).

ter δ affect the reconstruction of the distribution (c_1, c_2) . Other parameters are $\beta = 1, \sigma_1 = 0.1, \sigma_2 = 0.2, n = 12$.

3. Vary β which determines the correlation between c_1 and c_2 as well as the expansion order N in the spectral approximation in Eq. (4.5) to explore how the correlation between c_1 and c_2 and the dimensionality of the discretized ODE affect the reconstruction accuracy of (c_1, c_2) . Other parameters are $\sigma_1 = 0.1, \sigma_2 = 0.2, \sigma_3 = 0.2, \delta = 0.1$.

We plot the reconstructed numerical solution $\hat{u}_{n-1}(x, t; \hat{\theta})$ versus the ground truth numerical solution $u_{n-1}(x, t; \theta)$ in Eq. (4.5) as the numerical solution to Eq. (4.3) in Fig. 4(a), (b). The distribution of ground truth numerical solutions can be matched well by the distribution of the numerical solutions generated with the reconstructed $(\hat{c}_1, \hat{c}_2)$, shown in Fig. 4(c). From Fig. 4(d), the larger the variance in c_1 , the larger the errors in the predicted distribution of $(\hat{c}_1, \hat{c}_2)$. When the variance c_1 is too large, we might have close-to-zero or even negative c_1 in Eq. (4.3), which makes numerically solving Eq. (4.3) ill-posed and yields a poor reconstruction of the joint distribution of (c_1, c_2) . Thus, it is reasonable to reconstruct the distribution of model parameters for well-posed dynamical systems instead of ill-posed dynamical systems. From Fig. 4(e), when the variance σ_3^2 in the initial condition is small and the initial conditions are more densely distributed, a smaller δ leads to a more accurate reconstruction of (c_1, c_2) . Thus, it is necessary to consider uncertainties in the initial condition of an ODE or PDE for the accurate reconstruction of unknown model parameters and find the size of the neighborhood δ that is compatible with uncertainty in the initial state. Finally, from Fig. 4(f), the error in the reconstructed distribution of $(\hat{c}_1, \hat{c}_2)$ is independent of the dimensionality n of the discretized ODE system. Furthermore, the accuracy of the reconstructed $(\hat{c}_1, \hat{c}_2)$ is insensitive to the correlation between the two uncertain model parameters (c_1, c_2) , indicating the robustness of our proposed local temporally decoupled squared W_2 method for reconstructing the distribution of (c_1, c_2) in Eq. (4.3).

As an application of our proposed method for parameter inference problems in biophysics, we reconstruct reaction rates in an 8D ODE sys-

tem characterizing drug dynamics in an ocular model studied in Crawshaw et al. (2025).

Example 4.3. We consider an ODE system that arises from modeling studies of monoclonal antibodies (MABs) injected into the vitreous gel of the eye in the treatment of wet age-related macular degeneration. The monoclonal antibodies bind to the vascular endothelial growth factor (VEGF) with this binding inhibiting the latter's stimulation of pathological capillary growth through the retina (Chappelow & Kaiser, 2008; Mitchell et al., 2018). However, injected MABs are cleared from the eye, by passing into the aqueous compartment at the front of the eye. This in turn necessitates multiple injections of MABs into the eye. Hence, analyzing how long MABs are retained within the eye and the prospect of reducing injection frequency has motivated numerous modeling studies (e.g. (Caruso et al., 2019; Crawshaw et al., 2025; Hutton-Smith et al., 2018, 2016; Lamirande et al., 2024)). However, a common theme within these studies is the need to estimate the parameters (Crawshaw et al., 2025; Hutton-Smith et al., 2016; Mitchell et al., 2018).

Hence, we consider a simple well-mixed model in the literature for the MAB ranibizumab (Crawshaw et al., 2025; Hutton-Smith et al., 2016), with its concentration in the vitreous of the eye denoted by r_{vit} , and the vitreous VEGF concentration denoted by v_{vit} . A complex of one MAB bound to one VEGF in the vitreous has concentration c_{vit} and, noting that VEGF has two binding sites, a complex of two MABs and a VEGF is also considered, with concentration h_{vit} . The formation and disassociation of these species is represented by the law of mass action, with the reaction rates labeled by “ k ”, and all species can pass into the aqueous at the front of the eye, with concentrations in this compartment then distinguished by the subscript “aq” ($r_{aq}, v_{aq}, c_{aq}, h_{aq}$). Once in this latter compartment, all species are taken to pass into the bloodstream via Schlemm’s canal, with a clearance rate CL . In addition, there is a flux, V_{in} , of vitreal VEGF to reflect the rate of elevated levels of VEGF leaking into the vitreous in pathology. This generates the set of eight ODEs:

Vitreous

$$\begin{aligned}
\frac{dv_{vit}}{dt} &= (k_{off}c_{vit} - 2k_{on}v_{vit}r_{vit}) - k_v^{el}v_{vit} + \frac{V_{in}}{V_{vit}}, \\
\frac{dr_{vit}}{dt} &= (k_{off}c_{vit} - 2k_{on}v_{vit}r_{vit}) + (2k_{off}h_{vit} - k_{on}r_{vit}c_{vit}) - k_r^{el}r_{vit}, \\
\frac{dc_{vit}}{dt} &= -(k_{off}c_{vit} - 2k_{on}v_{vit}r_{vit}) + (2k_{off}h_{vit} - k_{on}r_{vit}c_{vit}) - k_c^{el}c_{vit}, \\
\frac{dh_{vit}}{dt} &= -(2k_{off}h_{vit} - k_{on}r_{vit}c_{vit}) - k_h^{el}h_{vit}.
\end{aligned} \quad (4.6)$$

Aqueous

$$\begin{aligned}
\frac{dv_{aq}}{dt} &= (k_{off}c_{aq} - 2k_{on}v_{aq}r_{aq}) + \frac{V_{vit}}{V_{aq}}k_v^{el}v_{vit} - \frac{CL}{V_{aq}}v_{aq}, \\
\frac{dr_{aq}}{dt} &= (k_{off}c_{aq} - 2k_{on}v_{aq}r_{aq}) + (2k_{off}h_{aq} - k_{on}r_{aq}c_{aq}) + \frac{V_{vit}}{V_{aq}}k_r^{el}r_{vit} \\
&\quad - \frac{CL}{V_{aq}}r_{aq}, \\
\frac{dc_{aq}}{dt} &= -(k_{off}c_{aq} - 2k_{on}v_{aq}r_{aq}) + (2k_{off}h_{aq} - k_{on}r_{aq}c_{aq}) + \frac{V_{vit}}{V_{aq}}k_c^{el}c_{vit} \\
&\quad - \frac{CL}{V_{aq}}c_{aq}, \\
\frac{dh_{aq}}{dt} &= -(2k_{off}h_{aq} - k_{on}r_{aq}c_{aq}) + \frac{V_{vit}}{V_{aq}}k_h^{el}h_{vit} - \frac{CL}{V_{aq}}h_{aq}, \quad t \in [0, 2] \text{ (unit: day)}.
\end{aligned} \quad (4.7)$$

From (Crawshaw et al., 2025), we set the vitreous and aqueous humor volumes $V_{vit} = 2.05\text{mL}$ and $V_{aq} = 0.105\text{mL}$, respectively, and $V_{in} = 5.408\text{pmol} \cdot \text{day}^{-1}$. We aim to reconstruct the distribution of the seven reaction rate parameters: k_{off} , k_{on} , k_v^{el} , k_r^{el} , k_c^{el} , k_h^{el} and CL in Eqs. (4.6) and (4.7), which are subject to uncertainties in the drug properties (Crawshaw et al., 2025; Hutton-Smith et al., 2016; Mitchell et al., 2018).

We generate a synthetic data set of model parameters by sampling the seven kinetic parameters: $\mathbf{k} := (k_{off}, k_{on}, k_v^{el}, k_r^{el}, k_c^{el}, k_h^{el}, CL)$ from the following model:

$$\mathbf{k} = \mathbf{k}_0 + \mathbf{c} \mathbf{k}_0 * \mathbf{A} \tilde{\mathbf{k}}, \quad (4.8)$$

where

$$\begin{aligned}
\mathbf{k}_0 &= (1.669\text{day}^{-1}, 0.00114\text{pM}^{-1} \cdot \text{day}^{-1}, 0.575\text{day}^{-1}, 0.293\text{day}^{-1}, \\
&0.259\text{day}^{-1}, 0.176\text{day}^{-1}, 2.505\text{mL} \cdot \text{day}^{-1})
\end{aligned} \quad (4.9)$$

is the vector of mean values of those kinetic parameters used in Crawshaw et al. (2025). In Eq. (4.8), $*$ is the Hadamard componentwise product. $\mathbf{A} \in \mathbb{R}^{7 \times 7}$ is a randomly generated matrix whose components are sampled independently from the distribution $\mathcal{U}(-\frac{1}{2}, \frac{1}{2})$. In Eq. (4.8), $\tilde{\mathbf{k}} := (k_1, k_2, \dots, k_7)$ is independently generated for each trajectory in the training dataset. Omitting the units for simplicity, we sample $k_1, k_2 \sim \mathcal{U}(0, 1)$, $k_3, k_4 \sim \mathcal{N}(0, 0.5^2)$, $k_5 \sim \text{Exp}(2)$, $k_6 \sim \text{B}(2, 5)$ (the Beta distribution with shape parameters $\alpha = 2, \beta = 5$), and $k_7 \sim \Gamma(2, 2)$ (the Gamma distribution with a shape parameter $\alpha = 2$ and scale parameter $\lambda = 2$). The initial condition of each trajectory is independently sampled from $\mathcal{N}(\mathbf{I}_7, 0.05^2 \mathbf{I}_{7 \times 7})$, where $\mathbf{I}_7 \in \mathbb{R}^7$ refers to a constant vector whose components are 1.

The distribution of trajectories of the four quantities $v_{vit}(t)$, $r_{vit}(t)$, $c_{vit}(t)$ and $h_{vit}(t)$ obtained with the parameter vector $\mathbf{k}$ sampled from the ground truth distribution Eq. (4.8) can be matched well by the distribution of trajectories obtained by using the parameter vector $\mathbf{k}$ sampled from the reconstructed distribution generated by the trained SNN in Fig. 5(a)–(d). Additionally, we plot the empirical joint distribution of any two parameters in $\mathbf{k}$ versus the empirical joint reconstructed distribution of the corresponding two variables in Fig. 5(e), and most pairwise joint distributions of any two components in $\mathbf{k}$ can be matched well by their reconstructed counterparts. However, the reconstruction of k_{on} is not accurate. A possible reason could be that the magnitude of k_{on} ($O(10^{-3})$) is much smaller than that of other parameters ($O(1)$), making it harder to reconstruct the distribution of k_{on} . Another possibility is a lack of practical identifiability for these parameters. Since only

Table 2

Errors in the reconstructed distribution of kinetic parameters in the ocular pharmacokinetic model. Here, ELU refers to using the ELU activation function for forward propagation and ResNet refers to using the ResNet technique (described in Fig. 1). When using the W_2 , MMD, MSE, and the Mean² + Var loss functions, the SNN has 3 hidden layers with 10 neurons in each layer.

Width	# of layers	Forward propagation	Initialization for weights & biases	Error
10	1	ELU	$\mathcal{N}(0, 0.03^2)$	7.79×10^{-3}
10	2	ELU	$\mathcal{N}(0, 0.03^2)$	1.39×10^{-3}
10	3	ELU	$\mathcal{N}(0, 0.03^2)$	1.10×10^{-3}
10	4	ELU	$\mathcal{N}(0, 0.03^2)$	8.93×10^{-4}
5	3	ELU	$\mathcal{N}(0, 0.03^2)$	8.85×10^{-4}
15	3	ELU	$\mathcal{N}(0, 0.03^2)$	9.20×10^{-4}
20	3	ELU	$\mathcal{N}(0, 0.03^2)$	0.156
10	2	ResNet	$\mathcal{N}(0, 0.03^2)$	8.25×10^{-4}
10	3	ResNet	$\mathcal{N}(0, 0.03^2)$	8.67×10^{-4}
10	4	ResNet	$\mathcal{N}(0, 0.03^2)$	8.71×10^{-4}
10	3	ELU	$\mathcal{N}(0, 0)$	0.609
10	3	ELU	$\mathcal{N}(0, 0.01^2)$	7.98×10^{-4}
10	3	ELU	$\mathcal{N}(0, 0.02^2)$	8.42×10^{-4}
W_2	3	ELU	$\mathcal{N}(0, 0.03^2)$	1.24×10^{-3}
MMD	3	ELU	$\mathcal{N}(0, 0.03^2)$	1.21×10^{-3}
MSE	3	ELU	$\mathcal{N}(0, 0.03^2)$	1.23×10^{-3}
Mean ² + Var	3	ELU	$\mathcal{N}(0, 0.03^2)$	1.19×10^{-3}

observed trajectories are available, the reconstruction of the distribution of model parameters using our method signifies a “worst case scenario” with no prior information on the magnitudes of model parameters.

We also investigate how the number of neurons in each layer, the number of hidden layers in the SNN model (Fig. 1), the initialization for the distribution of weights in the SNN model, as well as whether adopting the ResNet technique (He et al., 2016) for forward propagation would affect the accuracy of the reconstructed distribution of the kinetic parameters. From Table 2, SNNs with more than one hidden layer and 10 neurons in each layer can all reconstruct the distribution of $\mathbf{k}$ in Eq. (4.8) well. Thus, we do not need a wide or deep SNN model in Fig. 1 to reconstruct the distribution of $\mathbf{k}$. As an additional comparison with other loss functions, we train the SNN with 3 hidden layers and 10 neurons in each layer by minimizing the W_2 , MMD, MSE, and the Mean² + Var loss functions (defined in Appendix J) as was done in Example 4.1. Our results indicate that the SNN model, when equipped with appropriate numbers of hidden layers and neurons in each layer, has the ability to approximate the distribution of unknown model parameters well, and minimizing our local time-decoupled squared W_2 loss function can most efficiently train the SNN and yield the most accurate reconstructed distribution of unknown model parameters. Also, applying the ResNet technique can moderately increase the reconstruction accuracy when the number of hidden layers increases. Finally, initializing the weights and biases to small non-zero values is important, as the reconstruction error is huge if we set all weights and biases to zero. It is worth further investigation on how to optimally design the structure of the SNN model in Fig. 1, which is beyond the scope of this paper.

Finally, we consider reconstructing the distribution of parameters of a jump-diffusion process, in which intrinsic stochasticity from the Wiener process and Poisson process as well as uncertainty in the initial condition of the jump-diffusion process are coupled with uncertainty in the parameters governing the jump-diffusion process.

Example 4.4. We reconstruct the distribution of uncertain parameters governing a jump-diffusion process that can be used to describe the posited stock returns (Merton, 1976; Xia et al., 2024a). Instead of considering a deterministic jump magnitude function as in (Xia et al., 2024a, Example 2), we consider the following jump-diffusion process:

$$\begin{aligned}
dX_t &= 0.05dt + s\sqrt{|X_t|}dB_t + \int_U \xi X_t d\tilde{N}(\gamma(d\xi)dt), \quad t \in [0, 2], \\
s &\sim \sigma_0 \mathcal{N}(1, 1), \quad \xi \sim \mathcal{N}(\beta_0, \sigma_1^2), \quad X_0 \sim \mathcal{N}(2, \sigma_2^2).
\end{aligned} \quad (4.10)$$

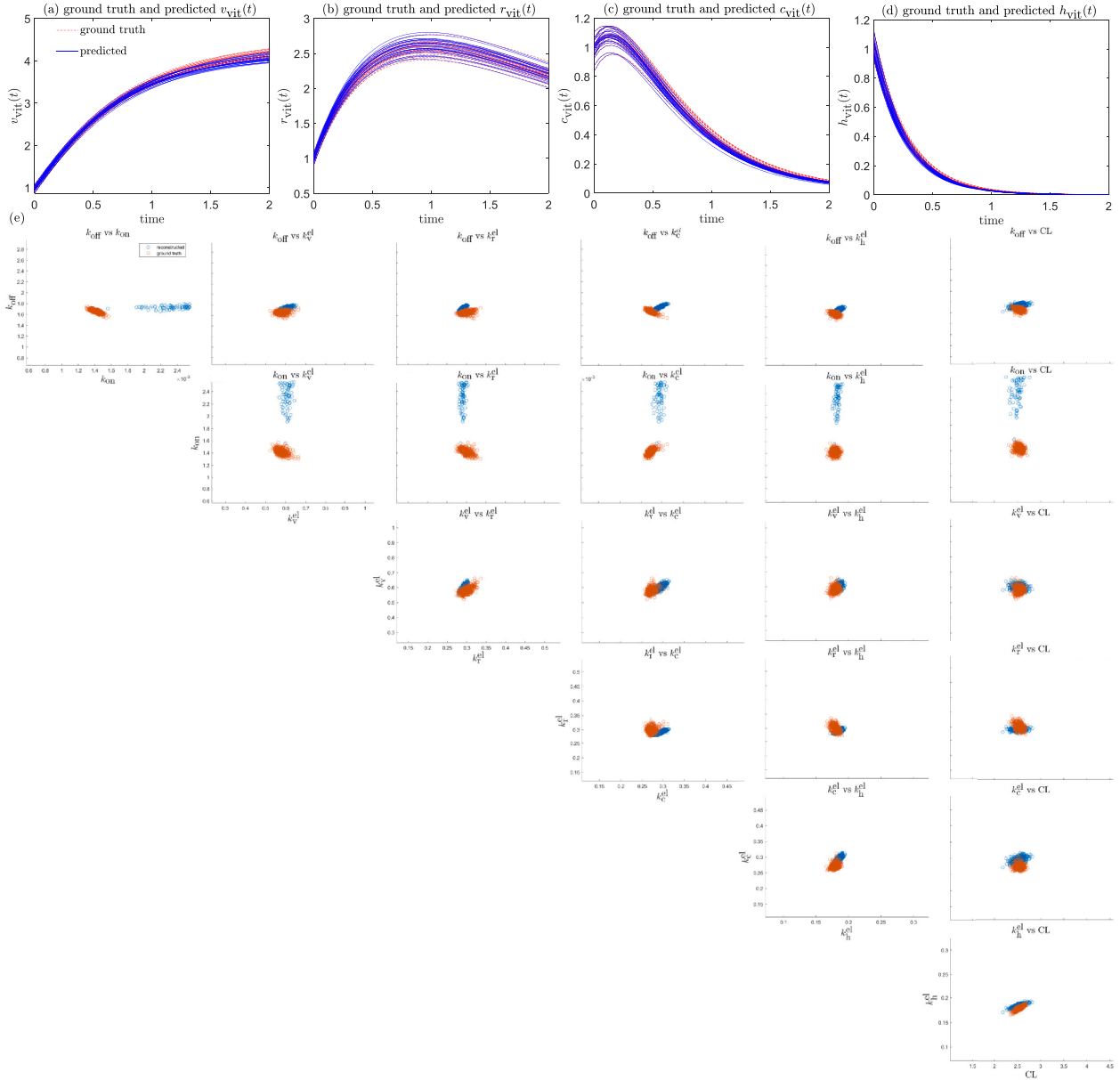


Fig. 5. (a)–(d) The first 50 out of 400 trajectories of the four quantities $v_{vit}(t)$, $r_{vit}(t)$, $c_{vit}(t)$, $h_{vit}(t)$ obtained with the parameter vector $\mathbf{k}$ sampled from the ground truth distribution Eq. (4.8) (red dashed lines) versus trajectories obtained by using the parameter vector $\mathbf{k}$ sampled from the reconstructed distribution generated by the trained SNN (blue solid lines). The SNN has 3 hidden layers and 10 neurons in each layer. The nodes and weights are initialized by independently sampling from $\mathcal{N}(0, 0.03^2)$. (e) The reconstructed joint distribution of any two kinetic parameters in Eq. (4.8). In all subplots, the red dots are sampled from the ground truth distribution while the blue dots are sampled from the reconstructed distribution. When using the ResNet technique for forward propagation, the scalar input of the SNN is 0.1 as inputting 1 leads to overflow during training.

In Eq. (4.10), the drift term represents the risk-free interest rate, the diffusion term stands for the fluctuation in the stock price, and the jump term represents events of paying dividends. In Eq. (4.10), $\tilde{N}$ is a compensated Poisson process defined in Eq. (2.11). We aim at reconstructing the distributions of s as well as ξ employing two separate SNNs using the distribution of $\hat{s}_0$ and $\hat{\xi}$ in the following approximate jump-diffusion process:

$$\begin{aligned} d\hat{X}_t &= 0.05dt + \hat{s}\sqrt{|\hat{X}_t|}dB_t + \int_U \hat{\xi}\hat{X}_t d\hat{N}(\gamma(d\xi)dt), \\ \hat{X}_0 &= X_0, \quad t \in [0, 2]. \end{aligned} \quad (4.11)$$

In Eq. (4.11), $\hat{N}$ is another compensated Poisson process that is independent of $\tilde{N}$. It has been shown in Xia et al. (2024a) that when $\xi \equiv 1$, larger values of s and β_0 make it more difficult to reconstruct the jump-diffusion process i.e., relative errors in the learned diffusion and jump functions get

larger because the trajectories are more sparsely distributed. Here, we carry out further experiments on the following cases:

1. Change the values of β_0 as well as the value of σ_1 in Eq. (4.10) to explore how the mean and variance of the jump magnitude function affect the reconstruction of distributions of $|s|$ and ξ (here we reconstruct the distribution of $|s|$ because $\hat{s}\sqrt{|\hat{X}_t|}dB_t$ and $|\hat{s}|\sqrt{|\hat{X}_t|}dB_t$ are identically distributed). Other parameters are $\sigma_0 = 0.3$, $\sigma_2 = 0.1$, and $\delta = 0.1$ (the neighborhood size in the loss function Eq. (J.1)).
2. Vary the value of σ_0 as well as the value of σ_1 to investigate how the uncertainty in the diffusion function and uncertainty in the jump magnitude affect the reconstruction of distributions of $|s|$ and ξ . Other parameters are $\beta_0 = 0.3$, $\sigma_2 = 0.1$, and $\delta = 0.1$.

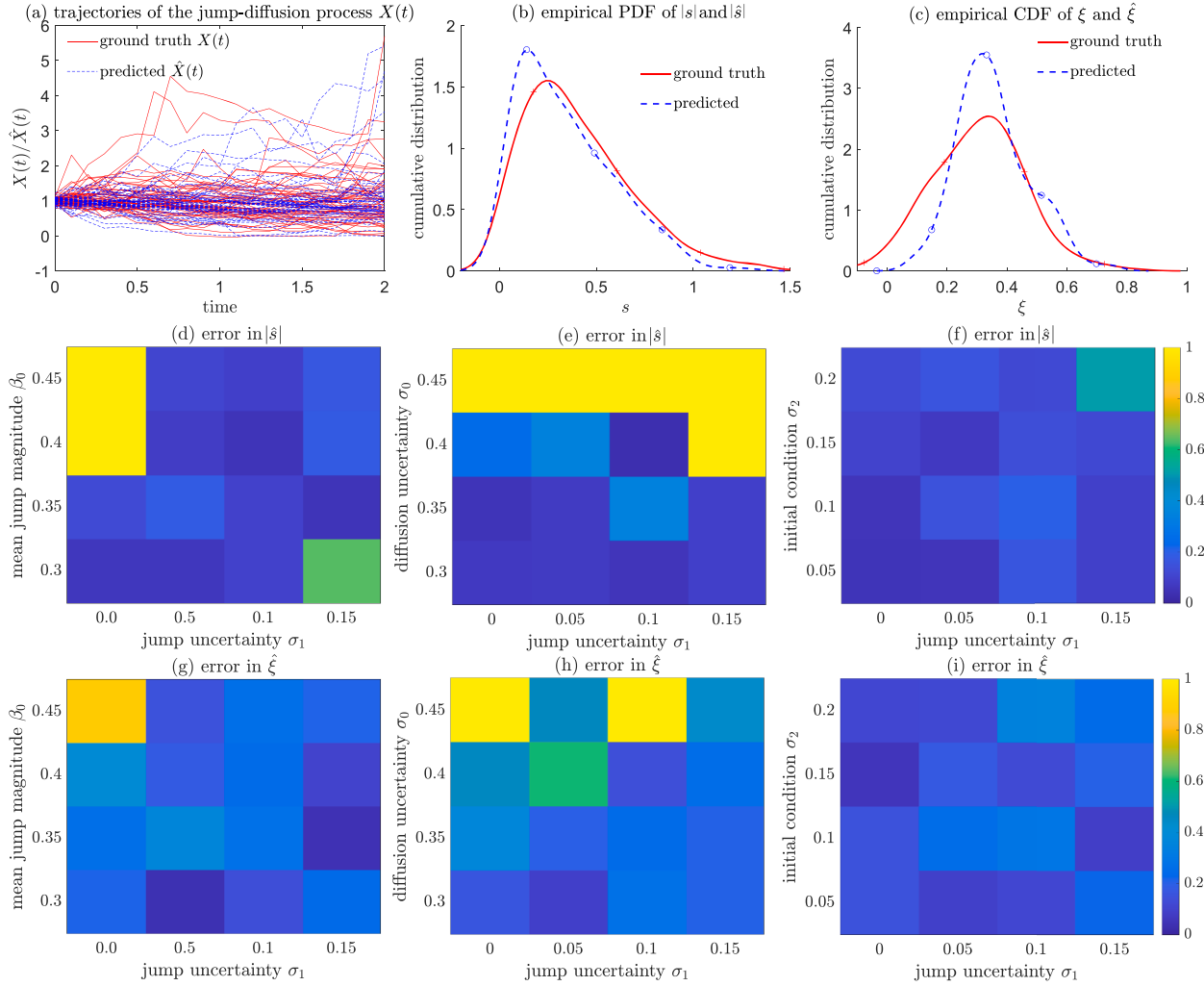


Fig. 6. (a) Ground truth trajectories generated from Eq. (4.10) versus the reconstructed trajectories generated from the approximate Eq. (4.11). For clarity, we plot 50 ground truth trajectories and 50 reconstructed trajectories. (b) The empirical probability density function ground truth $|s|$ versus the empirical probability density function of the reconstructed $|\hat{s}|$ in Eqs. (4.10) and (4.11). (c) The empirical probability density function ground truth ξ versus the empirical probability density function of the reconstructed $|\hat{\xi}|$ in Eqs. (4.10) and (4.11). In (a)–(c), $\sigma_0 = 0.3$, $\beta_0 = 0.35$, $\sigma_1 = 0.15$, $\sigma_2 = 0.1$, and $\delta = 0.1$ in the loss function Eq. (J.1). (d) and (g) The errors in the reconstructed distribution of $\hat{\sigma}$ and $\hat{\xi}$ for case 1 of Example 4.4 on Page 11, respectively. (e) and (h) The errors in the reconstructed distribution of $\hat{\sigma}$ and $\hat{\xi}$ for case 2 of Example 4.4 on Page 12, respectively. (f) and (i) The errors in the reconstructed distribution of $\hat{\sigma}$ and $\hat{\xi}$ for case 3 of Example 4.4 on Page 12, respectively.

3. Vary the value of σ_1 and σ_2 to determine how uncertainty in the initial condition and jump magnitude affects the reconstruction of $|s|$ and ξ . Other parameters are $\sigma_0 = 0.3$, $\beta_0 = 0.3$, $\sigma_2 = 0.1$, and $\delta = 0.1$.

From Fig. 6(a), the distribution of reconstructed trajectories of the approximate jump-diffusion process Eq. (4.11) match well with the distribution of ground truth jump-diffusion trajectories at each time t . Also, from Fig. 6(b), (c), the probability density functions of $|s|$ and ξ can be matched well by the probability density functions of $|\hat{s}|$ and $\hat{\xi}$ in Eq. (4.11), respectively. From Fig. 6(d), (g), when the mean and variance of ξ in the jump function of Eq. (4.10) become large, the reconstruction of the distributions of s and ξ is also less accurate. The increase in the mean β_0 of the jump magnitude impacts the reconstruction accuracy more than the increase in the variance σ_1 does. From Fig. 6(e), (h), a larger σ_0 characterizing greater uncertainty in the diffusion function of the jump-diffusion process Eq. (4.10) also leads to less accurate reconstructions of both s and ξ . Finally, from Fig. 6(f), (i), choosing a neighborhood size $\delta = 0.1$ works well for $\sigma_2 \in [0.05, 0.2]$ characterizing different levels of uncertainty in the initial condition and errors in both $|s|$ and $\hat{\xi}$ are well controlled.

When the form of the diffusion function in the jump-diffusion process in Eq. (4.10) is unknown but the diffusion function itself is deterministic, we can use a deterministic parameterized neural network to reconstruct the diffusion function as was done in Xia et al. (2024a) while simultaneously using an SNN to reconstruct the distribution of ξ in the jump function of Eq. (4.10). When the diffusion function is deterministic, the reconstruction accuracy of both the diffusion function and the distribution of ξ in Eq. (4.10) is good. The results are shown in Appendix K.

5. Summary and conclusions

In this paper, we proposed and analyzed a local time-decoupled squared W_2 distance method for reconstructing the distributions of parameters in specific dynamical systems from time-series data using an SNN model (Fig. 1). We also analyzed the SNN model and proved that it could approximate any continuous random variable as long as moderate assumptions are satisfied, making it a suitable model for reconstructing the distribution of parameters in a dynamical system. Our method took advantage of a previous local squared W_2 method

and a previous time-decoupled squared W_2 method so that both uncertainty in the initial state and intrinsic fluctuations such as the Wiener process in the dynamics could be considered. We showcased the effectiveness of our approach for reconstructing the distribution of model parameters in several dynamical systems such as ODEs, PDEs and SDEs. Our method outperformed several other benchmark statistical methods.

One limitation of our proposed method is that when the model parameters span across different magnitudes, the reconstruction of the distribution of model parameters whose magnitudes are smaller is less accurate. In practice, prior information on the distribution of model parameters might be available either from biophysical estimates or experimental assays. Thus, how to incorporate prior information of the model parameter distribution, such as confining the range of model parameters, is a potential future research direction. Specifically, it is illuminating to make comparisons with Bayesian approaches or Monte-Carlo simulation approaches, both of which can incorporate a prior distribution of model parameters, in terms of both accuracy and computational complexity. Also, it is helpful to consider applying our method to quantify uncertainties in model selection from time-series data or spatiotemporal data (Nardini et al., 2020). Taking into account measurement errors might also be necessary when such errors are not negligible (Nardini & Bortz, 2019). The computational complexity of evaluating our proposed local time-decoupled squared W_2 loss function is $O(N_T N_E[(N^\#(X_0; \delta))^3 \log(N^\#(X_0; \delta))])$, where $N^\#(X_0; \delta)$ refers to the number of trajectories in the data set such that their initial conditions satisfy $\|X(0) - X_0\| \leq \delta$ and N_T denotes the number of time steps. Therefore, considering utilizing entropic regularized Wasserstein distances and applying the Sinkhorn algorithm (Cuturi, 2013) to solve the corresponding optimal transport problems could be helpful in further reducing the computational complexity to $O(N_T N_E[(N^\#(X_0; \delta))^2])$. Additionally, it will be promising to apply our proposed local time-decoupled squared W_2 method for reconstructing the distribution of uncertain parameters in more complicated dynamical systems such as multidimensional SPDEs. It will be useful to analyze other stochastic neural networks' ability to approximate the distribution of uncertain model parameters, especially model parameters that take categorical values. On the other hand, more theoretical analysis on the Wasserstein-distance-based loss function, such as proving that minimizing a Wasserstein-distance-type loss function is sufficient for reconstructing the distribution of parameters in a dynamical system, can also be interesting, which would require the analysis of identifiability of model parameters, i.e., whether two different sets of model parameters would lead to different trajectories. Finally, more theoretical and empirical studies on how to design the optimal architecture, e.g. the depth and width, of the SNN model in Fig. 1 would be informative.

CRedit authorship contribution statement

Mingtao Xia: Writing – review & editing, Writing – original draft, Visualization, Validation, Supervision, Software, Project administration, Methodology, Investigation, Formal analysis, Conceptualization; **Qijing Shen:** Writing – review & editing, Visualization, Software, Methodology, Investigation; **Philip K. Maini:** Writing – review & editing, Resources, Investigation; **Eamonn A. Gaffney:** Writing – review & editing, Resources, Investigation; **Alex Mogilner:** Writing – review & editing, Resources, Investigation.

Data availability

All codes used in this manuscript are publicly available on GitHub at https://github.com/mtxia99/noisy_dynamical_system. In this study, no data were created or used.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgement

This work was supported in part through the NYU IT High Performance Computing resources, services, and staff expertise. The authors also acknowledge Dr. Jessica R. Crawshaw for helpful discussions.

Appendix A. The proof of Theorem 2.1

Here, we prove Theorem 2.1. Note that

$$W_{2,\delta}^{2,e}(X(t_i), \hat{X}(t_i)) = \frac{1}{N} \sum_{j=1}^N W_2^2(v_{X_{0,j},\delta}^e(t_i), \hat{v}_{X_{0,j},\delta}^e(t_i)), \quad (A.1)$$

where N is the number of observed trajectories and $X_{0,j}$ denotes the initial condition of the j th trajectory. Suppose $0 = t_0^1 < t_1^1 < \dots < t_{n_1}^1 = T$; $0 = t_0^2 < t_1^2 < \dots < t_{n_2}^2 = T$ are sets of grids on $[0, T]$. We define a third set of grids $0 = t_0^3 < \dots < t_{n_3}^3 = T$ such that $\{t_0^1, \dots, t_{n_1}^1\} \cup \{t_0^2, \dots, t_{n_2}^2\} = \{t_0^3, \dots, t_{n_3}^3\}$. Let $\Delta t := \max\{\max_i(t_{i+1}^1 - t_i^1), \max_j(t_{j+1}^2 - t_j^2), \max_k(t_{k+1}^3 - t_k^3)\}$. We denote $v_{X_{0,j},\delta}^e(t)$ and $\hat{v}_{X_{0,j},\delta}^e(t)$ to be the empirical conditional probability distributions of $X(t)$ and $\hat{X}(t)$ at time t conditioned on $|X(0) - X_{0,j}| \leq \delta$ and $|\hat{X}(0) - X_{0,j}| \leq \delta$, respectively. We shall show that:

$$\left| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1))(t_{i+1}^1 - t_i^1) - \sum_{i=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_i^3), \hat{X}(t_i^3))(t_{i+1}^3 - t_i^3) \right| \rightarrow 0, \quad (A.2)$$

as $\Delta t \rightarrow 0$.

First, suppose in the interval (t_i^1, t_{i+1}^1) , we have $t_i^1 = t_\ell^3 < t_{\ell+1}^3 < \dots < t_{\ell+s}^3 = t_{i+1}^1$, $s \geq 1$. For $s > 1$, since $t_{i+1}^1 - t_i^1 = \sum_{k=\ell}^{\ell+s-1} (t_{k+1}^3 - t_k^3)$, we have:

$$\begin{aligned} & \left| W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1))(t_{i+1}^1 - t_i^1) - \sum_{k=\ell}^{\ell+s-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3))(t_{k+1}^3 - t_k^3) \right| \\ & \leq \frac{1}{N} \sum_{j=1}^N \sum_{k=\ell+1}^{\ell+s-1} \left(W_2(v_{X_{0,j},\delta}^e(t_i^1), \hat{v}_{X_{0,j},\delta}^e(t_i^1)) + W_2(v_{X_{0,j},\delta}^e(t_k^3), \hat{v}_{X_{0,j},\delta}^e(t_k^3)) \right) \\ & \quad \times \left| W_2(v_{X_{0,j},\delta}^e(t_i^1), \hat{v}_{X_{0,j},\delta}^e(t_i^1)) - W_2(v_{X_{0,j},\delta}^e(t_k^3), \hat{v}_{X_{0,j},\delta}^e(t_k^3)) \right| (t_{k+1}^3 - t_k^3). \end{aligned} \quad (A.3)$$

Since $\|X\|$ and $\|\hat{X}\|$ are uniformly bounded, we have

$$W_2(v_{X_{0,j},\delta}^e(t_i^1), \hat{v}_{X_{0,j},\delta}^e(t_i^1)) \leq \sup_{X_0} \mathbb{E}[\|X(t)\|^2]^{\frac{1}{2}} + \mathbb{E}[\|\hat{X}(t)\|^2]^{\frac{1}{2}} \leq X + \hat{X} \quad (A.4)$$

and

$$W_2(v_{X_{0,j},\delta}^e(t_k^3), \hat{v}_{X_{0,j},\delta}^e(t_k^3)) \leq X + \hat{X}, \quad (A.5)$$

where $X, \hat{X}$ are the upper bounds in Eq. (2.7). We can take a specific coupling measure $\pi_{\delta, X_{0,j}}^*(X(t_i^1), X(t_k^3))$ such that if we regard

$$(X(t_i^1), X(t_k^3)) : \mathbb{R}^{2d} \rightarrow \mathbb{R}^{2d} \quad (A.6)$$

as a mapping from initial state space $(X_0, \tilde{X}_0) \in \mathbb{R}^{2d}$ to the solutions $(X(t_i^1), \tilde{X}(t_k^3)) \in \mathbb{R}^{2d}$ at times (t_i^1, t_k^3) satisfying $X(t_i^1) = X_0, \tilde{X}(t_k^3) = \tilde{X}_0$, then $\pi_{\delta, X_{0,j}}^*$ is the pushforward measure of the probability measure of the initial condition $v_{X_{0,j},\delta}^e \cdot \delta_{\tilde{X}_0=X_0}$ under the mapping $(X(t_i^1), X(t_k^3))$. Here,

$\nu_{X_{0,j},\delta}^e$ is the empirical conditional distribution of the initial condition X_0 conditioned on $|X_{0,j} - X_0| \leq \delta$. Then, for any $X_{0,j}$, we have

$$\begin{aligned} W_2^2(\nu_{X_{0,j},\delta}^e(t_i^1), \nu_{X_{0,j},\delta}^e(t_k^3)) &\leq \sup_{X_{0,j}} \mathbb{E}_{(X(t_i^1), X(t_k^3)) \sim \pi_{\delta, X_{0,j}}^*} [\|X(t_k^3) - X(t_i^1)\|_2^2] \\ &\leq \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i^1}^{t_{i+1}^1} \sum_{\ell=1}^d f_\ell^2(X(t), t; \theta) dt \right] (t_{i+1}^1 - t_i^1). \end{aligned} \quad (\text{A.7})$$

Similarly, we have

$$W_2^2(\hat{\nu}_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3)) \leq \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i}^{t_{i+1}} \sum_{\ell=1}^d f_\ell^2(X(t), t; \hat{\theta}) dt \right] (t_{i+1}^1 - t_i^1). \quad (\text{A.8})$$

Using the triangular inequality of the Wasserstein distance (Clement & Desch, 2008), we have

$$\begin{aligned} &|W_2(\nu_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_i^1)) - W_2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3))| \\ &\leq |W_2(\nu_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_i^1)) - W_2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3))| \\ &\quad + |W_2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3)) - W_2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3))| \\ &\leq W_2(\nu_{X_{0,j},\delta}^e(t_i^1), \nu_{X_{0,j},\delta}^e(t_k^3)) + W_2(\hat{\nu}_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3)). \end{aligned} \quad (\text{A.9})$$

Substituting Eqs. (A.4), (A.7), (A.8), and (A.9) into Eq. (A.3), we conclude that

$$\begin{aligned} &|W_2^2(\nu_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_i^1)) (t_{i+1}^1 - t_i^1) \\ &\quad - \sum_{k=\ell}^{\ell+s-1} W_2^2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3)) (t_{k+1}^3 - t_k^3)| \\ &\leq 2(X + \hat{X})(t_{i+1}^1 - t_i^1)(\sqrt{F_i \Delta t} + \sqrt{\hat{F}_i \Delta t}), \end{aligned} \quad (\text{A.10})$$

where

$$\begin{aligned} F_i &:= \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i^1}^{t_{i+1}^1} \sum_{\ell=1}^d f_\ell^2(X(t), t; \theta) dt \right], \\ \hat{F}_i &:= \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i^1}^{t_{i+1}^1} \sum_{\ell=1}^d f_\ell^2(X(t), t; \hat{\theta}) dt \right]. \end{aligned} \quad (\text{A.11})$$

Summing over i and j for the inequality (A.10), we have:

$$\begin{aligned} &| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1)) (t_{i+1}^1 - t_i^1) - \sum_{i=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3)) (t_{k+1}^3 - t_k^3) | \\ &\leq 2(X + \hat{X})T \max_i (\sqrt{F_i \Delta t} + \sqrt{\hat{F}_i \Delta t}). \end{aligned} \quad (\text{A.12})$$

Similarly,

$$\begin{aligned} &| \sum_{i=0}^{n_2-1} W_{2,\delta}^{2,e}(X(t_i^2), \hat{X}(t_i^2)) (t_{i+1}^2 - t_i^2) - \sum_{i=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3)) (t_{k+1}^3 - t_k^3) | \\ &\leq 2(X + \hat{X})T \max_i (\sqrt{F_i' \Delta t} + \sqrt{\hat{F}_i' \Delta t}), \end{aligned} \quad (\text{A.13})$$

where

$$\begin{aligned} F_i' &:= \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i^2}^{t_{i+1}^2} \sum_{\ell=1}^d f_\ell^2(X(t), t; \hat{\theta}) dt \right], \\ \hat{F}_i' &:= \sup_{X_{0,j}} \mathbb{E} \left[\int_{t_i^2}^{t_{i+1}^2} \sum_{\ell=1}^d f_\ell^2(X(t), t; \hat{\theta}) dt \right]. \end{aligned} \quad (\text{A.14})$$

Thus, as $\Delta t \rightarrow 0$,

$$| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(\nu(t_i^1), \hat{\nu}(t_i^1)) (t_{i+1}^1 - t_i^1) - \sum_{i=0}^{n_2-1} W_{2,\delta}^{2,e}(\nu(t_i^2), \hat{\nu}(t_i^2)) (t_{i+1}^2 - t_i^2) | \rightarrow 0,$$

which implies the limit

$$\lim_{\max(t_{i+1}^1 - t_i^1) \rightarrow 0} \sum_{i=0}^{N-1} W_{2,\delta}^{2,e}(\nu(t_i^1), \hat{\nu}(t_i^1)) (t_i^1 - t_{i-1}^1) \quad (\text{A.16})$$

exists. Therefore, the local time-decoupled squared W_2 distance in Eq. (2.6):

$$\bar{W}_{2,\delta}^{2,e}(X, \hat{X}) := \int_0^T W_{2,\delta}^{2,e}(X(t), \hat{X}(t)) dt \quad (\text{A.17})$$

is well-defined.

Appendix B. Proof of Corollary 2.1

The proof of Corollary 2.1 is similar to the proof of Theorem 2.1 and the proof of Theorem 3.1 in Xia et al. (2024a). Suppose $0 = t_0^1 < t_1^1 < \dots < t_{n_1}^1 = T$; $0 = t_0^2 < t_1^2 < \dots < t_{n_2}^2 = T$ are two sets of grids on $[0, T]$. We define a third set of grids $0 = t_0^3 < \dots < t_{n_3}^3 = T$ such that $\{t_0^1, \dots, t_{n_1}^1\} \cup \{t_0^2, \dots, t_{n_2}^2\} = \{t_0^3, \dots, t_{n_3}^3\}$. Let $\Delta t := \max\{\max_i(t_{i+1}^1 - t_i^1), \max_j(t_{j+1}^2 - t_j^2), \max_k(t_{k+1}^3 - t_k^3)\}$. We denote $\nu_{X_{0,j},\delta}^e(t)$ and $\hat{\nu}_{X_{0,j},\delta}^e(t)$ to be the empirical conditional probability distributions of $X(t)$ and $\hat{X}(t)$ at time t conditioned on $|X(0) - X_{0,j}| \leq \delta$ and $|\hat{X}(0) - X_{0,j}| \leq \delta$, respectively. We need to show that

$$\begin{aligned} &| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1)) (t_{i+1}^1 - t_i^1) - \sum_{i=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3)) (t_{k+1}^3 - t_k^3) | \rightarrow 0. \end{aligned} \quad (\text{B.1})$$

For $j = 1, 2$, we define:

$$\begin{aligned} F_i^j &:= \sup_{X_{0,\theta}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d f_\ell^2(X(t^-), t^-; \theta) dt \right], \\ \hat{F}_i^j &:= \sup_{X_{0,\hat{\theta}}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d f_\ell^2(\hat{X}(t^-), t^-; \hat{\theta}) dt \right], \\ \Sigma_i^j &:= \sup_{X_{0,\theta}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d \sum_{j=1}^m \sigma_{\ell,j}^2(X(t^-), t^-) dt \right], \\ \hat{\Sigma}_i^j &:= \sup_{X_{0,\hat{\theta}}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d \sum_{j=1}^m \sigma_{\ell,j}^2(\hat{X}(t^-), t^-; \hat{\theta}) dt \right], \\ B_i^j &:= \sup_{X_{0,\theta}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d \int_U \beta_\ell^2(X(t^-), \xi, t^-; \theta) \nu(d\xi) dt \right], \\ \hat{B}_i^j &:= \sup_{X_{0,\hat{\theta}}} \mathbb{E} \left[\int_{t_i^j}^{t_{i+1}^j} \sum_{\ell=1}^d \int_U \hat{\beta}_\ell^2(\hat{X}(t^-), \xi, t^-; \hat{\theta}) \nu(d\xi) dt \right]. \end{aligned} \quad (\text{B.2})$$

Similar to the proof of Theorem 2.1, we find that:

$$\begin{aligned} &|W_2^2(\nu_{X_{0,j},\delta}^e(t_i^1), \hat{\nu}_{X_{0,j},\delta}^e(t_i^1)) (t_{i+1}^1 - t_i^1) \\ &\quad - \sum_{k=\ell}^{\ell+s-1} W_2^2(\nu_{X_{0,j},\delta}^e(t_k^3), \hat{\nu}_{X_{0,j},\delta}^e(t_k^3)) (t_{k+1}^3 - t_k^3)| \\ &\leq 2(X + \hat{X})(t_{i+1}^1 - t_i^1) (\sqrt{F_i^1 \Delta t} + \Sigma_i^1 + B_i^1 + \sqrt{\hat{F}_i^1 \Delta t} + \hat{\Sigma}_i^1 + \hat{B}_i^1). \end{aligned} \quad (\text{B.3})$$

Summing over $i = 0, \dots, n_1 - 1$ and $j = 1, \dots, N$, we can obtain

$$\begin{aligned} &| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1)) (t_{i+1}^1 - t_i^1) - \sum_{k=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3)) (t_{k+1}^3 - t_k^3) | \\ &\leq 2(X + \hat{X})T \max_i (\sqrt{F_i^1 \Delta t} + \Sigma_i^1 + B_i^1 + \sqrt{\hat{F}_i^1 \Delta t} + \hat{\Sigma}_i^1 + \hat{B}_i^1). \end{aligned} \quad (\text{B.4})$$

Similarly, we have:

$$\left| \sum_{i=0}^{n_2-1} W_{2,\delta}^{2,e}(X(t_i^2), \hat{X}(t_i^2))(t_{i+1}^2 - t_i^2) - \sum_{k=0}^{n_3-1} W_{2,\delta}^{2,e}(X(t_k^3), \hat{X}(t_k^3))(t_{k+1}^3 - t_k^3) \right| \leq 2(X + \hat{X})T \max \left(\sqrt{F_i^1 \Delta t + \Sigma_i^2 + B_i^2} + \sqrt{\hat{F}_i^2 \Delta t + \hat{\Sigma}_i^2 + \hat{B}_i^2} \right). \quad (\text{B.5})$$

Thus,

$$\left| \sum_{i=0}^{n_1-1} W_{2,\delta}^{2,e}(X(t_i^1), \hat{X}(t_i^1))(t_{i+1}^1 - t_i^1) - \sum_{i=0}^{n_2-1} W_{2,\delta}^{2,e}(X(t_i^2), \hat{X}(t_i^2))(t_{i+1}^2 - t_i^2) \right| \leq 2(X + \hat{X})T \max \left(\sqrt{F_i^1 \Delta t + \Sigma_i^1 + B_i^1} + \sqrt{\hat{F}_i^2 \Delta t + \hat{\Sigma}_i^1 + \hat{B}_i^1} \right) + 2(X + \hat{X})T \max \left(\sqrt{F_i^2 \Delta t + \Sigma_i^2 + B_i^2} + \sqrt{\hat{F}_i^3 \Delta t + \hat{\Sigma}_i^2 + \hat{B}_i^2} \right). \quad (\text{B.6})$$

Since f, σ, β are uniformly bounded, $F_i^j, \Sigma_i^j, B_i^j, \hat{F}_i^j, \hat{\Sigma}_i^j, \hat{B}_i^j \rightarrow 0$ as $\Delta t \rightarrow 0$ uniformly for $j = 1, 2$. Thus, the limit

$$\tilde{W}_{2,\delta}^{2,e}(X, \hat{X}) := \int_0^T W_{2,\delta}^{2,e}(X(t), \hat{X}(t)) dt \quad (\text{B.7})$$

exists.

Appendix C. The proof of Theorem 2.2

In this section, we prove Theorem 2.2. First, consider $\tilde{X}$ that solves the following jump-diffusion process:

$$d\tilde{X}(t) = f(\tilde{X}(t), t; \hat{\theta})dt + \sigma(\tilde{X}(t), t; \hat{\theta})dB_t + \int_U \beta(\tilde{X}(t), \xi, t; \hat{\theta})\tilde{N}(dt, v(d\xi)), \quad \tilde{X}(0) = X(0). \quad (\text{C.1})$$

Using Theorem 2.1 in Xia et al. (2024a), we have:

$$\mathbb{E}[\|X(t) - \tilde{X}(t)\|^2] \leq \mathbb{E}[H(t)|X(0)] \exp\left((2C + 1 + (2C + 1)m + (2C + 1)\gamma(U))td\right), \quad (\text{C.2})$$

where $X(0)$ is the initial condition and $H(t)$ is defined as

$$H(t) := \mathbb{E}\left[\sum_{i=1}^d \int_0^t (f_i(X(s^-), s^-; \theta) - f_i(X(s^-), s^-; \hat{\theta}))^2 ds\right] + \mathbb{E}\left[\sum_{i=1}^d \int_0^t \sum_{j=1}^m (\sigma_{i,j}(X(s^-), s^-; \theta) - \sigma_{i,j}(X(s^-), s^-; \hat{\theta}))^2 ds\right] + \mathbb{E}\left[\sum_{i=1}^d \int_0^t \int_U (\beta_i(X(s^-), \xi, s^-; \theta) - \beta_i(X(s^-), \xi, s^-; \hat{\theta}))^2 \gamma(d\xi) ds\right] \leq Ct(1 + m + \gamma(U))\|\theta - \hat{\theta}\|^2. \quad (\text{C.3})$$

Since the probability distribution of $\hat{X}(t) \in \mathbb{R}^d$ is the same as the probability distribution of $\tilde{X}(t)$ for any $t \in [0, T]$, we conclude that

$$W_2^2(v_{X_0}(t), \hat{v}_{X_0}(t)) = W_2^2(v_{X_0}(t), \tilde{v}_{X_0}(t)), \quad (\text{C.4})$$

where $v_{X_0}(t)$, $\hat{v}_{X_0}(t)$, and $\tilde{v}_{X_0}(t)$ are the probability distributions of $X(t)$, $\hat{X}(t)$, and $\tilde{X}(t)$ given the same initial condition X_0 , respectively. Given any coupled distribution of $\theta, \hat{\theta}$ denoted by $\pi(\mu, \hat{\mu})$ such that its marginal distributions coincide with the probability distributions of θ and $\hat{\theta}$, we denote $\pi^*(X(t), \tilde{X}(t))$ to be its pushforward probability measure for $(X(t), \tilde{X}(t)) \in \mathbb{R}^{2d}$. We can easily check that the marginal distributions of π^* coincide with $v_{X_0}(t)$ and $\tilde{v}_{X_0}(t)$, respectively. From Eqs. (C.2), (C.3), we have

$$\mathbb{E}_{(X(t), \tilde{X}(t)) \sim \pi^*}[\|X(t) - \tilde{X}(t)\|^2] \leq Ct(1 + m + \gamma(U))\mathbb{E}_{(\theta, \hat{\theta}) \sim \pi(\mu, \hat{\mu})}[\|\theta - \hat{\theta}\|^2]$$

$$\times \exp\left((2C + 1 + (2C + 1)m + (2C + 1)\gamma(U))td\right). \quad (\text{C.5})$$

Taking the infimum of π over all coupled distributions of $(\theta, \hat{\theta})$ whose marginal distributions are μ and $\hat{\mu}$, respectively, we conclude that

$$W_2^2(v_{X_0}(t), \hat{v}_{X_0}(t)) \leq C_1 t \exp(C_0 t) W_2^2(\mu, \hat{\mu}), \quad (\text{C.6})$$

where

$$C_0 := (2C + 1 + (2C + 1)m + (2C + 1)\gamma(U))d, \quad C_1 := C(1 + m + \gamma(U)) \quad (\text{C.7})$$

are two constants.

Using the triangular inequality of the W_2 distance (Clement & Desch, 2008), we have

$$\begin{aligned} W_2(v_{X_0,\delta}^e(t), \hat{v}_{X_0,\delta}^e(t)) &\leq W_2(v_{X_0}^e(t), \hat{v}_{X_0}^e(t)) + W_2(v_{X_0}^e(t), v_{X_0,\delta}^e(t)) \\ &\quad + W_2(\hat{v}_{X_0}^e(t), \hat{v}_{X_0,\delta}^e(t)) \\ &\leq W_2(v_{X_0}(t), \hat{v}_{X_0}(t)) + W_2(\hat{v}_{X_0}(t), v_{X_0}^e(t)) \\ &\quad + W_2(v_{X_0}(t), v_{X_0,\delta}^e(t)) \\ &\quad + W_2(v_{X_0,\delta}^e(t), \hat{v}_{X_0,\delta}^e(t)) + W_2(\hat{v}_{X_0,\delta}^e(t), \hat{v}_{X_0}^e(t)), \end{aligned} \quad (\text{C.8})$$

where $v_{X_0,\delta}^e(t)$ and $\hat{v}_{X_0,\delta}^e(t)$ are the empirical conditional probability distributions of $X(t)$ and $\hat{X}(t)$ at time t conditioned on $|X(0) - X_0| \leq \delta$ and $|\hat{X}(0) - X_0| \leq \delta$, respectively. For any $\theta \in \mathbb{R}^\ell$, consider

$$\begin{aligned} dX(t; X_0) &= f(X(t; X_0), t; \theta)dt + \sigma(X(t; X_0), t; \theta)dB_t \\ &\quad + \int_U \beta(X(t; X_0), \xi, t; \theta)\tilde{N}(dt, \gamma(d\xi)), \\ dX(t; X'_0) &= f(X(t; X'_0), t; \theta)dt + \sigma(X(t; X'_0), t; \theta)dB_t \\ &\quad + \int_U \beta(X(t; X'_0), \xi, t; \theta)\tilde{N}(dt, \gamma(d\xi)), \\ X(0; X_0) &= X_0, \quad X(0; X'_0) = X'_0, \quad \|X_0 - X'_0\| \leq \delta. \end{aligned} \quad (\text{C.9})$$

Using the stochastic Gronwall lemma (Mehri & Scheutzwow, 2019, Theorem 2.2) and (Xia et al., 2024a, Theorem 2.1), we conclude that:

$$\begin{aligned} \mathbb{E}[\|X(t; X_0) - X(t; X'_0)\|_2^2] &\leq \exp\left((2dC + dCm + dC\gamma(U) + 1)t\right) \mathbb{E}[\|X_0 - X'_0\|^2]. \end{aligned} \quad (\text{C.10})$$

Therefore, we have:

$$W_2(v_{X_0}^e(t), v_{X_0,\delta}^e(t)) \leq \delta \exp\left(\frac{C_0 t}{2}\right). \quad (\text{C.11})$$

Similarly, we conclude that:

$$W_2(\hat{v}_{X_0}^e(t), \hat{v}_{X_0,\delta}^e(t)) \leq \delta \exp\left(\frac{C_0 t}{2}\right). \quad (\text{C.12})$$

Eq. (C.6) also holds if we replace $\hat{v}_{X_0}(t)$ on the LHS of Eq. (C.6) with the empirical distribution $\hat{v}_{X_0}^e(t)$ and then replace $\hat{\mu}$ with $\mu_{X_0}^e$ on the RHS, i.e.,

$$W_2(v_{X_0}(t), v_{X_0,\delta}^e(t)) \leq \sqrt{C_1 t} \exp\left(\frac{C_0 t}{2}\right) W_2(\mu, \mu_{X_0}^e). \quad (\text{C.13})$$

Similarly,

$$W_2^2(\hat{v}_{X_0}(t), \hat{v}_{X_0,\delta}^e(t)) \leq \sqrt{C_1 t} \exp\left(\frac{C_0 t}{2}\right) W_2(\hat{\mu}, \hat{\mu}_{X_0}^e). \quad (\text{C.14})$$

In Eqs. (C.13) and (C.14), $\mu_{X_0}^e$ and $\hat{\mu}_{X_0}^e$ denote the empirical distributions of θ and $\hat{\theta}$, respectively. Finally, by plugging Eqs. (C.11), (C.12), (C.13), (C.14), (C.6) into Eq. (C.8), we conclude that:

$$\begin{aligned} W_2(v_{X_0,\delta}^e(t), \hat{v}_{X_0,\delta}^e(t)) &\leq 2\delta \exp\left(\frac{C_0 t}{2}\right) \\ &\quad + \sqrt{C_1 t} \exp(C_0 t) (W_2(\mu, \mu_{X_0}^e) + W_2(\hat{\mu}, \hat{\mu}_{X_0}^e) + W_2(\mu, \hat{\mu})). \end{aligned} \quad (\text{C.15})$$

Squaring both sides of Eq. (C.15), integrating over time, and taking the expectation w.r.t. the empirical probability measure of the initial condition X_0 , we conclude that:

$$\mathbb{E}[\tilde{W}_{2,\delta}^{2,e}(X, \hat{X})] \leq 8T\delta^2 \exp(C_0 T)$$

$$\begin{aligned}
& + \sum_{j=1}^N \frac{1}{N} \frac{6C_1}{C_0} T \exp(C_0 T) (W_2^2(\mu, \hat{\mu}) \\
& + \mathbb{E}[W_2^2(\mu_{X_{0,j}}^e, \mu)] + \mathbb{E}[W_2^2(\mu_{X_{0,j}}^e, \hat{\mu})]). \quad (C.16)
\end{aligned}$$

Specifically, C_0 grows linearly with the dimensionality d . Finally, taking the expectation of Eq. (C.16) and applying the empirical error bound of the squared W_2 distance in Fournier and Guillin (2015), we conclude that:

$$\begin{aligned}
\sum_{j=1}^N \frac{1}{N} \mathbb{E}[W_2^2(\mu_{X_{0,j}}^e, \mu)] & \leq C_2 \mathbb{E}[h(N^\#(X_0; \delta), \ell)] \Theta_6^{\frac{1}{3}}, \\
\sum_{j=1}^N \frac{1}{N} \mathbb{E}[W_2^2(\mu_{X_{0,j}}^e, \hat{\mu})] & \leq C_2 \mathbb{E}[h(N^\#(X_0; \delta), \ell)] \hat{\Theta}_6^{\frac{1}{3}}, \quad (C.17)
\end{aligned}$$

where C_2 is a constant used in (Fournier & Guillin, 2015, Theorem 1) and h is defined in Eq. (2.19). This completes the proof of Theorem 2.2.

Appendix D. Reconstructing the distribution of parameters in an SPDE

We consider the following parabolic SPDE with a Dirichlet boundary condition studied in Grecksch and Kloeden (1996):

$$\begin{aligned}
dU(x, t; \theta) &= A(\theta)U(x, t; \theta) + f(U(x, t; \theta); \theta)dt + g(U(x, t; \theta); \theta)dB_t, \\
x &\in D, t \in [0, T], \\
U(x, 0; \theta) &\in H_0^{1,2}(\Omega), U(x, 0) \sim \nu_0, \quad U(x, t; \theta) = 0, x \in \partial D. \quad (D.1)
\end{aligned}$$

In Eq. (D.1), B_t is a standard scalar Wiener process and D is a bounded domain in $\mathbb{R}^d$ with sufficiently smooth boundary ∂D . $H_0^{1,2}(\Omega)$ is the space of functions $U : \Omega \rightarrow \mathbb{R}$ that vanish on $\partial\Omega$ such that U and its first-order generalized derivatives belong to $L^2(\Omega)$ equipped with the norm $\|U\|_{L^2} := \int_D U^2(x, t)dx$. ν_0 is a probability measure defined on the Sobolev space $B(H_0^{1,2}(\Omega))$. A is a linear operator densely defined in $L^2(\Omega)$ such that for $U \in H^{1,2}(\Omega)$, $AU \in L^2(\Omega)$. A is strongly monotone, i.e., there is a constant $\alpha > 0$ such that

$$(-AU, U) \geq \alpha \|U\|_{H^{1,2}}^2, \quad \forall U \in H_0^{1,2}(\Omega), \quad (D.2)$$

where $(\cdot, \cdot)$ is the inner product. In addition, f and g , which map either $L^2(\Omega)$ or $H_0^{1,2}(\Omega)$ into itself, are formed from real-valued functions of a real variable with uniformly bounded derivatives of appropriate order.

We reconstruct the distribution of the model parameter θ in Eq. (D.1) using an approximate model

$$\begin{aligned}
d\hat{U}(x, t; \hat{\theta}) &= A(\hat{\theta})\hat{U}(x, t; \hat{\theta}) + f(\hat{U}(x, t; \hat{\theta}); \hat{\theta})dt + g(\hat{U}(x, t; \hat{\theta}); \hat{\theta})d\hat{B}_t, \\
x &\in D, t \in [0, T], \\
\hat{U}(x, 0) &= U(x, 0; \theta) = U(x, 0), \quad U(x, t; \hat{\theta}) = 0, x \in \partial D. \quad (D.3)
\end{aligned}$$

In Eq. (D.3), $\hat{B}_t$ is another standard scalar Wiener process independent of B_t in Eq. (D.1). We use the Itô-Galerkin scheme for spatial discretization of Eqs. (D.1) and (D.3):

$$\begin{aligned}
dU_n(t; \theta) &= (A_n(U_n(t; \theta); \theta) + f_n(U_n(t; \theta); \theta))dt + g_n(U_n(t; \theta); \theta)dB_t, \\
d\hat{U}_n(t; \hat{\theta}) &= (A_n(\hat{U}_n(t; \hat{\theta}); \hat{\theta}) + f_n(\hat{U}_n(t; \hat{\theta}); \hat{\theta}))dt + g_n(\hat{U}_n(t; \hat{\theta}); \hat{\theta})d\hat{B}_t. \quad (D.4)
\end{aligned}$$

In Eq. (D.4),

$$U_n(x, t; \theta) := \sum_{j=1}^n u_j(t; \theta) \varphi_j(x) \in X_n, \quad \hat{U}_n(x, t; \hat{\theta}) := \sum_{j=1}^n \hat{u}_j(t; \hat{\theta}) \varphi_j(x) \in X_n \quad (D.5)$$

refer to the spectral approximations of $U(x, t; \theta)$ and $\hat{U}(x, t; \hat{\theta})$ in Eqs. (D.1) and (D.3), respectively. X_n is the n -dimensional subspace of $H_0^1(\Omega)$ spanned by the basis functions $\{\varphi_1, \dots, \varphi_n\}$. In Eq. (D.4),

$$A_n(U; \theta) := P_n(A(\theta)U), \quad f_n(U; \theta) := P_n(f(\theta)U), \quad g_n(U; \theta) := P_n(g(\theta)U),$$

where P_n denotes the projection of $L^2(\Omega)$ or $H_0^{1,2}(\Omega)$ onto X_n . In the spatiotemporal SPDEs (D.1) and (D.3), we assume that for all θ , the operator $-A$ has the same set of corresponding eigenfunctions

$$-A(\theta)\varphi_j = \lambda_j(\theta)\varphi_j, \quad j = 1, 2, \dots, n, \quad \lambda_j \leq \lambda_{j+1}.$$

$\varphi_i \in H_0^{1,2}(D), i = 1, 2, \dots$ forms an orthonormal basis in $L^2(D)$ with $\|\varphi_j\|_{L^2} = 1$ and $\lambda_j(\theta) \rightarrow \infty$ uniformly in θ as $j \rightarrow \infty$. We can prove the following result.

Corollary D.1. We assume that $A(\theta)$, interpreted as mappings of $H_0^1(\Omega)$ into $L^2(\Omega)$ and $f(\theta)$ and $g(\theta)$, interpreted as mappings of $H_0^1(\Omega)$ into itself, are Lipschitz continuous:

$$\begin{aligned}
\|A(\theta)(U, \theta) - A(\hat{\theta})(\hat{U}, \hat{\theta})\|_{L^2} &\leq L(\|U - \hat{U}\|_{H^{1,2}} + \|\theta - \hat{\theta}\|), \\
\|f(U; \theta) - f(\hat{U}; \hat{\theta})\|_{H^{1,2}} &\leq L(\|U - \hat{U}\|_{H^{1,2}} + \|\theta - \hat{\theta}\|), \\
\|g(U; \theta) - g(\hat{U}; \hat{\theta})\|_{H^{1,2}} &\leq L(\|U - \hat{U}\|_{H^{1,2}} + \|\theta - \hat{\theta}\|), \quad L \leq \infty. \quad (D.6)
\end{aligned}$$

Furthermore, we assume that

$$\mathbb{E}[\|\theta\|^6] \leq \Theta_6, \quad \mathbb{E}[\|\hat{\theta}\|^6] \leq \hat{\Theta}_6. \quad (D.7)$$

Then, we have the following bound for the local time-decoupled squared distance between the probability measures of U and $\hat{U}$:

$$\begin{aligned}
\bar{W}_{2,\delta}^{2,e}(U, \hat{U}) &\leq 3 \cdot \left(8C_0(\beta_n; n)\delta^2 T \exp(C_0(\beta_n; n)T) \right. \\
&\quad + \frac{6C_1(\beta_n)T}{C_0(\beta_n; n)} \exp(C_0(\beta_n; n)T) \\
&\quad \times (W_2^2(\mu, \hat{\mu}) + (\Theta_6^{\frac{1}{3}} + \hat{\Theta}_6^{\frac{1}{3}})2T\mathbb{E}[h(N^\#(U_n(0; \theta); \delta); \ell)]) \\
&\quad + 3T \sup_{\theta, U(x,0;\theta)} K_{T,U(\cdot,0),\theta} \lambda_{N+1}^{-1}(\theta) \\
&\quad \left. + 3T \sup_{\hat{\theta}, U(x,0;\hat{\theta})} K_{T,U(\cdot,0),\hat{\theta}} \lambda_{N+1}^{-1}(\hat{\theta}) \right) \quad (D.8)
\end{aligned}$$

In Eq. (D.8), $K_{T,U(\cdot,0),\theta}$ is a constant that depends on $T, U(x, 0)$ and θ , while $C_i(\beta_n)$ are constants in Theorem 2.2 that grow at most linearly with n . β_n is another constant depending on the eigenvalues $\{\lambda_i\}_{i=1}^n$. The vector $U_n(0; \theta) := (u_1(0; \theta), \dots, u_n(0; \theta))$ refers to the spectral expansion of the initial condition. $\bar{W}_{2,\delta}^{2,e}(U, \hat{U})$ is the local time-decoupled squared distance between the probability measures of U and $\hat{U}$:

$$\bar{W}_{2,\delta}^{2,e}(U, \hat{U}) := \int_0^T W_{2,\delta}^{2,e}(U(x, t; \theta), \hat{U}(x, t; \hat{\theta}))dt, \quad (D.9)$$

and

$$W_{2,\delta}^{2,e}(U(x, t; \theta), \hat{U}(x, t; \hat{\theta})) := \int W_2^2(\nu_{U_0,\delta}^e(t), \nu_{\hat{U}_0,\delta}^e(t))\nu_0^e(dU_0), \quad (D.10)$$

where $\nu_0^e(dU_0)$ is the empirical distribution of the initial condition $U(\cdot, 0)$, and $\nu_{U_0,\delta}^e(t)$ and $\nu_{\hat{U}_0,\delta}^e(t)$ are the empirical conditional distributions of $U(x, t; \theta)$ and $\hat{U}(x, t; \hat{\theta})$ at time t conditioned on $\|U(x, 0) - U_0\|_{L^2} \leq \delta$ and $\|\hat{U}(x, 0) - \hat{U}_0\|_{L^2} \leq \delta$, respectively. In Eq. (D.10), the W_2 distance between $\nu_{U_0,\delta}^e(t)$ and $\nu_{\hat{U}_0,\delta}^e(t)$ is defined as

$$W_2(\nu_{U_0,\delta}^e(t), \nu_{\hat{U}_0,\delta}^e(t)) := \inf_{\pi(v,v)} \mathbb{E}_{(U,\hat{U}) \sim \pi(\nu_{U_0,\delta}^e(t), \nu_{\hat{U}_0,\delta}^e(t))} [\|U - \hat{U}\|_{L^2}^2]^{\frac{1}{2}}. \quad (D.11)$$

Proof. First, we show that A_n is also Lipschitz continuous:

$$\begin{aligned}
\|A_n(U_n; \theta) - A_n(\hat{U}_n; \hat{\theta})\|_{L^2} &= \|P_n(A(U_n; \theta) - A(\hat{U}_n; \hat{\theta}))\|_{L^2} \\
&\leq \|A(U_n; \theta) - A(\hat{U}_n; \hat{\theta})\|_{L^2} \\
&\leq L(\|U_n - \hat{U}_n\|_{H^{1,2}} + \|\theta - \hat{\theta}\|). \quad (D.12)
\end{aligned}$$

Because X_n is a finite dimensional space, there exists a constant β_n depending on $\varphi_1, \dots, \varphi_n$ such that $\forall U_n \in X_n, \|U\|_{H^{1,2}} \leq \beta_n \|U\|_{L^2}$. Thus,

$$\begin{aligned}
\|A_n(U_n; \theta) - A_n(\hat{U}_n; \hat{\theta})\|_{H^{1,2}} &\leq \beta_n \|A_n(U_n; \hat{\theta}) - A_n(\hat{U}_n; \hat{\theta})\|_{L^2} \\
&\leq \beta_n L(\|U_n - \hat{U}_n\|_{L^2} + \|\theta - \hat{\theta}\|). \quad (D.13)
\end{aligned}$$

Similarly, f_n and g_n can also be shown to be Lipschitz continuous in U_n and θ :

$$\begin{aligned} \|f_n(U_n; \theta) - f_n(\hat{U}_n; \hat{\theta})\|_{L^2} &\leq \|f_n(U_n; \theta) - f_n(\hat{U}_n; \hat{\theta})\|_{H^{1,2}} \\ &\leq \beta_n L(\|U_n - \hat{U}_n\|_{L^2} + \|\theta - \hat{\theta}\|), \\ \|g_n(U_n; \theta) - g_n(\hat{U}_n; \hat{\theta})\|_{L^2} &\leq \|g_n(U_n; \theta) - g_n(\hat{U}_n; \hat{\theta})\|_{H^{1,2}} \\ &\leq \beta_n L(\|U_n - \hat{U}_n\|_{L^2} + \|\theta - \hat{\theta}\|). \end{aligned} \quad (D.14)$$

From Grecksch and Kloeden (1996, Section 3), for every θ , we have

$$\mathbb{E}[\|U(\mathbf{x}, k\Delta t; \theta) - U_n(\mathbf{x}, k\Delta t; \theta)\|^2] \leq K_{k\Delta t, U(\cdot, 0), \theta} \lambda_{N+1}^{-1}(\theta), \quad (D.15)$$

where $K_{k\Delta t, U(\cdot, 0), \theta}$ is a constant depending on time $k\Delta t$ and the initial condition $U(\cdot, 0)$, and θ . Without loss of generality, we assume that $K_{k\Delta t, U(\cdot, 0), \theta}$ is non-decreasing in k (otherwise we can replace $K_{k\Delta t, U(\cdot, 0), \theta}$ with $\tilde{K}_{k\Delta t, U(\cdot, 0), \theta} := \max_{1 \leq i \leq k} K_{i\Delta t, U(\cdot, 0), \theta}$). Given the initial condition $U(\mathbf{x}, 0)$ and $P_n U(\mathbf{x}, 0)$, we denote the probability measures of $U(\mathbf{x}, k\Delta t; \theta)$ and $U_n(\mathbf{x}, T; \theta)$ by $\nu_{U(\cdot, 0)}(k\Delta t)$ and $\nu_{n, U_n(\cdot, 0)}(k\Delta t)$, respectively. Moreover, the joint probability measure of $(U(\mathbf{x}, k\Delta t; \theta), U_n(\mathbf{x}, k\Delta t; \theta))$ has marginal distributions $\nu_{U(\cdot, 0)}(k\Delta t)$ and $\nu_{n, U_n(\cdot, 0)}(k\Delta t)$, respectively. From Eq. (D.15), we can deduce that:

$$\begin{aligned} W_2^2(\nu_{U(\cdot, 0)}(k\Delta t), \nu_{n, U_n(\cdot, 0)}(k\Delta t)) &\leq \mathbb{E}[\|U(\mathbf{x}, k\Delta t; \theta) - U_n(\mathbf{x}, k\Delta t; \theta)\|^2] \\ &\leq \sup_{\theta, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \theta} \lambda_{N+1}^{-1}(\theta). \end{aligned} \quad (D.16)$$

Furthermore, using the definition of the local squared W_2 distance in Eq. (2.4), we have:

$$\begin{aligned} W_{2, \delta}^{2, e}(U(\cdot, k\Delta t; \theta), U_n(\cdot, k\Delta t; \theta)) &\leq \sup_{\theta, U(\mathbf{x}, 0)} \mathbb{E}[\|U(\mathbf{x}, k\Delta t; \theta) - U_n(\mathbf{x}, k\Delta t; \theta)\|^2] \\ &\leq \sup_{\theta, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \theta} \lambda_{N+1}^{-1}(\theta). \end{aligned} \quad (D.17)$$

Similarly, we can conclude that:

$$W_{2, \delta}^{2, e}(\hat{U}(\cdot, k\Delta t; \hat{\theta}), \hat{U}_n(\cdot, k\Delta t; \hat{\theta})) \leq \sup_{\hat{\theta}, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \hat{\theta}} \lambda_{N+1}^{-1}(\hat{\theta}). \quad (D.18)$$

Given the same initial condition $U(\mathbf{x}, 0) = \hat{U}(\mathbf{x}, 0)$, using the triangle inequality of the Wasserstein distance (Clement & Desch, 2008), we have

$$\begin{aligned} W_{2, \delta}^{2, e}(U(\mathbf{x}, t; \theta), \hat{U}(\mathbf{x}, t; \hat{\theta})) &\leq 3W_{2, \delta}^{2, e}(U(\cdot, k\Delta t; \theta), U_n(\cdot, k\Delta t; \theta)) \\ &\quad + 3W_{2, \delta}^{2, e}(\hat{U}(\cdot, k\Delta t; \hat{\theta}), \hat{U}_n(\cdot, k\Delta t; \hat{\theta})) + 3W_{2, \delta}^{2, e}(U_n(\mathbf{x}, t; \theta), \hat{U}_n(\mathbf{x}, t; \hat{\theta})) \\ &\leq 3W_{2, \delta}^{2, e}(U_n(\mathbf{x}, t; \theta), \hat{U}_n(\mathbf{x}, t; \hat{\theta})) + 3 \sup_{\theta, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \theta} \lambda_{N+1}^{-1}(\theta) \\ &\quad + 3 \sup_{\hat{\theta}, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \hat{\theta}} \lambda_{N+1}^{-1}(\hat{\theta}). \end{aligned} \quad (D.19)$$

Integrating both sides of the inequality (D.19) over time, we have:

$$\begin{aligned} \tilde{W}_{2, \delta}^{2, e}(U, \hat{U}) &\leq 3\tilde{W}_{2, \delta}^{2, e}(U_n, \hat{U}_n) + 3T \sup_{\theta, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \theta} \lambda_{N+1}^{-1}(\theta) \\ &\quad + 3T \sup_{\hat{\theta}, U(\mathbf{x}, 0)} K_{T, U(\cdot, 0), \hat{\theta}} \lambda_{N+1}^{-1}(\hat{\theta}). \end{aligned} \quad (D.20)$$

Let

$$U_n(t; \theta) := (u_1(t; \theta), \dots, u_n(t; \theta)), \quad \hat{U}_n(t; \hat{\theta}) := (\hat{u}_1(t; \hat{\theta}), \dots, \hat{u}_n(t; \hat{\theta})) \quad (D.21)$$

be two vectors of the spectral expansion coefficients of $U_n(\mathbf{x}, t; \theta)$ and $\hat{U}_n(\mathbf{x}, t; \hat{\theta})$ in Eq. (D.5). We have $\|U_n(\mathbf{x}, t; \theta)\|_{L^2} = \|U_n(t)\|$ and $\|\hat{U}_n(\mathbf{x}, t; \hat{\theta})\|_{L^2} = \|\hat{U}_n(t)\|$ because $\|\varphi_i\|_{L^2} = 1, i = 1, \dots, n$. Furthermore, U_n and $\hat{U}_n$ are solutions to the two SDEs:

$$\begin{aligned} dU_n &= (A_n(U_n, t; \theta) + F_n(U_n, t; \theta))dt + G_n(U_n, t; \theta)dB_t, \\ d\hat{U}_n &= (A_n(\hat{U}_n, t; \hat{\theta}) + F_n(\hat{U}_n, t; \hat{\theta}))dt + G_n(\hat{U}_n, t; \hat{\theta})dB_t \end{aligned} \quad (D.22)$$

where $A_n(U_n, t; \theta)$, $F_n(U_n, t; \theta)$ and $G_n(U_n, t; \theta)$ are the n -dimensional vector of the coefficients in the spectral expansions of $A_n(U_n, \theta)$, $f_n(U_n, \theta)$

and $g_n(U_n, \theta)$ in Eqs. (D.13) and (D.14), respectively. A_n , F_n and G_n are also Lipschitz continuous in U_n and θ from Eqs. (D.13), (D.14). Applying Eq. (2.18) in Theorem 2.2, we obtain:

$$\begin{aligned} \mathbb{E}[\tilde{W}_{2, \delta}^{2, e}(U_n, \hat{U}_n)] &= \mathbb{E}[\tilde{W}_{2, \delta}^{2, e}(U_n, \hat{U}_n)] \leq 8C_0(\beta_n, n)T\delta^2 \exp(C_0(\beta_n; n)T) \\ &\quad + \frac{6C_1(\beta_n)}{C_0(\beta_n; n)}T \exp(C_0(\beta_n; n)T)(W_2^2(\mu, \hat{\mu}) + 2C_2\mathbb{E}[h(N^\#(U_n(\mathbf{x}, 0); \delta), \ell)]) \\ &\quad \cdot (\Theta_6^{\frac{1}{3}} + \hat{\Theta}_6^{\frac{1}{3}})) \end{aligned} \quad (D.23)$$

where $C_i, i = 0, \dots, 2$ are constants defined in Theorem (2.2). Taking the expectation of Eq. (D.20) and plugging in the inequality (D.23) for $\mathbb{E}[\tilde{W}_{2, \delta}^{2, e}(U_n, \hat{U}_n)]$, the inequality (D.8) is proved. $\square$

Appendix E. Proof of Theorem 3.1

Given a parameterized multivariate normal distribution f_x , we design an SNN described in Fig. 1 with the ReLU activation and the linear forward propagation. The probability density function of the output of this SNN, denoted by $\hat{f}_x$, can approximate f_x in the W_2 distance sense. Given a real number $0 < c < \epsilon_0$, we choose $\Delta x > 0$ such that:

$$W_2^2(f_x, \hat{f}_x) < c, \quad \forall \mathbf{x}, \tilde{\mathbf{x}} \in D, \quad \|\mathbf{x} - \tilde{\mathbf{x}}\| < \sqrt{d}\Delta x. \quad (E.1)$$

We consider a uniform equidistance grid set $X := \{\mathbf{x}_i\}_{i=1}^K, \mathbf{x}_i = (x_i^1, \dots, x_i^d)$ such that the distance between two adjacent points is Δx , $D \subseteq \cup_{i=1}^K \otimes_{j=1}^d [x_i^j, x_i^j + \Delta x]$, and $\otimes_{j=1}^d [x_i^j, x_i^j + \Delta x] \cap \otimes_{j=1}^d [x_{i_2}^j, x_{i_2}^j + \Delta x] = \emptyset$ if $i_1 \neq i_2$. Therefore, $\forall \mathbf{x} = (x^1, \dots, x^d) \in D$, there exists $\mathbf{x}_i \in X$ such that $\|\mathbf{x} - \mathbf{x}_i\| < \sqrt{d}\Delta x$. Let $0 < \epsilon < \frac{1}{2}$ be a small positive number. We set $4dK$ neurons in the first layer, grouped into dK groups. When inputting $\mathbf{x}$ into the SNN, the outputs of the four neurons in the $(i, j), i = 1, \dots, K, j = 1, \dots, d$ group are:

$$\begin{aligned} n_{i,j,1}^1 &= \text{ReLU}(\epsilon^{-1}(x^j - x_i^j - \Delta x)), \quad n_{i,j,2}^1 = \text{ReLU}(\epsilon^{-1}(x^j - x_i^j - \Delta x + \epsilon)), \\ n_{i,j,3}^1 &= \text{ReLU}(\epsilon^{-1}(x^j - x_i^j - \epsilon)), \quad n_{i,j,4}^1 = \text{ReLU}(\epsilon^{-1}(x^j - x_i^j)). \end{aligned} \quad (E.2)$$

The second hidden layer contain dK neurons labeled by $(i, j), i = 1, \dots, K, j = 1, \dots, d$. We set weights between the first layer and the second layer such that the input of the (i, j) neuron in the second layer is:

$$n_{i,j}^{2, \text{in}} := n_{i,j,2}^1 - n_{i,j,1}^1 - (n_{i,j,4}^1 - n_{i,j,3}^1). \quad (E.3)$$

It is easy to verify that $n_{i,j}^{2, \text{in}} \in [0, 1]$. Furthermore, when $x^j \in [x_i^j + \epsilon, x_i^j + \Delta x]$, then $n_{i,j}^{2, \text{in}} = 1, j = 1, \dots, d$. If $x^j \leq x_i^j$ or $x^j \geq x_i^j + \Delta x$, then $n_{i,j}^{2, \text{in}} = 0$. The output of the (i, j) neuron in the second hidden layer is designed as

$$n_{i,j}^2 = \text{ReLU}(\epsilon^{-1}(n_{i,j}^{2, \text{in}} - 1 + \epsilon)), \quad i = 1, \dots, K, \quad j = 1, \dots, d. \quad (E.4)$$

The third layer contains K neurons, and the output of each neuron in the third hidden layer is:

$$n_i^3 = \text{ReLU}\left(\sum_{j=1}^d \epsilon^{-1}\left(n_{i,j}^2 - 1 + \frac{\epsilon}{d}\right)\right), \quad i = 1, \dots, K. \quad (E.5)$$

Thus, $n_i^3 \in [0, 1]$ and $n_i^3 = 1$ when $\mathbf{x} \in \otimes_j [x_{j,i} + \epsilon, x_{j,i} + \Delta x - \epsilon]$. If there is a $j = 1, \dots, d$ such that $x^j < x_i^j$ or $x^j > x_i^j + \Delta x$, then $n_i^3 = 0$. For $\mathbf{x} \in D$, there exists at most one i such that $n_i^3 \neq 0$. We denote $D(\epsilon) = \{x \in D : \exists! i, n_i^3 = 1\}$. It is easy to check that as $\epsilon \rightarrow 0$, $D(\epsilon) \rightarrow D$.

We set $d'K$ neurons in the fourth layer. Each neuron is labeled with $(i, j), i = 1, \dots, K, k = 1, \dots, d'$. The input and output of the (i, j) neuron in the fourth layer are identical (i.e., each neuron in the fourth layer outputs its input):

$$n_i^3(\omega_{i,k}^4 + (A_i^{-1}b_i)_k), \quad i = 1, \dots, K, \quad k = 1, \dots, d'. \quad (E.6)$$

Here, $\omega_{i,k} \sim \mathcal{N}(0, 1)$ are independent random variables, while $\mathbf{b}_i$ and A_i are the mean and covariance matrix of $f_{x_i}(\mathbf{y}) = \mathcal{N}(\mathbf{b}_i, A_i)$, respectively. $(A_i^{-1} \mathbf{b}_i)_k$ refers to the k th component of the vector $A_i^{-1} \mathbf{b}_i$. The output of the fifth layer is d' dimensional:

$$\sum_{i=1}^K n_i^3 (A_i \omega_i^4 + \mathbf{b}_i), \quad (\text{E.7})$$

where $\omega_i := (\omega_{i,1}, \dots, \omega_{i,d'})$. Because for any $\mathbf{x} \in D$, there exists at most one i such that $n_i^3 \neq 0$, for any $\mathbf{x} \in D$, we have:

$$\begin{aligned} \sup_{\hat{\mathbf{y}} \sim f_{\mathbf{x}}} \mathbb{E}[\|\mathbf{y}\|^2] &\leq \sup_i \mathbb{E}[\|A_i \omega_i^4 + \mathbf{b}_i\|^2] \\ &\leq \sup_i (\|A_i^T A_i\|_F^2 + \|\mathbf{b}_i\|_2^2). \end{aligned} \quad (\text{E.8})$$

For each $\mathbf{x} \in D(\epsilon)$, if $n_i^3 = 1$ i.e., $\|\mathbf{x} - \mathbf{x}_i\| < \Delta x$, then $\hat{f}_{\mathbf{x}}$ is the probability density function of a multivariate normal distribution, and $\hat{f}_{\mathbf{x}}$ is identical to f_{x_i} . Therefore, we have

$$\begin{aligned} \int_D W_2^2(f_{\mathbf{x}}, \hat{f}_{\mathbf{x}}) \gamma(\mathbf{d}\mathbf{x}) &< \int_{D(\epsilon)} W_2^2(f_{\mathbf{x}}, \hat{f}_{\mathbf{x}}) \gamma(\mathbf{d}\mathbf{x}) \\ &+ 2\gamma(D - D(\epsilon)) \cdot \left(\sup_{\mathbf{y} \sim f_{\mathbf{x}}} \mathbb{E}[\|\mathbf{y}\|^2] + \sup_{\hat{\mathbf{y}} \sim f_{\mathbf{x}}} \mathbb{E}[\|\hat{\mathbf{y}}\|^2] \right) \\ &\leq c\gamma(D) + \gamma(D - D(\epsilon)) 2 \left(\sup_{\hat{\mathbf{y}} \sim f_{\mathbf{x}}} \mathbb{E}[\|\mathbf{y}\|^2] + \sup_i (\|A_i^T A_i\|_F^2 + \|\mathbf{b}_i\|_2^2) \right) \\ &\leq c + 4\gamma(D - D(\epsilon))Y. \end{aligned} \quad (\text{E.9})$$

Choosing $c < \epsilon_0$ and a small ϵ such that $\gamma(D - D(\epsilon)) \leq \frac{\epsilon_0 - c}{4Y}$, we have proved [Theorem 3.1](#).

Appendix F. Proof of [Corollary 3.1](#)

The proof of [Corollary 3.1](#) is based on the proof of [Theorem 3.1](#). Given any positive number $c > 0$, we can find a $\Delta x > 0$ such that:

$$W_2^2(f_{\mathbf{x}}, f_{\tilde{\mathbf{x}}}) < c, \quad \forall \|\mathbf{x} - \tilde{\mathbf{x}}\| < \sqrt{d}\Delta x, \quad \mathbf{x}, \tilde{\mathbf{x}} \in D. \quad (\text{F.1})$$

We establish an equidistance grid point set $X := \{\mathbf{x}_i\}_{i=1}^K$ on D such that the distance between two adjacent points is Δx , $D \subseteq \cup_{i=1}^K \otimes_{j=1}^d [x_{i,j}^j, x_{i,j}^j + \Delta x]$ and $\otimes_{j=1}^d [x_{i_1,j}^j, x_{i_1,j}^j + \Delta x] \cap \otimes_{j=1}^d [x_{i_2,j}^j, x_{i_2,j}^j + \Delta x] = \emptyset$ if $i_1 \neq i_2$. Thus, $\forall \mathbf{x} \in D$, there exists $\mathbf{x}_i \in X$ such that $\|\mathbf{x} - \mathbf{x}_i\| \leq \sqrt{d}\Delta x$. Denote $\Phi(x)$ to be the cumulative distribution function of a standard normal random variable. Suppose $-M = h_{i,0} < h_{i,1} < \dots < h_{i,s} = M$ such that:

$$\begin{aligned} \Phi(h_{i,r+1}) - \Phi(h_{i,r}) &= p_{i,r+1}, \quad r = 1, \dots, s-2, \quad \Phi(h_{i,1}) = p_{i,1}, \\ \Phi(h_{i,s-1}) &= 1 - p_{i,s}, \end{aligned} \quad (\text{F.2})$$

where $p_{i,r} := p_r(\mathbf{x}_i)$. We design an SNN of the form in [Fig. 1](#) that satisfies [Eq. \(3.12\)](#). We set the first three hidden layers of this SNN to be the same as the first three layers of the SNN designed in the proof of [Theorem 3.1](#) in [Appendix E](#). We also refer to the outputs of the third hidden layer as $n_i^3, i = 1, \dots, K$. For $\mathbf{x}_i \in X$, we denote $\mathbf{x}_i = (x_i^1, \dots, x_i^d)$. From [Appendix E](#), $n_i^3 \in [0, 1]$, $n_i^3 = 1$ when $\mathbf{x} = (x^1, \dots, x^d) \in \otimes_{j=1}^d [x_i^j - \epsilon, x_i^j + \Delta x - \epsilon]$, and $n_i^3 = 0$ when $\mathbf{x} \in D - \otimes_{j=1}^d [x_i^j, x_i^j + \Delta x]$. $\epsilon > 0$ is a small number to be determined. If there is a $j = 1, \dots, d$ such that $x^j < x_i^j$ or $x^j > x_i^j + \Delta x$, then $n_i^3 = 0$. For $\mathbf{x} \in D$, there is at most one i such that $n_i^3 = 1$.

The fourth layer contains $K(s+1)$ groups of neurons with each group containing 2 neurons. The outputs of the $(i, r), i = 1, \dots, K, r = 0, \dots, s$ group are:

$$n_{i,r}^4 = \text{ReLU}(\tilde{w}_i^3 n_i^3 - h_{i,r} - M_0), \quad \tilde{n}_{i,r}^4 = \text{ReLU}(\tilde{w}_i^3 n_i^3 - h_{i,r} - \epsilon_0 - M_0). \quad (\text{F.3})$$

$\tilde{w}_i^3 \sim \mathcal{N}(M_0, 1)$ are independent random variables. $M_0 > |h_{i,r}|, \forall i, r$ is a large number and $\epsilon_0 > 0$ is a small number satisfying $\epsilon_0 < \min_{i,r} \frac{h_{i,r+1} - h_{i,r}}{2}$. Both M_0 and ϵ_0 are to be determined.

The fifth layer contains sK neurons and the output of the (i, r) neuron is:

$$n_{i,r+1}^5 = \text{ReLU}\left(\epsilon_0^{-1} \left(n_{i,r}^4 - \tilde{n}_{i,r}^4 - \left(n_{i,r+1}^4 - \tilde{n}_{i,r+1}^4 \right) \right)\right), \quad r = 0, \dots, s-1. \quad (\text{F.4})$$

Therefore, when $n_i^3 = 1$, for $r = 0, \dots, s-1$, we have:

$$\begin{aligned} n_{i,r+1}^5 &= 1, \omega_i^3 \in [M_0 + h_{i,r} + \epsilon_0, M_0 + h_{i,r+1}], \\ n_{i,r+1}^5 &= 0, \omega_i^3 < M_0 + h_{i,r} \text{ or } \omega_i^3 > M_0 + h_{i,r+1} + \epsilon_0, \\ n_{i,r+1}^5 &\in (0, 1), \text{ otherwise}; \end{aligned} \quad (\text{F.5})$$

Furthermore, we can see that for any $n_i^3 \in [0, 1]$

$$0 \leq \sum_{r=1}^s n_{i,r}^5 \leq 1, \quad (\text{F.6})$$

and there exist at most two non-zero $n_{i,s}^5, n_{i+1,s}^5 > 0$ in $(n_{i,1}^5, \dots, n_{i,s}^5)$.

For any $\epsilon_1 > 0$, there exist a small $\epsilon_0 > 0$ and a large $M_0 > 0$ such that when $n_i^3 = 1$:

$$0 \leq p_{i,r} - \hat{p}_{i,r} \leq \frac{\epsilon_1}{s}, \quad \hat{p}_{i,r} := p(n_{i,r}^5 = 1), \quad i = 1, \dots, K. \quad (\text{F.7})$$

The sixth hidden layer contains sKd' neurons, each of whose input and output are both

$$n_{i,r,k}^6 = n_{i,r}^5 (w_{i,r,k} + (A_{i,r}^{-1} \mathbf{b}_{i,r})_k), \quad i = 1, \dots, K, r = 1, \dots, s, k = 1, \dots, d'. \quad (\text{F.8})$$

Here, $w_{i,r,k} \sim \mathcal{N}(0, 1)$ are independent random variables. $\mathbf{b}_{i,s}$ and $A_{i,s}$, respectively, are the mean vector and covariance matrix such that

$$f_{x_i}(\mathbf{y}) = \sum_{r=1}^s p_r(\mathbf{x}_i) \mathcal{N}(\mathbf{b}_{i,r}, A_{i,r}^T A_{i,r}). \quad (\text{F.9})$$

The seventh layer contains d' neurons whose output is:

$$\sum_{i=1}^K \sum_{r=1}^s n_{i,r}^5 (A_{i,r} \mathbf{w}_{i,r} + \mathbf{b}_{i,r}), \quad (\text{F.10})$$

where $\mathbf{w}_{i,r} := (w_{i,r,1}, \dots, w_{i,r,d'})$. When $\mathbf{x} \in D(\epsilon) := \{\mathbf{x} \in D : \exists i, n_i^3 = 1\}$, if $n_i^3 = 1$, then the probability density function of the output of the seventh layer can be written as:

$$\sum_{r=1}^s \hat{p}_{i,r} \mathcal{N}(\mathbf{b}_{i,r}, A_{i,r}^T A_{i,r}) + p_i(\mathbf{y}), \quad \int_{\mathbb{R}^{d'}} p_i(\mathbf{y}) \mathbf{d}\mathbf{y} := 1 - \sum_{r=1}^s \hat{p}_{i,r} \leq \epsilon_1, \quad p_i(\mathbf{y}) \geq 0. \quad (\text{F.11})$$

Furthermore, since there exist at most two non-zero consecutive $n_{i,s}^5, n_{i+1,s}^5 > 0$, we have:

$$\begin{aligned} \int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 p_i(\mathbf{y}) \mathbf{d}\mathbf{y} &\leq \int_{\mathbb{R}^{d'}} p_i(\mathbf{y}) \mathbf{d}\mathbf{y} \cdot \left(2 \max_{i,r} \mathbb{E}_{\mathbf{y} \sim \mathcal{N}(\mathbf{b}_{i,r}, A_{i,r}^T A_{i,r})} [\|\mathbf{y}\|^2] \right. \\ &\quad \left. + 2 \max_{i,r} \mathbb{E}_{\mathbf{y} \sim \mathcal{N}(\mathbf{b}_{i,r}, A_{i,r}^T A_{i,r})} [\|\mathbf{y}\|^2] \right) \\ &= (1 - \sum_{r=1}^s \hat{p}_{i,r}) 4 \max_r \left(\|\mathbf{b}_{i,r}\|^2 + \|A_{i,r}^T A_{i,r}\|_F^2 \right) \\ &\leq 4\epsilon_1 \max_r \left(\|\mathbf{b}_{i,r}\|^2 + \|A_{i,r}^T A_{i,r}\|_F^2 \right). \end{aligned} \quad (\text{F.12})$$

Applying [Lemma 3.1](#), we have

$$\begin{aligned} W_2^2(\hat{f}_{\mathbf{x}}, f_{x_i}) &\leq 2 \max_r \frac{\epsilon_1}{s} \cdot s \left(\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2 \right) \\ &\quad + 4\epsilon_1 \max_r \left(\|\mathbf{b}_{i,r}\|^2 + \|A_{i,r}^T A_{i,r}\|_F^2 \right), \end{aligned} \quad (\text{F.13})$$

when $\mathbf{x} \in D^i(\epsilon) := \{\mathbf{x} \in D(\epsilon) \mid \|\mathbf{x}_i - \mathbf{x}\| \leq \|\mathbf{x}_j - \mathbf{x}\|, \forall j \neq i\}$. Finally, we have:

$$\begin{aligned} & \int_D W_2^2(f_{\mathbf{x}}, \hat{f}_{\mathbf{x}}) \gamma(d\mathbf{x}) \\ & \leq \int_{D(\epsilon)} W_2^2(f_{\mathbf{x}}, \hat{f}_{\mathbf{x}}) \gamma(d\mathbf{x}) \\ & \quad + 2(1 - \gamma(D(\epsilon))) \left(\mathbb{E}[\|\mathbf{y}\|^2] + 4 \sup_{i,r} (\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2) \right) \\ & \leq 2 \sum_i \int_{D^i(\epsilon)} W_2^2(\hat{f}_{\mathbf{x}}, f_{\mathbf{x}_i}) \gamma(d\mathbf{x}) + 2 \sum_i \int_{D^i(\epsilon)} c \gamma(d\mathbf{x}) \quad (\text{F.14}) \\ & \quad + 2(1 - \gamma(D(\epsilon))) \left(\mathbb{E}[\|\mathbf{y}\|^2] + 4 \sup_{i,r} (\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2) \right) \\ & \leq 2 \sup_{i,r} \left(6\epsilon_1 (\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2) \right) + 2c \\ & \quad + 2(1 - \gamma(D(\epsilon))) \left(\mathbb{E}[\|\mathbf{y}\|^2] + 4 \sup_{i,r} (\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2) \right). \end{aligned}$$

Note that $\sup_{i,r} (\|A_{i,r}^T A_{i,r}\|_F^2 + \|\mathbf{b}_{i,r}\|^2)$ is uniformly bounded from the assumption Eq. (3.11). Letting $c, \epsilon, \epsilon_1 \rightarrow 0^+$ in Eq. (F.14), we have proved the inequality (3.12).

Appendix G. Proof of Theorem 3.2

First, we define an auxiliary function with an additional parameter σ :

$$f_{\sigma^2}(\mathbf{y}) := \int_{\mathbb{R}^{d'}} f(\mathbf{y}') \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}'. \quad (\text{G.1})$$

Since f is uniformly continuous and uniformly bounded, $\forall \epsilon_0 > 0$, there exist a small $\delta > 0$ and a small $\sigma_0 > 0$ such that for any $\sigma < \sigma_0$, we have i) $|f(\mathbf{y}) - f(\tilde{\mathbf{y}})| < \epsilon_0, \forall \|\tilde{\mathbf{y}} - \mathbf{y}\| < \delta$ and ii)

$$\int_{B(0, \delta)} \mathcal{N}(\mathbf{y}; \sigma^2 I_{d' \times d'}) d\mathbf{y} > 1 - \epsilon_0. \quad (\text{G.2})$$

Therefore, we conclude that $\lim_{\sigma \rightarrow 0} f_{\sigma^2}(\mathbf{y}) = f(\mathbf{y})$ uniformly on $\mathbb{R}^{d'}$.

Letting $\{\mathbf{y}_j\}_{j=1}^{(n_0+1)^{d'}}$ and $\{w_j\}_{j=1}^{(n_0+1)^{d'}}$ be the multidimensional Hermite collocation points and weights on $\mathbb{R}^{d'}$ described in Shen et al. (2011), we have

$$\begin{aligned} & \int_{\mathbb{R}^{d'}} I_{n_0} f(\mathbf{y}') \cdot I_{n_0} \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}' \\ & = \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}) w_i. \end{aligned} \quad (\text{G.3})$$

Here, I_{n_0} is the interpolation operator such that

$$f(\mathbf{y}_j) = I_{n_0} f(\mathbf{y}_j) \in P_{n_0}, \quad j = 1, \dots, (n_0 + 1)^{d'}, \quad (\text{G.4})$$

where P_{n_0} is the space spanned by the generalized Hermite functions $\hat{H}_n(\mathbf{y})$ such that $|\mathbf{n}|_{\infty} \leq n_0$ and

$$\hat{H}_n(\mathbf{y}) := \prod_{i=1}^{d'} \hat{H}_{n_i}(y_i), \quad \mathbf{n} = (n_1, \dots, n_{d'}), \quad \mathbf{y} = (y_1, \dots, y_{d'}) \quad (\text{G.5})$$

is the multidimensional generalized Hermite function defined in Shen et al. (2011) ($\hat{H}_{n_i}$ is the 1D generalized Hermite function of order n_i). We denote

$$\begin{aligned} f_{\sigma^2, n_0}(\mathbf{y}) & := \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) \cdot \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}) w_i \\ & = \int_{\mathbb{R}^{d'}} I_{n_0} f(\mathbf{y}') \cdot I_{n_0} \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}', \end{aligned} \quad (\text{G.6})$$

where f_{σ^2, n_0} is nonnegative because the collocation weights $w_j > 0$. Furthermore,

$$\left| \int_{\mathbb{R}^{d'}} f(\mathbf{y}') \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) - I_{n_0} f(\mathbf{y}') I_{n_0} \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}' \right|$$

$$\begin{aligned} & \leq \|f - I_{n_0} f\|_{L^2} \cdot \|\mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'})\|_{L^2} \\ & \quad + \|I_{n_0} f\|_{L^2} \cdot \|\mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) - I_{n_0} \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'})\|_{L^2}. \end{aligned} \quad (\text{G.7})$$

Using (Shen et al., 2011, Theorem 7.18, Theorem 8.6), we have:

$$\begin{aligned} \|f - I_{n_0} f\| & \leq \|f - I_{n_0}^1 f\|_{L^2} + \|I_{n_0}^2 \circ \dots \circ I_{n_0}^{d'} f - f\|_{L^2} \\ & \quad + \|(I_{n_0}^1 - \mathbb{I}) \circ (I_{n_0}^2 \circ \dots \circ I_{n_0}^{d'} f - f)\|_{L^2}, \\ & \leq C n_0^{-\frac{1}{3}} \|\partial_{y_1} f\|_{L^2} + \|I_{n_0}^2 \circ \dots \circ I_{n_0}^{d'} f - f\|_{L^2} \\ & \quad + \|I_{n_0}^2 \circ \dots \circ (I_{n_0}^1 - \mathbb{I})(I_{n_0}^{d'} f - f)\|_{L^2}, \\ & \leq C n_0^{-\frac{1}{3}} \|\partial_{y_1} f\|_{L^2} + \|I_{n_0}^2 \circ \dots \circ I_{n_0}^{d'} f - f\|_{L^2} \\ & \quad + C n_0^{-\frac{1}{3}} \|(I_{n_0}^2 \circ \dots \circ I_{n_0}^{d'} \partial_{y_1} f - \partial_{y_1} f)\|_{L^2} \\ & \leq \dots \\ & \leq C n_0^{-\frac{1}{3}} |f|_{\text{mix}} \end{aligned} \quad (\text{G.8})$$

Here, C is a constant and n_0 is taken to be large enough such that $C n_0^{-\frac{1}{3}} < 1$. $\mathbb{I}$ is the identity operator and $I_{n_0}^i, i = 1, \dots, d'$ is the projection operator in the i th direction, i.e., if we denote $X_{n_0} := \{y_i\}_{i=0}^{n_0}$ to be the 1D Hermite collocation points, then

$$I_{n_0}^i f(\mathbf{y}) = f(\mathbf{y}), \quad \forall \mathbf{y} = (y_1, \dots, y_{d'}) \text{ if } y_i \in X_{n_0}. \quad (\text{G.9})$$

Similarly, for any fixed $\mathbf{y} \in \mathbb{R}^{d'}$:

$$\begin{aligned} & \|\mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) - I_{n_0} \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'})\|_{L^2} \\ & \leq \sum_{|\mathbf{n}|_0 \leq n_0} C n_0^{-\frac{1}{3}} \|\partial_{\mathbf{n}} \mathcal{N}\|_{L^2}. \end{aligned} \quad (\text{G.10})$$

Combining Eqs. (G.3), (G.7), (G.8) and (G.10), for every $\mathbf{y}$ and every $\sigma > 0$, as $n_0 \rightarrow \infty$,

$$\begin{aligned} & \left| \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}) w_i - \int_{\mathbb{R}^{d'}} f(\mathbf{y}') \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}' \right| \\ & \leq C n_0^{-\frac{1}{3}} \left(|f|_{\text{mix}} \|\mathcal{N}(\mathbf{y}, \sigma^2 I_{d' \times d'})\|_{L^2} + (\|f\|_{L^2} + C n_0^{-\frac{1}{3}} |f|_{\text{mix}}) \right) \\ & \quad \cdot |\mathcal{N}(\mathbf{y}, \sigma^2 I_{d' \times d'})|_{\text{mix}}, \end{aligned} \quad (\text{G.11})$$

which implies that for any fixed $\sigma > 0$,

$$\begin{aligned} f_{\sigma^2, n_0}(\mathbf{y}) & = \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}) w_i \\ & \rightarrow \int_{\mathbb{R}^{d'}} f(\mathbf{y}') \mathcal{N}(\mathbf{y} - \mathbf{y}'; \sigma^2 I_{d' \times d'}) d\mathbf{y}' \end{aligned} \quad (\text{G.12})$$

uniformly in $\mathbf{y} \in \mathbb{R}^{d'}$ as $n_0 \rightarrow \infty$.

Additionally, we have:

$$\begin{aligned} & \int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) w_i \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}) d\mathbf{y} \\ & = \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) (|\mathbf{y}_i|^2 + d' \sigma^2) w_i \\ & = \sum_{i=1}^{(n_0+1)^{d'}} f(\mathbf{y}_i) (|\mathbf{y}_i|^2 + d' \sigma^2)^2 \frac{1}{|\mathbf{y}_i|^2 + d' \sigma^2} w_i \\ & = \int_{\mathbb{R}^{d'}} I_{n_0} (f(\mathbf{y}) (\|\mathbf{y}\|^2 + d' \sigma^2)^2) \cdot I_{n_0} \left(\frac{1}{(\|\mathbf{y}\|^2 + d' \sigma^2)} \right) d\mathbf{y} \\ & \leq \int_{\mathbb{R}^{d'}} f(\mathbf{y}) (\|\mathbf{y}\|^2 + d' \sigma^2)^2 \cdot \frac{1}{(\|\mathbf{y}\|^2 + d' \sigma^2)} d\mathbf{y} \\ & \quad + \|(I_{n_0} - \mathbb{I})(f(\mathbf{y}) (\|\mathbf{y}\|^2 + d' \sigma^2)^2)\|_{L^2} \cdot \left\| \frac{1}{\|\mathbf{y}\|^2 + d' \sigma^2} \right\|_{L^2} \end{aligned}$$

$$\begin{aligned}
& + \|I_{n_0}(f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2)\|_{L^2} \cdot \|(\mathbb{I} - I_{n_0})\left(\frac{1}{\|\mathbf{y}\|^2 + d'\sigma^2}\right)\|_{L^2} \\
& \leq \left(\sum_{i,j=1}^{d'} 2Cn_0^{-\frac{1}{3}} |f(\mathbf{y})y_i^2 y_j^2|_{\text{mix}} + 2Cn_0^{-\frac{1}{3}} d'\sigma^2 \sum_{i=1}^{d'} |f(\mathbf{y})y_i^2|_{\text{mix}}\right. \\
& \quad + Cn_0^{-\frac{1}{3}} (d')^2 \sigma^4 |f|_{\text{mix}} \cdot \sigma^{-\frac{1}{2}} C_1(d') \\
& \quad \left. + \|I_{n_0}(f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2)\|_{L^2} \cdot Cn_0^{-\frac{1}{3}} \left|\frac{1}{\|\mathbf{y}\|^2 + d'\sigma^2}\right|_{\text{mix}} + \mathbb{E}[\|\mathbf{y}\|^2] + d'\sigma^2\right),
\end{aligned} \tag{G.13}$$

where $C_1(d')$ is another constant depending on the dimensionality d' . Since

$$\begin{aligned}
& \|I_{n_0}(f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2) - f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2\|_{L^2} \\
& \leq \sum_{i,j=1}^{d'} 2Cn_0^{-\frac{1}{3}} |f(\mathbf{y})y_i^2 y_j^2|_{\text{mix}} + 2Cn_0^{-\frac{1}{3}} d'\sigma^2 \sum_{i=1}^{d'} |f(\mathbf{y})y_i^2|_{\text{mix}} \\
& \quad + Cn_0^{-\frac{1}{3}} (d')^2 \sigma^4 |f|_{\text{mix}} \leq \infty,
\end{aligned} \tag{G.14}$$

we conclude that

$$\int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 \sum_{i=1}^{(n_0+1)d'} f(\mathbf{y})w_i \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2) d\mathbf{y} < \infty. \tag{G.15}$$

Furthermore, from the inequality (G.14), $\|I_{n_0}(f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2)\|_{L^2} \rightarrow \|f(\mathbf{y})(\|\mathbf{y}\|^2 + d'\sigma^2)^2\|_{L^2}$ as $n_0 \rightarrow \infty$. We denote

$$\tilde{f}_{\sigma^2, n_0}(\mathbf{y}) := \frac{1}{\sum_{j=1}^{(n_0+1)d'} f(\mathbf{y}_j)w_j} f_{\sigma^2, n_0}(\mathbf{y}). \tag{G.16}$$

Therefore, $\int_{\mathbb{R}^{d'}} \tilde{f}_{\sigma^2, n_0}(\mathbf{y}) d\mathbf{y} = 1$. Note that:

$$\begin{aligned}
& \left| \int_{\mathbb{R}^{d'}} f(\mathbf{y}) d\mathbf{y} - \sum_{i=1}^{(n_0+1)d'} f(\mathbf{y}_i)w_i \right| = \left| \int_{\mathbb{R}^{d'}} f - I_{n_0} \sqrt{f} \cdot I_{n_0} \sqrt{f} d\mathbf{y} \right| \\
& \leq \|\sqrt{f}\|_{L^2} Cn_0^{-\frac{1}{3}} \|\sqrt{f}\|_{\text{mix}} + \|I_{n_0} \sqrt{f}\|_{L^2} Cn_0^{-\frac{1}{3}} \|\sqrt{f}\|_{\text{mix}}, \\
& \leq Cn_0^{-\frac{1}{3}} \|\sqrt{f}\|_{\text{mix}} (2\|\sqrt{f}\|_{L^2} + Cn_0^{-\frac{1}{3}} \|\sqrt{f}\|_{\text{mix}}).
\end{aligned} \tag{G.17}$$

Therefore, we can write

$$\sum_{i=1}^{n_0^{d'}} f(\mathbf{y}_i)w_i := 1 + n_0^{-\frac{1}{3}} c(\sqrt{f}) \leq \infty, \tag{G.18}$$

where $c(\sqrt{f})$ is a constant depending on $\|\sqrt{f}\|_{\text{mix}}$ ($\|\sqrt{f}\|_{L^2} = 1$). We assume n_0 is large enough such that $n_0^{-\frac{1}{3}} c(\sqrt{f}) < \frac{1}{2}$. From the inequality (G.13) and the definition of $\tilde{f}_{\sigma^2, n_0}$ in Eq. (G.16), we can bound $\mathbb{E}_{\mathbf{y} \sim \tilde{f}_{\sigma^2, n_0}}[\|\mathbf{y}\|^2]$ by

$$\begin{aligned}
\mathbb{E}_{\mathbf{y} \sim \tilde{f}_{\sigma^2, n_0}}[\|\mathbf{y}\|^2] & \leq \frac{1}{1 - |c(f)|n_0^{-\frac{1}{3}}} \left[\mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + d'\sigma^2 \right. \\
& \quad \left. + n_0^{-\frac{1}{3}} (C_2(\sigma; f) + \mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + 1) \right] \\
& \leq \left[\mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + n_0^{-\frac{1}{3}} (C_2(\sigma; f) + \mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + 1) + d'\sigma^2 \right] \\
& \quad + 2|c(\sqrt{f})|n_0^{-\frac{1}{3}} \left[\mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + n_0^{-\frac{1}{3}} (C_2(\sigma; f) \right. \\
& \quad \left. + \mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + 1) + d'\sigma^2 \right] \\
& = \mathbb{E}[\|\mathbf{y}\|^2] + d'\sigma^2 + n_0^{-\frac{1}{3}} C_3(\sigma; f),
\end{aligned} \tag{G.19}$$

where $C_2(\sigma; f)$ is a constant that depends on f and σ , and $C_3(\sigma; f)$ is another constant depending on f , $\sqrt{f}$ and σ .

Consider the special coupling measure

$$\begin{aligned}
\pi(f, \tilde{f}_{\sigma^2, n_0})(\mathbf{y}, \hat{\mathbf{y}}) & := \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y})) \delta(\mathbf{y} - \hat{\mathbf{y}}) \\
& \quad + \frac{1}{A} (f(\mathbf{y}) - \min(f, \tilde{f}_{\sigma^2, n_0})(\mathbf{y})) \cdot (\tilde{f}_{\sigma^2, n_0}(\hat{\mathbf{y}}) - \min(f, \tilde{f}_{\sigma^2, n_0})(\hat{\mathbf{y}})), \\
& \text{if } \int_{\mathbb{R}^{d'}} \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y})) d\mathbf{y} < 1, \\
\pi(f, \tilde{f}_{\sigma^2, n_0})(\mathbf{y}, \hat{\mathbf{y}}) & := f(\mathbf{y}) \delta(\mathbf{y} - \hat{\mathbf{y}}), \text{ if } \int_{\mathbb{R}^{d'}} \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y})) d\mathbf{y} = 1,
\end{aligned} \tag{G.20}$$

where $A := \int_{\mathbb{R}^{d'}} \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y})) d\mathbf{y}$ and δ is the Dirac delta measure. The marginal probability densities of $\pi(f, \tilde{f}_{\sigma^2, n_0})$ are $f(\mathbf{y})$ and $\tilde{f}_{\sigma^2, n_0}(\mathbf{y})$, respectively. Furthermore, we have:

$$\begin{aligned}
& \mathbb{E}_{(\mathbf{y}, \hat{\mathbf{y}}) \sim \pi(f, \tilde{f}_{\sigma^2, n_0})}[\|\mathbf{y} - \hat{\mathbf{y}}\|^2] \\
& \leq 2 \int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 (f(\mathbf{y}) - \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y}))) d\mathbf{y} \\
& \quad + 2 \int_{\mathbb{R}^{d'}} \|\hat{\mathbf{y}}\|^2 (\tilde{f}_{\sigma^2, n_0}(\mathbf{y}) - \min(f(\mathbf{y}), \tilde{f}_{\sigma^2, n_0}(\mathbf{y}))) d\mathbf{y} \\
& \leq 4 \int_{\mathbb{R}^{d'}} \|\mathbf{y}\|^2 |f(\mathbf{y}) - \tilde{f}_{\sigma^2, n_0}(\mathbf{y})| d\mathbf{y}.
\end{aligned} \tag{G.21}$$

Fixing $\sigma > 0$, from Eq. (G.12), $f_{\sigma^2, n_0} \rightarrow f_{\sigma^2}$ uniformly as $n_0 \rightarrow \infty$; furthermore, from the definition of $\tilde{f}_{\sigma^2, n_0}$ in Eq. (G.16), $\tilde{f}_{\sigma^2, n_0} \rightarrow f_{\sigma^2, n_0}$ uniformly as $n_0 \rightarrow \infty$. Finally, since $\lim_{\sigma \rightarrow 0} f_{\sigma^2}(\mathbf{y}) = f(\mathbf{y})$ uniformly as $\sigma \rightarrow 0$ for any $\mathbf{y} \in \mathbb{R}^{d'}$, we conclude that

$$\tilde{f}_{\sigma^2, n_0(\sigma)} \rightarrow f \tag{G.22}$$

uniformly as $\sigma \rightarrow 0$ and $n_0(\sigma) \rightarrow \infty$ in $\mathbb{R}^{d'}$.

Since $\mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] < \infty$ and $\mathbb{E}_{\mathbf{y} \sim \tilde{f}_{\sigma^2, n_0(\sigma)}}[\|\mathbf{y}\|^2] \rightarrow \mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + d'\sigma^2$ as $n_0(\sigma) \rightarrow \infty$ from Eq. (G.19), for any $\epsilon > 0$, there exists a measurable set $A \subseteq \mathbb{R}^{d'}$ such that:

1. $|\int_A \|\mathbf{y}\|^2 f(\mathbf{y}) d\mathbf{y} - \mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2]| < \epsilon$
2. we can find a sufficiently small σ and a sufficiently large $n_0(\sigma)$ such that $d'\sigma^2 < \epsilon$ and:

$$\int_A \|\mathbf{y}\|^2 \cdot |f(\mathbf{y}) - \tilde{f}_{\sigma^2, n_0(\sigma)}(\mathbf{y})| d\mathbf{y} \leq \epsilon. \tag{G.23}$$

Using the inequality (G.19), we have

$$\begin{aligned}
& \mathbb{E}_{(\mathbf{y}, \hat{\mathbf{y}}) \sim \pi(f, \tilde{f}_{\sigma^2, n_0(\sigma)})}[\|\mathbf{y} - \hat{\mathbf{y}}\|^2] \leq 4 \int_A \|\mathbf{y}\|^2 |f(\mathbf{y}) - \tilde{f}_{\sigma^2, n_0(\sigma)}(\mathbf{y})| d\mathbf{y} \\
& \quad + 4 \int_{\mathbb{R}^{d'} - A} \|\mathbf{y}\|^2 f(\mathbf{y}) d\mathbf{y} + 4 \int_{\mathbb{R}^{d'} - A} \|\mathbf{y}\|^2 \tilde{f}_{\sigma^2, n_0(\sigma)}(\mathbf{y}) d\mathbf{y} \\
& \leq 4\epsilon + 4\epsilon + 4 \left(\mathbb{E}_{\mathbf{y} \sim f}[\|\mathbf{y}\|^2] + d'\sigma^2 + n_0(\sigma)^{-\frac{1}{3}} C_3(\sigma; f) \right. \\
& \quad \left. - \left(\int_A \|\mathbf{y}\|^2 f(\mathbf{y}) d\mathbf{y} - \epsilon \right) \right) \\
& \leq 16\epsilon + 4d'\sigma^2 + 4n_0(\sigma)^{-\frac{1}{3}} C_3(\sigma; f) \leq 24\epsilon,
\end{aligned} \tag{G.24}$$

if we take an $n_0(\sigma)$ large enough such that $n_0(\sigma)^{-\frac{1}{3}} C_3(\sigma; f) \leq \epsilon$. From Eq. (G.16), we have:

$$\tilde{f}_{\sigma^2, n_0(\sigma)}(\mathbf{y}) = \sum_{i=1}^{(n_0(\sigma)+1)d'} \frac{f(\mathbf{y}_i)w_i}{\sum_{j=1}^{(n_0(\sigma)+1)d'} f(\mathbf{y}_j)w_j} \cdot \mathcal{N}(\mathbf{y} - \mathbf{y}_i; \sigma^2 I_{d' \times d'}), \tag{G.25}$$

which is indeed the probability density function of a Gaussian mixture model, and this completes the proof of Theorem 3.2.

Appendix H. The approximation ability of the SNN model in Fig. 1

In this subsection, we analyze the capability of the SNN model to approximate a family of probability density functions for a random variable $\mathbf{x} \sim f_{\mathbf{x}}$, $\mathbf{y} \in \mathbb{R}^{d'}$ characterized by $\mathbf{x} \in D \subseteq \mathbb{R}^d$, $\mathbf{x} \sim \gamma(\cdot)$. We assume the following conditions hold:

1. f_x is uniformly continuous such that for any ϵ , there exists a Δx satisfying $W_2^2(f_x, f_{\tilde{x}}) \leq \epsilon$, $\forall x, \tilde{x} \in D, \|x - \tilde{x}\| \leq \Delta x$.
2. For any x , f_x satisfies the conditions in [Theorem 3.2](#).

Let $\epsilon > 0$ be a small positive number. We can find a $\Delta x > 0$ such that $W_2^2(f_x, f_{\tilde{x}}) < \frac{\epsilon}{4}$ if $\|x - \tilde{x}\| \leq \sqrt{d}\Delta x, \forall x, \tilde{x} \in D$. We can find an equidistance grid point set $X := \{x_i\}_{i=1}^K$ such that the distance between two adjacent points is Δx , $D \subseteq \bigcup_{i=1}^K \bigotimes_{j=1}^d [x_i^j, x_i^j + \Delta x)$, and $\bigotimes_{j=1}^d [x_{i_1}^j, x_{i_1}^j + \Delta x) \cap \bigotimes_{j=1}^d [x_{i_2}^j, x_{i_2}^j + \Delta x) = \emptyset$ if $i_1 \neq i_2$. Thus, for any $x \in D$, there exists $x_i \in X$ satisfying $W_2^2(f_x, f_{x_i}) < \frac{\epsilon}{4}$. Furthermore, for any $x_i = (x_{i_1}^1, \dots, x_{i_d}^d) \in X$, from [Theorem 3.2](#), there exists a probability density function of a Gaussian mixture model:

$$\tilde{f}_{n_{0,i}, \sigma_i^2}(y_{x_i}) = \sum_{r=1}^{n_{0,i}} p_{i,r} \mathcal{N}(y_{x_i} - b_{i,r}, A_{i,r}^T A_{i,r}), \sum_{r=1}^{n_{0,i}} p_{i,r} = 1, \quad (\text{H.1})$$

such that $W_2^2(f_{x_i}, \tilde{f}_{n_{0,i}, \sigma_i^2}) < \frac{\epsilon}{4}, i = 1, \dots, K$. We denote $n_0 := \max_{1 \leq i \leq K} n_{0,i}$ and

$$f_{n_0, \sigma_i^2}(y_{x_i}) = \sum_{r=1}^{n_{0,i}-1} p_{i,r} \mathcal{N}(y_{x_i} - b_{i,r}, A_{i,r}^T A_{i,r}) + \sum_{r=n_{0,i}}^{n_0} \frac{p_{i,n_{0,i}}}{n_0 - n_{0,i} + 1} \mathcal{N}(y_{x_i} - b_{i,n_{0,i}}, A_{i,n_{0,i}}^T A_{i,n_{0,i}}). \quad (\text{H.2})$$

$W_2^2(f_{x_i}, f_{n_{0,i}, \sigma_i^2}) < \frac{\epsilon}{4}, i = 1, \dots, K$ for $f_{n_{0,i}, \sigma_i^2}$ defined in [Eq. \(H.2\)](#) because $W_2^2(f_{x_i}, f_{n_{0,i}, \sigma_i^2}) = W_2^2(\tilde{f}_{x_i}, f_{n_{0,i}, \sigma_i^2})$.

We define a new continuous random variable $\tilde{y}_x$ with a probability distribution $\tilde{f}_x, x \in D$ such that:

$$\tilde{f}_x = f_{n_0, \sigma_i^2}, \text{ if } x \in D \cap \bigotimes_{j=1}^d [x_i^j, x_i^j + \Delta x). \quad (\text{H.3})$$

Therefore, we have

$$\int_D W_2^2(\tilde{f}_x, f_x) \gamma(dx) < \sum_{i=1}^K \int_{D \cap \bigotimes_{j=1}^d [x_i^j, x_i^j + \Delta x)} 2(W_2^2(\tilde{f}_x, f_{x_i}) + W_2^2(f_{x_i}, f_{x_i})) \gamma(dx) = \epsilon. \quad (\text{H.4})$$

We denote

$$Y(\epsilon, X) := \sup_{i=1, \dots, K, s=1, \dots, n_0} (\|b_{i,s}\|^2 + \|A_{i,s}^T A_{i,s}\|_F^2), \quad (\text{H.5})$$

where $b_{i,s}$ and $A_{i,s}^T A_{i,s}$ are the mean vectors and covariance matrices in [Eq. \(H.2\)](#), respectively. Similar to the proof of [Corollary 3.1](#) in [Appendix F](#), for any $\epsilon_1 > 0$, there exists

$$D(\epsilon_1) := D \cap (\bigcup_{i=1}^K \bigotimes_{j=1}^d [x_i^j, x_i^j + \Delta x - \epsilon_1)) \quad (\text{H.6})$$

such that $\gamma(D - D(\epsilon_1)) := \int_{D-D(\epsilon_1)} 1 \gamma(dx) \leq \epsilon$. Additionally, similar to the estimates [Eqs. \(F.13\)](#) and [\(F.14\)](#), we can find an SNN whose output obeys a distribution $\hat{f}_x$ satisfying:

$$\begin{aligned} W_2^2(\hat{f}_x, \tilde{f}_{x_i}) &\leq 6\epsilon Y(\epsilon, X), \quad x \in D(\epsilon_1) := \{x \in D(\epsilon_1) | \|x_i - x\| \leq \|x_j - x\|\} \\ W_2^2(\hat{f}_x, f_{x_i}) &\leq 2 \max_i \mathbb{E}_{y_x \sim f_{n_{0,i}, \sigma_i^2}} [\|y_x\|^2] + 2 \mathbb{E}_{y_x \sim \hat{f}_x} [\|y_x\|^2] \\ &\leq 10 \max_i \mathbb{E}_{y_x \sim f_{n_{0,i}, \sigma_i^2}} [\|y_x\|^2] \leq 10Y(\epsilon, X), \quad x \notin D(\epsilon_1). \end{aligned} \quad (\text{H.7})$$

Therefore, we have

$$\begin{aligned} \int_D W_2^2(\hat{f}_x, \tilde{f}_x) \gamma(dx) &\leq \int_{D(\epsilon_1)} W_2^2(\hat{f}_x, \tilde{f}_x) \gamma(dx) + \int_{D-D(\epsilon_1)} W_2^2(\hat{f}_x, \tilde{f}_x) \gamma(dx) \\ &= 6\epsilon Y(\epsilon, X) + 10\epsilon Y(\epsilon, X) = 16\epsilon Y(\epsilon, X), \end{aligned} \quad (\text{H.8})$$

and

$$\begin{aligned} \int_D W_2^2(\hat{f}_x, f_x) \gamma(dx) &\leq 2 \int_D (W_2^2(\hat{f}_x, \tilde{f}_x) + W_2^2(\tilde{f}_x, f_x)) \gamma(dx) \\ &\leq 32Y(\epsilon, X)\epsilon + 2\epsilon. \end{aligned} \quad (\text{H.9})$$

In [Eq. \(H.9\)](#), $Y(\epsilon, X)$ also depends on ϵ and the choice of the grid point set $X = \{x_i\}_{i=1}^K$. Specifically, as $\epsilon \rightarrow 0$, if we can design a grid point set $\{x_i\}_{i=1}^K$ such that the quantity $\epsilon \cdot Y(\epsilon, X) \rightarrow 0$ (e.g., when y is uniformly bounded for any $x \in D$, by the construction of f_{n_0, σ_i^2} in [Appendix G](#), $\sup_{i=1, \dots, K, s=1, \dots, n_0} (\|b_{i,s}\|^2 + \|A_{i,s}^T A_{i,s}\|_F^2)$ in [Eq. H.5](#) is uniformly bounded and its upper bound is independent of the choice of the grid point set), then there exists an SNN model in [Fig. 1](#) such that the distribution of its output can approximate f_x , the probability distribution of y given the input x , in the squared W_2 sense.

Appendix I. Default training settings

We list the hyperparameters and settings for training the SNN model in [Fig. 1](#) of each example in [Table I.3](#) below.

Table I.3

Training settings for each example.

Hyperparameters	Example 4.1	Example 4.2	Example 4.3	Example 4.4
Gradient descent method	Adam	Adam	Adam	Adam
Learning rate	0.001	0.0005	0.0003	0.001
Weight decay	0.01	0	0.005	0.02
No. of epochs	500	400	2000	400
No. of training trajectories	200	200	400	300
Hidden layers	1	3	3	2
Activation function	ELU	ELU	ELU	ReLU
Neurons in each layer	40	50	10	50
time step Δt	0.1	0.1	0.05	0.1
Number of timesteps N_T	81	21	41	21
Initialization for biases	2	$\mathcal{N}(0, 0.01)$ $(\mathcal{N}((1.5, 0.5)^T, 0.1I_2))$ for the output layer)	$\mathcal{N}(0, 0.03^2)$	0.01
Initialization for means of weights	$\mathcal{N}(0, 0.01^2)$	$\mathcal{N}(0, 0.01^2)$	$\mathcal{N}(0, 0.03^2)$	$\mathcal{N}(0, 0.01^2)$
Initialization for variances of weights	$\mathcal{N}(0, 0.01^2)$	$\mathcal{N}(0, 0.01^2)$	$\mathcal{N}(0, 0.03^2)$	$\mathcal{N}(0, 0.01^2)$

Appendix J. Definitions of different loss metrics

Below, we provide descriptions and definitions for the different loss functions used for comparison in this study. In the following, N denotes the number of samples, and $t_i = i \frac{T}{N_T}, i = 0, \dots, N_T$ denotes a uniform mesh in $[0, T]$.

1. A scaled local time-decoupled squared W_2 distance:

$$\frac{1}{N_T} \sum_{i=1}^{N_T} W_{2,\delta}^{2,e}(X(t_i), \hat{X}(t_i)) = \frac{1}{N_T N} \sum_{i=1}^{N_T} \sum_{j=1}^N W_2^2(v_{X_{0,j},\delta}^e(t_i), \hat{v}_{X_{0,j},\delta}^e(t_i)), \quad (\text{J.1})$$

where $W_{2,\delta}^{2,e}(X(t_i), \hat{X}(t_i))$ is the local squared W_2 distance defined in [Eq. \(2.4\)](#). $v_{X_{0,j},\delta}^e(t)$ and $\hat{v}_{X_{0,j},\delta}^e(t)$ are the empirical conditional probability distributions of X and $\hat{X}$ at time t conditioned on $|X(0) - X_j(0)| \leq \delta$ and $|\hat{X}(0) - X_j(0)| \leq \delta$, respectively ($X_j(0)$ denotes the initial state of the j th trajectory in the training set). $W_2^2(v_{X_{0,j},\delta}^e, \hat{v}_{X_{0,j},\delta}^e)$ is estimated by

$$W_2^2(v_{X_{0,j},\delta}^e, \hat{v}_{X_{0,j},\delta}^e) \approx \text{ot.emd2}\left(\frac{1}{N_j} \tilde{I}_{N_j}, \frac{1}{N_j} \tilde{I}_{N_j}, C_j(t_i)\right), \quad (\text{J.2})$$

where ot.emd2 is the function for solving the earth movers distance problem in the PoT package of Python in [Flamary et al. \(2021\)](#). N_j is

the number of elements in the set $X_j := \left\{ \{X_i\}_{i=0}^T \mid \|X_i(0) - X_j(0)\| \leq \delta, i = 1, \dots, N \right\}$, $\tilde{I}_{N_j}$ is an N_j -dimensional vector whose elements are all 1, and $C_j(t_i) \in \mathbb{R}^{N_j \times N_j}$ is a matrix with entries $(C_j(t_i))_{sr} = \|X_s(t_i) - \hat{X}_r(t_i)\|^2$. $X_s(t_i)$ and $\hat{X}_r(t_i)$ are the states of the s th ground truth trajectory at time t_i in the set X_j and the states of the r th predicted trajectory at time t_i in the set $\hat{X}_j := \left\{ \{\hat{X}_i\}_{i=0}^T \mid \|\hat{X}_i(0) - X_j(0)\| \leq \delta, i = 1, \dots, N \right\}$, respectively.

2. A scaled time-decoupled squared W_2 distance

$$\tilde{W}_2^2(X, \hat{X}) \approx \frac{1}{N_T} \sum_{i=1}^{N_T} W_2^2(v^e(t_i), \hat{v}^e(t_i)),$$

where $v^e(t_i)$ and $\hat{v}^e(t_i)$ are the empirical distributions of $X(t_i)$ and $\hat{X}(t_i)$, respectively. It is estimated by

$$W_2^2(v_N^e(t_i), \hat{v}_N^e(t_i)) \approx \text{ot.emd2}\left(\frac{1}{N} \tilde{I}_N, \frac{1}{N} \tilde{I}_N, C(t_i)\right), \quad (\text{J.3})$$

where `ot.emd2` is the function for solving the earth movers distance problem in the `ot` package of Python in Flamary et al. (2021). N is the number of ground truth and predicted samples, $\tilde{I}_N$ is an N -dimensional vector whose elements are all 1, and $C(t_i) \in \mathbb{R}^{N \times N}$ is a matrix with entries $(C(t_i))_{sj} = \|X_s(t_i) - \hat{X}_j(t_i)\|^2$. $X_s(t_i)$ and $\hat{X}_j(t_i)$ are the states of the s th ground truth trajectory at time t_i and the states of the j th predicted trajectory at time t_i , respectively.

3. MMD (maximum mean discrepancy) (Li et al., 2015):

$$\begin{aligned} \text{MMD}(\{X\}, \{\hat{X}\}) &= \frac{1}{N_T} \sum_{i=1}^{N_T} \mathbb{E}[K(\{X(t_i)\}, \{X(t_i)\})] \\ &\quad - 2\mathbb{E}[K(\{X(t_i)\}, \{\hat{X}(t_i)\})] + \mathbb{E}[K(\{\hat{X}(t_i)\}, \{\hat{X}(t_i)\})], \end{aligned} \quad (\text{J.4})$$

where K is the standard radial basis function (or Gaussian kernel) with the multiplier equal to 2 and the number of kernels equal to 5. $\{X(t_i)\}$ and $\{\hat{X}(t_i)\}$ denote the sets of ground truth observations and predicted trajectories at time t_i , respectively.

4. Mean squared error (MSE):

$$\text{MSE}(X, \hat{X}) = \frac{1}{N_T N} \sum_{i=1}^{N_T} \sum_{s=1}^N \|X_s(t_i) - \hat{X}_s(t_i)\|^2. \quad (\text{J.5})$$

$X_s(t_i)$ and $\hat{X}_s(t_i)$ are the states of the s th ground truth trajectory at time t_i and the states of the s th predicted trajectory at time t_i , respectively.

5. Mean² + Var loss function:

$$\begin{aligned} (\text{Mean}^2 + \text{Var})(X, \hat{X}) &= \frac{1}{N_T} \sum_{i=1}^{N_T} \left(\frac{1}{N} \sum_{s=1}^N \|X_s(t_i) - \hat{X}_s(t_i)\|^2 \right. \\ &\quad \left. + |\text{Var}(X(t_i)) - \text{Var}(\hat{X}(t_i))| \right) \end{aligned} \quad (\text{J.6})$$

where

$$\text{Var}(X(t_i)) = \sum_{s=1}^N \left\| X_s(t_i) - \sum_{i=1}^N \frac{1}{N} X_i(t_i) \right\|^2. \quad (\text{J.7})$$

$X_s(t_i)$ and $\hat{X}_s(t_i)$ are the states of the s th ground truth trajectory at time t_i and the states of the s th predicted trajectory at time t_i , respectively.

Appendix K. Using a neural network to approximate the diffusion function in Example 4.4

Here, we apply a feedforward neural network with deterministic weights and biases to approximate the diffusion function in a jump-diffusion process when the form of the diffusion function is unknown. We consider the following jump-diffusion process:

$$\begin{aligned} dX_t &= 0.05dt + \sigma_0 \sqrt{|X_t|} dB_t + \int_U \xi X_t d\tilde{N}(v(d\xi)dt), \quad t \in [0, 2], \\ \xi &\sim \mathcal{N}(\beta_0, \sigma_1^2), \quad X_0 \sim \mathcal{N}(2, \sigma_2^2). \end{aligned} \quad (\text{K.1})$$

In Eq. (K.1), σ_0 is a positive constant, and B_t and $\tilde{N}_t$ are an independent Wiener process and a compensated Poisson process, respectively. We use the following approximate jump-diffusion process to approximate Eq. (K.1):

$$d\hat{X}_t = 0.05dt + \hat{\sigma}(\hat{X}_t) d\hat{B}_t + \int_U \hat{\xi} \hat{X}_t d\hat{N}(v(d\xi)dt), \quad t \in [0, 2], \quad \hat{X}_0 = X_0. \quad (\text{K.2})$$

In Eq. (K.2), $\hat{\sigma}(\hat{X}_t)$ is a deterministic feedforward parameterized neural network that takes the state $\hat{X}_t$ as the input. We aim to approximate the ground truth diffusion function $|\sigma(X_t)| := |\sigma_0| \sqrt{|X_t|}$ in Eq. (K.1) using the approximate $|\hat{\sigma}(X_t)|$ in Eq. (K.2). $\hat{\xi}$ is the output of an SNN when the input is 1, which aims at approximating the distribution of ξ , and $\hat{B}_t$ is another Wiener process independent of the Wiener process B_t in Eq. (K.1) while $\hat{N}$ is another compensated Poisson process independent of B_t , $\hat{B}_t$ and $\tilde{N}_t$. We train both the deterministic neural network $\hat{\sigma}(\hat{X}_t)$ and the SNN that approximates the distribution of ξ simultaneously by minimizing our local time-decoupled squared W_2 loss function Eq. (2.6). The parameterized neural network $\hat{\sigma}(\hat{X}_t)$ consists of two hidden layers with fifty neurons in each layer. All other hyperparameters are the same as those used in Example 4.4, described in Table I.3. We vary the variance of ξ as well as the value of σ_0 when reconstructing the diffusion function and the distribution of ξ in the jump function.

The trajectories generated by the reconstructed jump-diffusion process align well with ground truth trajectories (shown in Fig. K.7(a)). Furthermore, the error in the reconstructed diffusion function, as well as the error in the reconstructed distribution of $\hat{\xi}$, is small. Thus, when

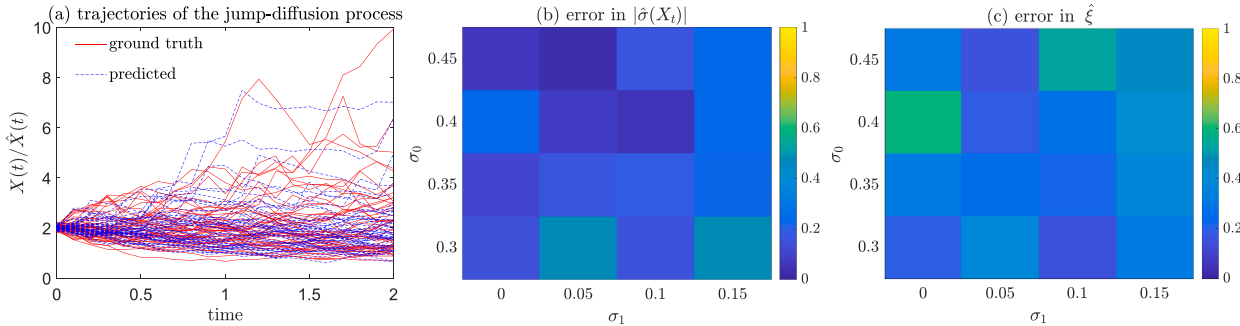


Fig. K.7. (a) Ground truth trajectories versus the reconstructed trajectories when $\sigma_0 = 0.3$, $\beta_0 = 0.3$, $\sigma_1 = 0.15$, $\sigma_2 = 0.1$, and $\delta = 0.1$ in the loss function Eq. (2.6). For clarity, 50 ground truth trajectories and 50 reconstructed trajectories are plotted. (b) The error in the reconstructed diffusion function $\hat{\sigma}$. Here, the error denotes the relative squared L^2 error: $\frac{\int_0^T (|\hat{\sigma}(X_t)| - \sigma_0 \sqrt{|X_t|})^2 dt}{\int_0^T (\sigma_0 \sqrt{|X_t|})^2 dt}$. (c) The error in the reconstructed distribution of $\hat{\xi}$ (Eq. (4.1)). In (b) and (c), the values of other parameters are: $\beta_0 = 0.3$, $\sigma_2 = 0.1$, and the size of neighborhood $\delta = 0.1$ in the loss function Eq. (J.1). Errors are the average error over three repeated experiments.

the form of the diffusion function is unknown while the diffusion function itself is deterministic, we could consider using a deterministic feed-forward neural network to reconstruct it while using an SNN (Fig. 1) to simultaneously reconstruct the distribution of ξ in the jump magnitude function in Eq. (K.1) by minimizing the loss function Eq. (J.1).

References

- Arjovsky, M., Chintala, S., & Bottou, L. (2017). Wasserstein generative adversarial networks. In *International conference on machine learning* (pp. 214–223). PMLR.
- Balasubramanian, K., Goldstein, L., Ross, N., & Salim, A. (2024). Gaussian random field approximation via Stein's method with applications to wide random neural networks. *Applied and Computational Harmonic Analysis*, 72, 101668.
- Bernton, E., Jacob, P. E., Gerber, M., & Robert, C. P. (2019). On parameter estimation with the Wasserstein distance. *Information and Inference: A Journal of the IMA*, 8(4), 657–676.
- Blanchet, J., & Kang, Y. (2021). Sample out-of-sample inference based on Wasserstein distance. *Operations Research*, 69(3), 985–1013.
- Blundell, C., Cornebise, J., Kavukcuoglu, K., & Wierstra, D. (2015). Weight uncertainty in neural network. In *International conference on machine learning* (pp. 1613–1622). PMLR.
- Bayesian neural network derivation (2020). <https://www.zhihu.com/tardis/bd/art/263053978>.
- Böhm, V., Lanusse, F., & Seljak, U. (2019). Uncertainty quantification with generative models. arXiv preprint arXiv:1910.10046.
- Boukaraichi, H., Akkari, N., Casenave, F., & Ryckelynck, D. (2022). Uncertainty quantification in a mechanical submodel driven by a Wasserstein-GAN. *IFAC-PapersOnLine*, 55(20), 469–474.
- Brunton, S. L., Proctor, J. L., & Kutz, J. N. (2016). Discovering governing equations from data by sparse identification of nonlinear dynamical systems. *Proceedings of the National Academy of Sciences*, 113(15), 3932–3937. <https://doi.org/10.1073/pnas.1517384113>
- Carrasi, A., Bocquet, M., Bertino, L., & Evensen, G. (2018). Data assimilation in the geosciences: An overview of methods, issues, and perspectives. *Wiley Interdisciplinary Reviews: Climate Change*, 9(5), e535. <https://doi.org/10.1002/wcc.535>
- Caruso, A., Füh, M., Alvarez-Sánchez, R., Belli, S., Diack, C., Maass, K. F., Schwab, D., Kettenberger, H., & Mazer, N. A. (2019). Ocular half-life of intravitreal biologics in humans and other species: Meta-analysis and model-based prediction. *Molecular Pharmacology*, 17(2), 695–709.
- Chappelow, A. V., & Kaiser, P. K. (2008). Neovascular age-related macular degeneration: potential therapies. *Drugs*, 68, 1029–1036.
- Chen, R. T. Q., Rubanova, Y., Bettencourt, J., & Duvenaud, D. K. (2018). Neural ordinary differential equations. *Advances in Neural Information Processing Systems*, 31.
- Chib, S., & Greenberg, E. (2002). Bayesian inference for regression models with ARIMA errors. *Journal of Econometrics*, 109, 385–408. [https://doi.org/10.1016/S0304-4076\(01\)00136-7](https://doi.org/10.1016/S0304-4076(01)00136-7)
- Clement, P., & Desch, W. (2008). An elementary proof of the triangle inequality for the Wasserstein metric. *Proceedings of the American Mathematical Society*, 136(1), 333–339.
- Crawshaw, J., Gaffney, E., Gertz, M., Maini, P., & Caruso, A. (2025). Modelling the ocular pharmacokinetics and pharmacodynamics of ranibizumab for improved understanding and data collection strategies in ocular diseases. *Investigative Ophthalmology & Visual Science*, 66(6), 20.
- Cresswell, W. (2015). Maximum likelihood estimation for dynamical models: A guide to implementation. *Journal of Theoretical Biology*, 382, 123–132. <https://doi.org/10.1016/j.jtbi.2015.06.028>
- Cuturi, M. (2013). Sinkhorn distances: Lightspeed computation of optimal transport. *Advances in Neural Information Processing Systems*, 26.
- Deng, Y., Shao, S. S., Mogilner, A., & Xia, M. (2025). Adaptive hyperbolic-cross-space mapped Jacobi method on unbounded domains with applications to solving multidimensional spatiotemporal integrodifferential equations. *Journal of Computational Physics*, 520, 113492.
- Flamary, R., Courty, N., Gramfort, A., Alaya, M. Z., Boisbunon, A., Chambon, S., Chapel, L., Corenflos, A., Fatras, K., Fournier, N., Gautheron, L., Gayraud, N. T. H., Janati, H., Rakotomamonjy, A., Redko, I., Rolet, A., Schutz, A., Seguy, V., Sutherland, D. J., ... Vayer, T. (2021). POT: Python optimal transport. *Journal of Machine Learning Research*, 22(78), 1–8.
- Fournier, N., & Guillin, A. (2015). On the rate of convergence in Wasserstein distance of the empirical measure. *Probability Theory and Related Fields*, 162(3), 707–738.
- Gal, Y., & Ghahramani, Z. (2016). Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. *Proceedings of the 33rd International Conference on Machine Learning (ICML)*, 48, 1050–1059. <https://proceedings.mlr.press/v48/gal16.html>
- Gonon, L., Grigoryeva, L., & Ortega, J.-P. (2023). Approximation bounds for random neural networks and reservoir systems. *The Annals of Applied Probability*, 33(1), 28–69.
- Grecksch, W., & Kloeden, P. E. (1996). Time-discretised Galerkin approximations of parabolic stochastic PDEs. *Bulletin of the Australian Mathematical Society*, 54(1), 79–85.
- Gulrajani, I., Ahmed, F., Arjovsky, M., Dumoulin, V., & Courville, A. C. (2017). Improved training of Wasserstein GANs. *Advances in Neural Information Processing Systems*, 30.
- He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 770–778).
- Hoffer, E., Hubara, I., & Soudry, D. (2017). Train longer, generalize better: closing the generalization gap in large batch training of neural networks. *Advances in Neural Information Processing Systems*, 30.
- Hornik, K., Stinchcombe, M., & White, H. (1989). Multilayer feedforward networks are universal approximators. *Neural Networks*, 2(5), 359–366.
- Hutton-Smith, L. A., Gaffney, E., Byrne, H. M., Caruso, A., Maini, P., & Mazer, N. A. (2018). Theoretical insights into the retinal dynamics of VEGF in patients treated with ranibizumab, based on an ocular pharmacokinetic/pharmacodynamic model. *Molecular Pharmacology*, 15(7), 2770–2784.
- Hutton-Smith, L. A., Gaffney, E. A., Byrne, H. M., Maini, P. K., Schwab, D., & Mazer, N. A. (2016). A mechanistic model of the intravitreal pharmacokinetics of large molecules and the pharmacodynamic suppression of ocular vascular endothelial growth factor levels by ranibizumab in patients with neovascular age-related macular degeneration. *Molecular Pharmacology*, 13(9), 2941–2950.
- Kloeden, P. E., & Platen, E. (1992). Numerical Solution of Stochastic Differential Equations. Berlin: Springer Nature.
- Lamirande, P., Gaffney, E. A., Gertz, M., Maini, P. K., Crawshaw, J. R., & Caruso, A. (2024). A first-passage model of intravitreal drug delivery and residence time—influence of ocular geometry, individual variability, and injection location. *Investigative Ophthalmology & Visual Science*, 65(12), 21.
- Leshno, M., Lin, V. Y., Pinkus, A., & Schocken, S. (1993). Multilayer feedforward networks with a nonpolynomial activation function can approximate any function. *Neural Networks*, 6(6), 861–867.
- Li, Y., Swersky, K., & Zemel, R. (2015). Generative moment matching networks. In *International conference on machine learning* (pp. 1718–1727). PMLR.
- Liu, W., & Röckner, M. (2015). Stochastic Partial Differential Equations: an Introduction. Springer.
- Loskot, P., Atitey, K., & Mihaylova, L. (2019). Comprehensive review of models and methods for inferences in bio-chemical reaction networks. *Frontiers in Genetics*, 10, 549.
- Lu, Y., & Lu, J. (2020). A universal approximation theorem of deep neural networks for expressing probability distributions. *Advances in Neural Information Processing Systems*, 33, 3094–3105.
- Manita, O. A., Peletier, M. A., Portegies, J. W., Sanders, J., & Senen-Cerda, A. (2022). Universal approximation in dropout neural networks. *Journal of Machine Learning Research*, 23(19), 1–46.
- Marin, J.-M., Pudlo, P., Robert, C. P., & Ryder, R. J. (2021). Approximate Bayesian computational methods. *Statistics and Computing*, 31, 239–259. <https://doi.org/10.1007/s11222-020-09968-y>
- Mehri, S., & Scheutzw, M. (2019). A stochastic Gronwall lemma and well-posedness of path-dependent SDEs driven by martingale noise. arXiv preprint arXiv:1908.10646.
- Merton, R. C. (1976). Option pricing when underlying stock returns are discontinuous. *Journal of Financial Economics*, 3(1–2), 125–144.
- Mitchell, P., Liew, G., Gopinath, B., & Wong, T. Y. (2018). Age-related macular degeneration. *The Lancet*, 392(10153), 1147–1159.
- Mullachery, V., Khera, A., & Husain, A. (2018). Bayesian neural networks. arXiv preprint arXiv:1801.07710.
- Nardini, J. T., & Bortz, D. M. (2019). The influence of numerical error on parameter estimation and uncertainty quantification for advective PDE models. *Inverse Problems*, 35(6), 065003.
- Nardini, J. T., Lagergren, J. H., Hawkins-Daarud, A., Curtin, L., Morris, B., Rutter, E. M., Swanson, K. R., & Flores, K. B. (2020). Learning equations from biological data with limited time samples. *Bulletin of Mathematical Biology*, 82, 1–33.
- Neal, R. M. (2012). Bayesian Learning for Neural Networks (118). Springer Science & Business Media.
- NYU HPC (2025). NYU HPC hardware specs. <https://sites.google.com/nyu.edu/nyu-hpc/hpc-systems/greene/hardware-specs>.
- Panaretos, V. M., & Zemel, Y. (2019). Statistical aspects of Wasserstein distances. *Annual Review of Statistics and its Application*, 6(1), 405–431.
- Pathak, J., Lu, Z., Hunt, B., Girvan, M., & Ott, E. (2018). Model-free prediction of large spatiotemporally chaotic systems from data: A reservoir computing approach. *Physical Review Letters*, 120(2), 024102. <https://doi.org/10.1103/PhysRevLett.120.024102>
- Pillow, J. W., Shlens, J., Paninski, L., Sher, A., Litke, A. M., Chichilnisky, E. J., & Simoncelli, E. P. (2008). Spatio-temporal correlations and visual signalling in a complete neuronal population. *Nature*, 454(7207), 995–999.
- Platen, E., & Bruti-Liberati, N. (2010). Numerical solution of stochastic differential equations with jumps in finance (vol. 64). Springer Science & Business Media.
- Raissi, M., Perdikaris, P., & Karniadakis, G. E. (2017). Machine learning of linear differential equations using gaussian processes. *Journal of Computational Physics*, 348, 683–693. <https://doi.org/10.1016/j.jcp.2017.07.050>
- Raissi, M., Perdikaris, P., & Karniadakis, G. E. (2019). Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378, 686–707.
- Rezende, D. J., Mohamed, S., & Wierstra, D. (2014). Stochastic backpropagation and approximate inference in deep generative models. *Proceedings of the 31st International Conference on Machine Learning (ICML)*, 32, 1278–1286. <https://proceedings.mlr.press/v32/rezende14.html>
- Roda, W. C., Varughese, M. B., Han, D., & Li, M. Y. (2020). Why is it difficult to accurately predict the COVID-19 epidemic? *Infectious Disease Modelling*, 5, 271–281.
- Senan, S., Ali, M. S., Vadel, R., & Arik, S. (2017). Decentralized event-triggered synchronization of uncertain Markovian jumping neutral-type neural networks with mixed delays. *Neural Networks*, 86, 32–41.
- Shen, J., Tang, T., & Wang, L.-L. (2011). Spectral Methods: Algorithms, Analysis and Applications (41). Springer Science & Business Media.
- Shen, J., & Wang, L. (2010). Sparse spectral approximations of high-dimensional problems based on hyperbolic cross. *SIAM Journal on Numerical Analysis*, 48, 1087–1109.
- Tang, C., & Salakhutdinov, R. R. (2013). Learning stochastic feedforward neural networks. *Advances in Neural Information Processing Systems*, 26.

- Vadivel, R., Ali, M. S., Alzahrani, F., Cao, J., & Joo, Y. H. (2020). Synchronization of decentralized event-triggered uncertain switched neural networks with two additive time-varying delays. *Nonlinear Analysis: Modelling and Control*, 25(2), 183–205.
- Villani, C. (2009). *Optimal Transport: Old and New* (338). Springer.
- Xia, M., Li, X., Shen, Q., & Chou, T. (2024a). An efficient Wasserstein-distance approach for reconstructing jump-diffusion processes using parameterized neural networks. *Machine Learning: Science and Technology*, 5, 045052.
- Xia, M., Li, X., Shen, Q., & Chou, T. (2024b). Squared Wasserstein-2 distance for efficient reconstruction of stochastic differential equations. *arXiv preprint arXiv:2401.11354*
- Xia, M., & Shen, Q. (2024). A local squared Wasserstein-2 method for efficient reconstruction of models with uncertainty. *arXiv preprint arXiv:2406.06825*
- Yu, T., Yang, Y., Li, D., Hospedales, T., & Xiang, T. (2021). Simple and effective stochastic neural networks. In *Proceedings of the AAAI conference on artificial intelligence* (pp. 3252–3260). (35).