

A TRANSITION-DENSITY-BASED OPERATOR LEARNING METHOD FOR FOKKER-PLANCK EQUATIONS WITH VARIOUS INITIAL CONDITIONS*

LI ZENG[†], XIAOLIANG WAN[‡], YAOBIN WANG[§], FABIO NOBILE[¶], AND TAO ZHOU^{||}

Abstract. Solving Fokker-Planck equations (FPEs) for multiple initial conditions typically requires repeated computations, leading to substantial computational costs. In this work, we propose a transition-density-based operator learning method to efficiently approximate the solution operator of FPEs with various initial conditions. The core idea is to learn the transition probability density function (PDF) of the underlying stochastic differential equation (SDE), from which the solution associated with a new initial distribution can be obtained through the Chapman–Kolmogorov equation without retraining the model. A major challenge in learning the transition PDF lies in the singular behavior induced by the Dirac initial condition. To address it, we introduce a conditional normalizing flow whose base distribution is given by the explicit transition PDF of a linearized SDE. This base distribution captures the short-time behavior of the target transition PDF and allows the normalizing flow to learn a near-identity transformation at small times. We further incorporate a time-weighted loss function to stabilize training near the initial time and develop an importance-sampling strategy for evaluating solutions associated with general initial conditions. A variety of numerical experiments are presented to illustrate the effectiveness and robustness of the proposed method.

Key words. Conditional Normalizing flow, Fokker-Planck equation, Operator learning

MSC codes. 68T07, 62G07, 65M99

1. Introduction. The FPE describes the time evolution of the PDF associated with a stochastic dynamical system of diffusion type, serving as a deterministic representation of the stochastic process. Numerically solving the FPE is of high interest in many fields, such as statistical physics, chemistry, biology, finance, and data science, as it enables uncertainty quantification and statistical inference. Moreover, in many of these applications, such as ensemble forecasting [24] and data assimilation [30], the governing stochastic dynamic is fixed, whereas the initial distribution varies across scenarios or posterior updates. Solving the FPE repeatedly for each initial condition can therefore lead to substantial computational costs.

Classical numerical discretization schemes for the FPE include the finite element method [9], the finite difference method [22], the path integral method [45], and the Jordan-Kinderlehrer-Otto (JKO) scheme [21], among others. Nevertheless, these traditional approaches incur high computational cost and suffer from poor scalability in high-dimensional settings. To overcome these challenges, deep learning methods [11, 33, 37] bring a new paradigm. In general, deep learning-based methods for solving

*Submitted to the editor's DATE.

Funding: The first author was supported by the NSF of China (under grant 12401566). The second author was supported by NSF grant DMS-2513234. The last author was supported by the NSF of China (under grants 12288201 and 12461160275).

[†]School of Mathematics and Statistics, Fuzhou University, Fuzhou, China (lzeng@fzu.edu.cn).

[‡]Department of Mathematics and Center for Computation and Technology, Louisiana State University, Baton Rouge 70803, USA (xlwan@lsu.edu).

[§]Faculty of Science and Technology, Beijing Normal-Hong Kong Baptist University, Zhuhai 519087, China (wangyaobin@bnu.edu.cn).

[¶]CSQI, École Polytechnique Fédérale de Lausanne, Lausanne, Switzerland (Fabio.Nobile@epfl.ch).

^{||}Institute of Computational Mathematics and Scientific/Engineering Computing, Academy of Mathematics and Systems Science, Chinese Academy of Sciences, Beijing, China (tzhoul@lsec.cc.ac.cn).

the FPE can be categorized into two main classes: those that directly represent the solution with neural networks [7, 31, 42, 46, 47, 51], and those that employ generative modeling techniques, such as normalizing flows [12, 13, 29, 39, 49, 50] and score-based approaches [3, 19, 52]. Alternatively, Bruna et al. [6, 46] approximate the solution of the interacting-particle FPE using a Gaussian mixture model, and numerically solve the ordinary differential equation (ODE) system for the mixture coefficients using the Neural Galerkin method, subsequently normalizing the solution at each time step to enforce the PDF constraints. In parallel, several techniques have been developed for some special FPEs that exploit the reinterpretation as the L^2 -Wasserstein gradient flow of the Kullback-Leibler divergence [23, 26].

While a variety of numerical methods have shown promising performance in solving FPEs with prescribed initial conditions, there remains a lack of methods capable of efficiently handling multiple initial conditions. The aforementioned methods have to solve the FPE anew every time a new initial condition arrives. This challenge can be viewed from the perspective of operator learning, where the objective is to learn a mapping from an initial distribution to the corresponding time-dependent solution of the FPE. Recently, Wang et al. [43] proposed a Gaussian-mixture latent surrogate that directly learns the evolution from initial to corresponding transient PDF representation, while focusing on Gaussian-mixture initial distributions and a truncated computational domain. Yang et al. [48] proposed a pseudoreversible normalizing flow and leveraged samples generated from an SDE to learn the terminal solution of the corresponding transition FPE. In the general data-driven setting, generative models are also widely used to learn the conditional PDF, such as conditional Glow [28], monotone generative adversarial networks [2], the gradient of a partially input-convex neural network, and Neural ODE [44], to name a few. For general parametric PDEs, some techniques have been proposed to learn the solution operator, such as the deep operator network (DeepONet) [27], the Fourier neural operator (FNO) [25], and their variants. Gaby et al. [15] learned the solution operator of the evolution PDEs for various initial conditions by learning a control vector field in the parameter space and solving the ODE system of the neural network parameters. There are also works that learn the Green's function to construct the solution operator, leveraging the explicit representation of the solution in terms of the Green's function [1, 5, 40]. However, direct application of these methods to FPEs requires additional treatments of boundary conditions and PDF constraints. In this work, we take a different route: rather than learning the dependence of the FPE solution on the initial distribution directly, we learn the transition PDF of the underlying stochastic dynamics. Solutions for new initial distributions are then recovered through the Chapman-Kolmogorov equation. Our framework is built on the following ideas.

First, we reformulate operator learning for the FPE as the problem of approximating the transition PDF, i.e., the solution of the FPE starting at initial time from a Dirac mass at an arbitrary location $\mathbf{x}_0$. We develop a conditional normalizing flow parametric in time and $\mathbf{x}_0$ to learn the transition PDF. Benefiting from the inherent property of a normalizing flow, the proposed approach naturally satisfies the PDF constraints without requiring additional enforcement. We mention that the work [35, 36] developed in parallel to ours, also proposes a conditional normalizing flow to approximate the transition PDF, however, in a Neural Galerkin framework.

Second, we propose a time- and initial state-dependent base distribution to improve the conditioning of the conditional normalizing flow. In contrast to the commonly used standard normal distribution, our base distribution is given by the explicit solution of a linearized SDE parametric with respect to the initial state $\mathbf{x}_0$. The lin-

earized SDE is obtained by applying a first-order Taylor expansion to the drift term and a zero-order expansion to the diffusion term of the target SDE. Using the solution of the linearized SDE as the base random process not only respects the causality for evolution PDEs [8] while bypassing the singular initial condition, but also alleviates training difficulties. We show, in particular, that the total variation distance between the transition PDF of the linearized process and that of the true process is of order $\mathcal{O}(t)$ while the Kullback–Leibler divergence scales as $\mathcal{O}(t^2)$ as $t \rightarrow 0$. The corresponding transformation tends to the identity map as $t \rightarrow 0$, making it easier to learn than transformation based on a fixed standard normal base distribution at all times.

Third, we propose a time-weighted loss function to mitigate numerical instabilities arising at small times when the transition PDF approaches a Dirac measure. The weight function takes the form t^{d+2} or $t^{\frac{d}{2}+2}$ depending on the distribution of the training points. It is designed to capture the temporal scaling of the residual when a small perturbation around the base PDF is considered. Introducing this weighting strategy can be interpreted as a trade-off between maintaining causality and alleviating optimization difficulties.

Fourth, we incorporate the learned conditional normalizing flow and an importance-sampling Monte Carlo estimator to approximate the final solution of the original FPE.

The remainder of this paper is organized as follows. Section 2 formalizes the operator learning problem for the FPE under different initial conditions. In section 3, we introduce the conditional normalizing flow with base distribution provided by the linearized SDE. In section 4, we analyze the behavior of the loss function as $t \rightarrow 0$, and we propose a time-weighted loss and an adaptive sampling algorithm. Importance sampling is discussed in section 5 to efficiently evaluate $p(\mathbf{x}, t)$. Section 6 presents a series of numerical experiments that illustrate the effectiveness and robustness of the proposed method, followed by concluding remarks in section 7. Additional details of certain proofs and calculations are provided in the Appendix.

2. Problem setup and reformulation. Consider the state variable $\mathbf{X}_t$ modeled by the following non-degenerate SDE,

$$(2.1) \quad d\mathbf{X}_t = \mathbf{f}(\mathbf{X}_t) dt + \mathbf{g}(\mathbf{X}_t) d\mathbf{W}_t, \quad t > 0,$$

where $\mathbf{f}(\mathbf{X}_t) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ and $\mathbf{g}(\mathbf{X}_t) : \mathbb{R}^d \rightarrow \mathbb{R}^{d \times m}$, $m \geq d$ denote the drift and diffusion terms, respectively, and $\mathbf{W}_t$ is an m -dimensional standard Wiener process. The PDF $p(\mathbf{x}, t)$ for $\mathbf{X}_t$ satisfies the corresponding time-dependent FPE:

$$(2.2) \quad \frac{\partial p}{\partial t} = \mathcal{L}_{\mathbf{f}, \mathbf{g}}^* p := -\nabla \cdot (p \mathbf{f}) + \frac{1}{2} \nabla \cdot \nabla \cdot (\mathbf{g} \mathbf{g}^T p).$$

In general, equation (2.2) is defined on $\mathbb{R}^d$ with the following boundary condition

$$(2.3) \quad p(\mathbf{x}) \rightarrow 0 \quad \text{as} \quad |\mathbf{x}|_2 \rightarrow \infty.$$

Meanwhile, the solution as a probability density function should be conservative and non-negative, i.e.,

$$(2.4) \quad \int_{\mathbb{R}^d} p(\mathbf{x}, t) d\mathbf{x} \equiv 1, \quad \text{and} \quad p(\mathbf{x}, t) \geq 0, \quad \forall t \geq 0.$$

In this work, we address scenarios with varying initial conditions and develop a transition-density-based operator learning method that generalizes across them,

thereby eliminating the need to solve the corresponding FPE repeatedly for different initial conditions, i.e., we consider the problem

$$(2.5) \quad \begin{cases} \frac{\partial p(\mathbf{x}, t)}{\partial t} = \mathcal{L}_{\mathbf{f}, \mathbf{g}}^* p(\mathbf{x}, t), & \mathbf{x} \in \mathbb{R}^d, t > 0, \\ p(\mathbf{x}, 0) = p_0(\mathbf{x}), \quad p_0 \in \mathcal{G}, & \mathbf{x} \in \mathbb{R}^d \end{cases}$$

where $\mathcal{G}$ is a functional class including all possible initial distributions of interest.

To achieve this, we turn to the transition PDF $p(\mathbf{x}, t | \mathbf{x}_0)$ given an initial state. The corresponding transition PDF satisfies the same governing equation as the original FPE (2.5), but with a Dirac Delta function as the initial condition,

$$(2.6) \quad \begin{cases} \frac{\partial p(\mathbf{x}, t | \mathbf{x}_0)}{\partial t} = \mathcal{L}_{\mathbf{f}, \mathbf{g}}^* p(\mathbf{x}, t | \mathbf{x}_0), & \mathbf{x} \in \mathbb{R}^d, t > 0, \\ p(\mathbf{x}, 0 | \mathbf{x}_0) = \delta(\mathbf{x} - \mathbf{x}_0), & \mathbf{x} \in \mathbb{R}^d. \end{cases}$$

Applying the Chapman-Kolmogorov equation yields the solution to the FPE (2.5),

$$(2.7) \quad p(\mathbf{x}, t) = \int_{\Omega_0} p(\mathbf{x}, t | \mathbf{x}_0) p_0(\mathbf{x}_0) d\mathbf{x}_0.$$

This motivates us to solve the FPE for different initial conditions by approximating the transition PDF $p(\mathbf{x} | \mathbf{x}_0, t)$. In this work, we extend the normalizing flow framework [13, 49] to handle the transition FPE with Dirac delta functions centered at various points as initial conditions.

2.1. Notation. For convenience, we introduce here some notations that will be used throughout the paper. For a vector $x \in \mathbb{R}^d$, $|x|$ denotes its Euclidean norm. For a matrix $A \in \mathbb{R}^{d \times d}$, $\|A\|$ denotes its spectral norm induced by the Euclidean norm. Likewise, for a bilinear map $B : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^d$, $\|B\|$ denotes its operator norm $\|B\| := \sup_{x, y \in \mathbb{R}^d, x \neq 0, y \neq 0} \frac{|B(x, y)|}{|x| |y|}$.

Given two probability measures ν and μ on $\mathbb{R}^d$, $\nu \ll \mu$, $H(\nu | \mu)$ denotes the relative entropy of ν w.r.t. μ , and $\|\nu - \mu\|_{\text{TV}}$ denotes the total variation distance between ν and μ .

$$H(\nu | \mu) := \int_{\mathbb{R}^d} \log \left(\frac{d\nu}{d\mu} \right) d\nu, \quad \|\nu - \mu\|_{\text{TV}} := \sup_{A \in \mathcal{B}(\mathbb{R}^d)} |\nu(A) - \mu(A)|.$$

Furthermore, if $\mu \ll \rho$, and $d\nu = p_\nu d\rho$, $d\mu = p_\mu d\rho$, then

$$H(\nu | \mu) = \int_{\mathbb{R}^d} p_\nu(\mathbf{x}) \log \left(\frac{p_\nu(\mathbf{x})}{p_\mu(\mathbf{x})} \right) d\rho(\mathbf{x}), \quad \|\nu - \mu\|_{\text{TV}} = \frac{1}{2} \int_{\mathbb{R}^d} |p_\nu(\mathbf{x}) - p_\mu(\mathbf{x})| d\rho(\mathbf{x}).$$

Moreover, given a measurable map $T : \mathbb{R}^d \rightarrow \mathbb{R}^d$, $T_\# \nu$ denotes the pushforward measure of ν through T , i.e., $(T_\# \nu)(A) := \nu(T^{-1}(A)) = \nu(\{\mathbf{x} \in \mathbb{R}^d : T(\mathbf{x}) \in A\})$ for any $A \in \mathcal{B}(\mathbb{R}^d)$. By $\mathcal{N}(\cdot; \boldsymbol{\mu}, \boldsymbol{\Sigma})$ we denote the PDF of a multivariate normal distribution with mean $\boldsymbol{\mu}$ and covariance matrix $\boldsymbol{\Sigma}$. For $\mathbf{f} \in L^2(\mathbb{R}^d, \nu; \mathbb{R}^m)$, $\|\mathbf{f}\|_{L^2(\mathbb{R}^d, \nu; \mathbb{R}^m)}$ denotes its L^2 -norm w.r.t. the measure ν , i.e., $\|\mathbf{f}\|_{L^2(\mathbb{R}^d, \nu; \mathbb{R}^m)} = \left(\int_{\mathbb{R}^d} |\mathbf{f}(\mathbf{x})|^2 d\nu(\mathbf{x}) \right)^{1/2}$.

3. The conditional normalizing flow. In this section, we introduce a conditional normalizing flow to approximate the transition PDF $p(\mathbf{x}, t | \mathbf{x}_0)$ given the initial state $\mathbf{x}_0$. Before proceeding to the conditional normalizing flow, we briefly review the idea of normalizing flows.

A normalizing flow seeks an invertible transformation T from the unknown random vector $\mathbf{X} \in \mathbb{R}^d$ to a given random vector $\mathbf{Z} \in \mathbb{R}^d$ with a well-known and easy to sample distribution. According to the change of variables formula, the PDF of $\mathbf{X} = T^{-1}(\mathbf{Z})$ can be expressed by,

$$(3.1) \quad p_{\mathbf{X}}(\mathbf{x}) = p_{\mathbf{Z}}(T(\mathbf{x})) \left| \det \nabla_{\mathbf{x}} T(\mathbf{x}) \right|.$$

Furthermore, by allowing the distribution of the given random variable $\mathbf{Z}$ and the transformation T to depend on some parameters, the initial state $\mathbf{x}_0$ and time t in our setting, we can derive a conditional normalizing flow that maps the current state $\mathbf{X}_t$ of (2.1) to a base random variable $\mathbf{Z}_t(\mathbf{x}_0)$ conditioned on an initial state $\mathbf{X}_0 = \mathbf{x}_0$, i.e.,

$$(3.2) \quad \begin{aligned} \mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0 &= T^{-1}(\mathbf{Z}_t(\mathbf{x}_0); \theta(\mathbf{x}_0, t)), \\ p_{\mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0}(\mathbf{x}) &= p_{\mathbf{Z}_t(\mathbf{x}_0)}(T(\mathbf{x}; \theta(\mathbf{x}_0, t))) \left| \det \nabla_{\mathbf{x}} T(\mathbf{x}; \theta(\mathbf{x}_0, t)) \right|, \end{aligned}$$

where $\theta(\mathbf{x}_0, t)$ represents the parameter set of the invertible transformation T conditioned on the initial state $\mathbf{x}_0$ and time t . This transformation T maps an unknown stochastic process $\mathbf{X}_t$ to another stochastic process $\mathbf{Z}_t$ given the initial state $\mathbf{X}_0 = \mathbf{x}_0$.

3.1. What makes a good base stochastic process for the conditional normalizing flow? The transition PDF provides a way to approximate the solution of the FPE under different initial conditions. However, directly approximating the solution of the FPE with a Dirac delta function as the initial condition is challenging. The challenge lies in the fact that the Dirac delta function lacks regularity, and there is no invertible transformation that maps it to a regular distribution (e.g., normal or uniform distributions) used in normalizing flows. To address this challenge, we linearize the original SDE and use the solution of the resulting linear SDE as the base stochastic process of the normalizing flow. The linearization of the original SDE is based on the assumption that the drift term $\mathbf{f}$ and diffusion term $\mathbf{g}$ are smooth functions. This allows us to approximate the original SDE by a linear SDE, which admits an analytical solution. Specifically, we approximate the original drift term by the first-order Taylor expansion and the diffusion term by the zero-order Taylor expansion at the initial state $\mathbf{x}_0$, i.e.,

$$(3.3) \quad \begin{cases} d\widehat{\mathbf{X}}_t = (\mathbf{f}(\mathbf{x}_0) + \nabla \mathbf{f}(\mathbf{x}_0)(\widehat{\mathbf{X}}_t - \mathbf{x}_0))dt + \mathbf{g}(\mathbf{x}_0)d\mathbf{W}_t, & t > 0, \\ \widehat{\mathbf{X}}_0 = \mathbf{x}_0. \end{cases}$$

The corresponding solution admits an analytical form,

$$(3.4) \quad \widehat{\mathbf{X}}_t = \mathbf{x}_0 + \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{f}(\mathbf{x}_0) ds + \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{g}(\mathbf{x}_0) d\mathbf{W}_s,$$

which is a multivariate Gaussian process with mean and covariance given by

$$(3.5) \quad \begin{aligned} \mathbb{E}[\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0] &= \mathbf{x}_0 + \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{f}(\mathbf{x}_0) ds, \\ \Sigma[\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0] &= \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{g}(\mathbf{x}_0) \mathbf{g}(\mathbf{x}_0)^\top e^{\nabla \mathbf{f}(\mathbf{x}_0)^\top (t-s)} ds. \end{aligned}$$

The two integrals with respect to time t in equation (3.5) can be effectively approximated by Gauss-Legendre quadrature, thus,

$$(3.6) \quad \begin{aligned} \mathbb{E} [\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0] &\approx \mathbf{x}_0 + \sum_{i=1}^m w_i e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s_i)} \mathbf{f}(\mathbf{x}_0), \\ \Sigma [\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0] &\approx \sum_{i=1}^m w_i e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s_i)} \mathbf{g}(\mathbf{x}_0) \mathbf{g}(\mathbf{x}_0)^\top e^{\nabla \mathbf{f}(\mathbf{x}_0)^\top (t-s_i)}, \\ w_i &= \frac{t}{2} \hat{w}_i, \quad s_i = \frac{t}{2} (\hat{s}_i + 1), \end{aligned}$$

where $\{(\hat{w}_i, \hat{s}_i)\}$ are the weights and nodes associated with the Legendre polynomial defined on $[-1, 1]$. Moreover, if $\nabla \mathbf{f}(\mathbf{x}_0)$ is invertible, we can rewrite the solution as

$$(3.7) \quad \widehat{\mathbf{X}}_t = \mathbf{x}_0 + (\nabla \mathbf{f}(\mathbf{x}_0))^{-1} \left(e^{\nabla \mathbf{f}(\mathbf{x}_0)t} - \mathbf{I} \right) \mathbf{f}(\mathbf{x}_0) + \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{g}(\mathbf{x}_0) d\mathbf{W}_s.$$

The corresponding conditional mean admits the equivalent expression,

$$\mathbb{E} [\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0] = \mathbf{x}_0 + (\nabla \mathbf{f}(\mathbf{x}_0))^{-1} \left(e^{\nabla \mathbf{f}(\mathbf{x}_0)t} - \mathbf{I} \right) \mathbf{f}(\mathbf{x}_0).$$

Thus, instead of approximating the integral, one may compute the inverse of $\nabla \mathbf{f}(\mathbf{x}_0)$.

We note that the linear SDE system (3.3) provides a close pathwise approximation to the original SDE (2.1) for small t . The corresponding transition PDF associated with the linear SDE provides a good approximation of the target transition PDF in the short-time regime, as verified by Proposition 3.1 below.

PROPOSITION 3.1. *Suppose that the drift term $\mathbf{f} = \mathbf{f}_\nu$ in the SDE (2.1) is differentiable, and the diffusion term $\mathbf{g}$ is constant, $\mathbf{G} = (G^{ij})_{i,j \leq d} = \frac{1}{2} \mathbf{g} \mathbf{g}^\top$ is non-degenerate, $\{p_{\nu_t}\}_{t \in [0, t_f]}$ is the PDF solution on the time interval $[0, t_f]$ to the corresponding FPE with a Dirac delta function $\delta_{\mathbf{x}_0}$ as the initial condition, $\{p_{\sigma_t}\}_{t \in [0, t_f]}$ is the PDF solution to the FPE corresponding to the linearized SDE (3.3) with drift term $\mathbf{f}_\sigma$ and the same initial condition $\delta_{\mathbf{x}_0}$. Denote ν_t and σ_t the probability measures corresponding to p_{ν_t} , p_{σ_t} , respectively. If $\|\nabla \mathbf{f}_\nu(\mathbf{x})\| \leq C_{\mathbf{f}}$, $\forall \mathbf{x} \in \mathbb{R}^d$ for some $C_{\mathbf{f}} \geq 0$, then for any $t \in (0, t_f)$,*

$$\begin{aligned} H(\sigma_t | \nu_t) &\leq \frac{1}{2} \int_0^t \int_{\mathbb{R}^d} \left| \mathbf{G}^{-1/2} (\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})) \right|^2 d\sigma_s ds, \\ \|\nu_t - \sigma_t\|_{\text{TV}}^2 &\leq \frac{1}{4} \int_0^t \int_{\mathbb{R}^d} \left| \mathbf{G}^{-1/2} (\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})) \right|^2 d\sigma_s ds. \end{aligned}$$

Both the relative entropy $H(\sigma_t | \nu_t)$ and the square of total variation $\|\nu_t - \sigma_t\|_{\text{TV}}^2$ are $\mathcal{O}(t^2)$ as $t \rightarrow 0$. Moreover, if $\|\nabla^2 \mathbf{f}_\nu(\mathbf{x})\| \leq H_{\mathbf{f}}$, $\forall \mathbf{x} \in \mathbb{R}^d$ for some $H_{\mathbf{f}} \geq 0$, then $H(\sigma_t | \nu_t)$, $\|\nu_t - \sigma_t\|_{\text{TV}}^2$ are both $\mathcal{O}(t^3)$ as $t \rightarrow 0$.

Proof. See Appendix A.

Statistical quantities evaluated under the target measure can be approximated by those under the approximate measure, with an error bounded by the total variation distance between the two measures, as established in the next proposition.

PROPOSITION 3.2. *Suppose that ν, σ are two probability measures on $\mathbb{R}^d$. Then for every function $\mathbf{f} \in L^2(\mathbb{R}^d, \nu; \mathbb{R}^m) \cap L^2(\mathbb{R}^d, \sigma; \mathbb{R}^m)$, we have*

$$\left| \int_{\mathbb{R}^d} \mathbf{f}(\mathbf{x}) d\nu(\mathbf{x}) - \int_{\mathbb{R}^d} \mathbf{f}(\mathbf{x}) d\sigma(\mathbf{x}) \right| \leq \sqrt{2\|\nu - \sigma\|_{\text{TV}}} (\|\mathbf{f}\|_{L^2(\mathbb{R}^d, \nu; \mathbb{R}^m)} + \|\mathbf{f}\|_{L^2(\mathbb{R}^d, \sigma; \mathbb{R}^m)}).$$

Proof. See Appendix B.

In particular, if we take $\sigma = T_{\#}\nu$ for some measurable map $T : \mathbb{R}^d \rightarrow \mathbb{R}^d$ and $f(\mathbf{x}) = \mathbf{x}$, assuming ν, σ have finite second moments, we have

$$\left| \int_{\mathbb{R}^d} \mathbf{x} d\nu(\mathbf{x}) - \int_{\mathbb{R}^d} T(\mathbf{x}) d\nu(\mathbf{x}) \right| \leq \sqrt{2\|\nu - \sigma\|_{\text{TV}}} (\|\mathbf{x}\|_{L^2(\mathbb{R}^d, \nu; \mathbb{R}^d)} + \|\mathbf{x}\|_{L^2(\mathbb{R}^d, \sigma; \mathbb{R}^d)}).$$

We also state that if $\nu_i \rightarrow \mu$ in total variation and their second moments converge, then the Knothe-Rosenblatt (KR) map T_i satisfying $(T_i)_{\#}\mu = \nu_i$ converges to the identity map in $L_2(\mu)$.

PROPOSITION 3.3. *Let $\{\nu_i\}_{i=1}^{\infty}$ be a family of absolutely continuous probability measures on $\mathbb{R}^d$ w.r.t. the Lebesgue measure, each with finite second moments, and assume that μ is an absolutely continuous measure, $\|\nu_i - \mu\|_{\text{TV}} \rightarrow 0$, $\|\mathbf{x}\|_{L^2(\mathbb{R}^d, \nu_i)} \rightarrow \|\mathbf{x}\|_{L^2(\mathbb{R}^d, \mu)}$ as $i \rightarrow \infty$. Then, let T_i be the KR map between μ and ν_i . We have $T_i \rightarrow \text{Id}$ in $L^2(\mu)$ as $i \rightarrow \infty$.*

Proof. See Appendix C.

Proposition 3.1 shows that, for small t , the transition PDF of the linearized SDE is close to that of the original SDE in total variation. Proposition 3.3 implies that, under certain conditions, KR maps associated with a sequence of measures and its limiting measure converge to the identity map. This observation motivates our design choice: the conditional normalizing flow is constrained to be the identity at $t = 0$ and to deform gradually as t increases. In this way, we avoid directly mapping from a Dirac delta initial condition and reduce the training difficulty near $t = 0$.

3.2. Architecture of the conditional normalizing flow. The transition-density-based method reformulates the original problem of finding $p(\mathbf{x}, t | \mathbf{x}_0)$ as the task of identifying an invertible transformation $T(\cdot; \theta(\mathbf{x}_0, t))$, mapping a stochastic process $\mathbf{X}_t$ to $\mathbf{Z}_t$ given initial state $\mathbf{X}_0 = \mathbf{x}_0$. In practice, we can construct such a complex invertible mapping by stacking a sequence of simple bijections, i.e.,

$$\mathbf{Z}_t | \mathbf{Z}_0 = \mathbf{z}_0(\mathbf{x}_0) = T(\mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0; \theta(\mathbf{x}_0, t)) = T_{[K]} \circ T_{[K-1]} \circ \cdots \circ T_{[1]}(\mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0),$$

where $T_{[i]} = T_{[i]}(\cdot; \theta_{[i]}(\mathbf{x}_0, t))$ is based on shallow neural networks. The inverse and Jacobian determinant are given as

$$\begin{aligned} \mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0 &= T^{-1}(\mathbf{Z}_t | \mathbf{Z}_0 = \mathbf{z}_0; \theta(\mathbf{x}_0, t)) = T_{[1]}^{-1} \circ \cdots \circ T_{[K-1]}^{-1} \circ T_{[K]}^{-1}(\mathbf{Z}_t | \mathbf{Z}_0 = \mathbf{z}_0(\mathbf{x}_0)), \\ |\det(\nabla_{\mathbf{x}} T(\cdot; \theta(\mathbf{x}_0, t)))| &= \prod_{i=1}^L |\det(\nabla_{\mathbf{x}_{[i-1]}} T_{[i]}(\cdot))|, \end{aligned}$$

which motivates the use of techniques that enable efficient computation of both the inverse and the Jacobian matrix of $T_{[i]}(\cdot)$. In this work, we adopt the structure with descending active transformation dimensions in the standard KRnet [38], which is inspired by the KR rearrangement [34] and affine coupling layers [10], to enhance accuracy and reduce model complexity in high-dimensional settings.

Let $\mathbf{x}^{\top} = ((\mathbf{x}^{(1)})^{\top}, (\mathbf{x}^{(2)})^{\top}, \dots, (\mathbf{x}^{(K)})^{\top})$ be a partition of $\mathbf{x}^{\top}$, where $\mathbf{x}^{(i)} = (x_1^{(i)}, \dots, x_{d_i}^{(i)})^{\top}$ with $1 \leq K \leq d$, $1 \leq d_i \leq d$ and $\sum_{i=1}^K d_i = d$. Denote $\mathbf{x}^{(i:j)} = ((\mathbf{x}^{(i)})^{\top}, (\mathbf{x}^{(i+1)})^{\top}, \dots, (\mathbf{x}^{(j)})^{\top})^{\top}$. Our conditional normalizing flow $T_{\text{C-KRnet}}$ is

defined as follows,

$$T_{[1]} = \tilde{T}_K, \quad T_{[k]} = \begin{pmatrix} \tilde{T}_{K+1-k} \\ \text{Id}_{K+2-k:K} \end{pmatrix}, \quad k = 2, \dots, K,$$

$$\mathbf{Z}_t | \mathbf{Z}_0 = \mathbf{z}_0 = T_{\text{C-KRnet}}(\mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0; \theta(\mathbf{x}_0, t)) = \begin{pmatrix} \tilde{T}_1 \\ \text{Id}_{(2:K)} \end{pmatrix} \circ \dots \circ \begin{pmatrix} \tilde{T}_{K-1} \\ \text{Id}_{(K)} \end{pmatrix} \circ \tilde{T}_K(\mathbf{X}_t | \mathbf{X}_0 = \mathbf{x}_0),$$

where $\text{Id}_{(i:j)}(\mathbf{x}) = \mathbf{x}^{(i:j)}$ denotes the projection operator, $\tilde{T}_k : [-1, 1]^{\sum_{i=1}^k d_i} \rightarrow [-1, 1]^{\sum_{i=1}^k d_i}$ is an invertible mapping of $\mathbf{x}^{(1:k)}$ that is constructed by the composition of conditional affine coupling layers defined in the Section 3.2.1.

The conditional normalizing flow $T_{\text{C-KRnet}}$ is primarily constructed using two nested loops: an outer loop and an inner loop. The outer loop consists of K stages, each corresponding to a mapping $\tilde{T}_k$. The inner loop $\tilde{T}_k$ contains l_k stages, representing the number of coupling layers. Following the application of $\tilde{T}_k$, the k th partition is kept fixed in the subsequent outer stages. As the outer loop progresses, the active dimension gradually decreases, which motivates reducing the depth l_k for later stages.

3.2.1. Conditional affine coupling layer. In this section, we focus on the construction of the inner loop $\tilde{T}_k = \tilde{T}_{k,l_k} \circ \dots \circ \tilde{T}_{k,2} \circ \tilde{T}_{k,1}$. To this end, we first introduce a conditional affine coupling layer, given some information $\boldsymbol{\xi}$ and time t , applied to a vector $\mathbf{y} = (\mathbf{y}_1, \mathbf{y}_2) \in \mathbb{R}^m$, where $\mathbf{y}_1 \in \mathbb{R}^{m_1}$, and $\mathbf{y}_2 \in \mathbb{R}^{m-m_1}$ as follows:

$$(3.8) \quad \begin{cases} \hat{\mathbf{y}}_1 = \mathbf{y}_1, \\ \hat{\mathbf{y}}_2 = \mathbf{y}_2 \odot (\mathbf{1}_{m-m_1} + \beta \tanh(t \cdot \mathbf{s})) + e^{\boldsymbol{\zeta}} \odot \tanh(t \cdot \mathbf{q}), \end{cases}$$

where $\beta \in (0, 1)$ is user defined, $\boldsymbol{\zeta}$ is trainable and $\mathbf{s}$ and $\mathbf{q}$ are outputs of a neural network that takes as input the unchanged portion $\mathbf{y}_1$, the given information $\boldsymbol{\xi}$ and time t , i.e.,

$$(3.9) \quad (\mathbf{s}, \mathbf{q}) = \text{NN}(\mathbf{y}_1, \boldsymbol{\xi}, t).$$

To emphasize the influence of the initial state, we adopt the approach proposed in [18], applying a random Fourier feature transformation [41] before the fully connected layers of the neural network used in equation (3.9). The resulting neural network is

$$(3.10) \quad \begin{aligned} \mathbf{h}_0 &= [\mathbf{y}_1, \boldsymbol{\xi}, t]^\top, \quad \mathbf{h}_1 = \left[\sin\left(\frac{1}{e^\gamma} \mathbf{F} \mathbf{h}_0 + \mathbf{b}_0\right), \cos\left(\frac{1}{e^\gamma} \mathbf{F} \mathbf{h}_0 + \mathbf{b}_0\right), \mathbf{h}_0 \right]^\top, \\ \mathbf{h}_j &= \text{SiLU}(\mathbf{W}_{j-1} \mathbf{h}_{j-1} + \mathbf{b}_{j-1}), \text{ for } j = 2, \dots, M, \quad (\mathbf{s}, \mathbf{q}) = \mathbf{W}_M \mathbf{h}_M + \mathbf{b}_M, \end{aligned}$$

where $\mathbf{F} \in \mathbb{R}^{r_h/2 \times \dim(\mathbf{h}_0)}$, $\mathbf{b}_0 \in \mathbb{R}^{r_h/2}$ are sampled from a standard normal distribution and a uniform distribution over $[0, 2\pi]^{r_h/2}$, respectively and r_h denotes the dimension of the random features. Both $\mathbf{F}$ and $\mathbf{b}_0$ are fixed after initialization. $\gamma, \{\mathbf{W}_j, \mathbf{b}_j\}_{j=1}^M$ are trainable parameters, and the sigmoid linear unit (SiLU) function is used as the activation function.

Since the conditional affine coupling layer (3.8) updates only $\mathbf{y}_2$, we alternate the roles of $\mathbf{y}_1$ and $\mathbf{y}_2$ in subsequent layers to ensure a full update of $\mathbf{y}$. Note that the transformation defined by the conditional affine coupling layer (3.8) reduces to the identity when $t = 0$, which is consistent with using the PDF corresponding to the linearized SDE as the base distribution of the conditional normalizing flow.

Now we are ready to construct $\tilde{T}_k$, which consists of l_k conditional affine coupling layers. Denote $\mathbf{x}_{[0]} = \mathbf{x}$, $\mathbf{x}_{[k]} = T_{[k]}(\mathbf{x}_{[k-1]})$, $\boldsymbol{\xi} = \mathbf{x}_0$, and $\mathbf{y}_{[k]} = \mathbf{x}_{[k]}^{(1:K-k)}$ the active part of $\mathbf{x}_{[k]}$ for $k = 1, \dots, K-1$. We apply a sequence of conditional affine coupling transformations (3.8) to $\mathbf{y}_{[k]}$ using a half-half split.

4. Deep adaptive algorithm. Given an SDE, we solve the corresponding transition FPE with Dirac delta functions centered at various points as initial conditions, which is considered as offline training. Let $p(\mathbf{x}, t|\mathbf{x}_0)$ be a conditional probability density function defined on $\mathbb{R}^d \times [0, t_f] \times \Omega_0$. It satisfies the transition FPE (2.6). By using the solution of the linearized SDE as the base stochastic process and enforcing the transformation to be the identity mapping at $t = 0$, the initial condition is naturally satisfied. Applying PINNs yields

$$(4.1) \quad \mathcal{L}(\boldsymbol{\theta}) = \int_{\mathbb{R}^d \times (0, t_f] \times \Omega_0} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 d\rho(\mathbf{x}, \mathbf{x}_0, t), \quad r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}}) = \frac{\partial p_{\boldsymbol{\theta}}}{\partial t} - \mathcal{L}_{\mathbf{f}, \mathbf{g}}^*[p_{\boldsymbol{\theta}}],$$

where $\rho(\mathbf{x}, \mathbf{x}_0, t) = \rho(\mathbf{x}|\mathbf{x}_0, t)\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$ is the sampling measure to be used in the training process. For the case where $\mathbf{x}_0$ lies within a bounded domain Ω_0 and t_f is finite, we define $\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$ to be the uniform measure over $\Omega_0 \times [0, t_f]$. For the remaining term, $\rho(\mathbf{x}|\mathbf{x}_0, t)$, we employ an adaptive sampling strategy similar to that used for specific FPEs in our previous work [13, 49]. The only difference is that, instead of sampling solely from the uniform distribution at the start, we are motivated to sample from the base stochastic process, as it approximates the target transition PDF particularly well when the time t is small. After some training epochs, we obtain a conditional normalizing flow that approximates the target solution and can be used to generate samples from it. These samples are then used to update the collocation points. The procedure offers an adaptive strategy for designing the sampling measure $\rho_k(\mathbf{x}|\mathbf{x}_0, t)$ at stage k , by combining the uniform measure, the measure from the previous stage, and the one induced by the current normalizing flow, i.e.,

$$(4.2) \quad \begin{aligned} \rho_0(\mathbf{x}|\mathbf{x}_0, t) &= \frac{\gamma_1}{\gamma_1 + \gamma_3} \mu(\mathbf{x}) + \frac{\gamma_3}{\gamma_1 + \gamma_3} \rho_{\mathbf{Z}}(\mathbf{x}|\mathbf{x}_0, t), \quad \gamma_1 + \gamma_2 + \gamma_3 = 1, \\ \rho_{k+1}(\mathbf{x}|\mathbf{x}_0, t) &= \gamma_1 \mu(\mathbf{x}) + \gamma_2 \rho_k(\mathbf{x}|\mathbf{x}_0, t) + \gamma_3 \rho_{\text{C-KRnet}, \boldsymbol{\theta}}(\mathbf{x}|\mathbf{x}_0, t), \end{aligned}$$

where μ denotes the uniform measure over a user-defined bounded domain, $\rho_{\mathbf{Z}}(\mathbf{x}, t|\mathbf{x}_0)$ denotes the measure induced by the base random process, $\rho_{\text{C-KRnet}, \boldsymbol{\theta}}(\mathbf{x}, t|\mathbf{x}_0)$ denotes the measure induced by the current conditional normalizing flow and $\gamma_1, \gamma_2, \gamma_3$ are hyperparameters used to tune the sampling measure.

The transition PDF of the base stochastic process serves as a good approximation of the target transition PDF when t is small. However, as $t \rightarrow 0$, the target solution approaches a Dirac delta function, and even small perturbations may lead to the divergence of the residual. To illustrate this, we consider a small perturbation around the transition PDF of the linearized SDE and analyze the asymptotic behavior of the corresponding residual as $t \rightarrow 0$.

PROPOSITION 4.1. *Suppose the drift term $\mathbf{f}$ in the SDE (2.1) is differentiable, and the diffusion term $\mathbf{g}$ is constant, $\mathbf{G} = (G^{ij})_{i,j \leq d} = \frac{1}{2} \mathbf{g} \mathbf{g}^\top$ is non-degenerate. Assume that $\|\nabla \mathbf{f}(\mathbf{x})\| \leq C_{\mathbf{f}}$, $\forall \mathbf{x} \in \mathbb{R}^d$ for some $C_{\mathbf{f}} \geq 0$. Given an initial state $\mathbf{x}_0$ and time t , let $\tilde{p}(\mathbf{x}, t|\mathbf{x}_0) = \mathcal{N}(\mathbf{x}; \mathbf{m}(t), \boldsymbol{\Sigma}(t))$ where $\boldsymbol{\Sigma}(t) \sim t\mathbf{G}$, and $\frac{d}{dt} \mathbf{m}(t)$ is*

bounded. Then,

$$(4.3) \quad \int_{\mathbb{R}^d} |r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2 d\mathbf{x} = \mathcal{O}(t^{-(\frac{d}{2}+2)}),$$

$$(4.4) \quad \mathbb{E}_{\mathbf{x} \sim \tilde{p}} [|r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2] = \int_{\mathbb{R}^d} |r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2 \tilde{p}(\mathbf{x}, t | \mathbf{x}_0) d\mathbf{x} = \mathcal{O}(t^{-(d+2)}).$$

Proof. See Appendix D.

According to Proposition 4.1, although the approximation may be close to the target transition PDF for small t , the integral of the residual over the spatial domain can scale as $t^{-(d+2)}$ or $t^{-(\frac{d}{2}+2)}$ as $t \rightarrow 0$ depending on the sampling measure $\rho(\mathbf{x} | \mathbf{x}_0, t)$, which potentially leads to a blow-up. This motivates us to define a time-weighted loss function. Incorporating the adaptive sampling strategy, we rewrite the loss function at the k -th stage as a sum of three parts:

$$\begin{aligned} \mathcal{L}^k(\boldsymbol{\theta}) = & \gamma_1 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 d\mu(\mathbf{x}) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0) \\ & + \gamma_2 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 d\rho_{k-1}(\mathbf{x} | \mathbf{x}_0, t) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0) \\ & + \gamma_3 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 d\rho_{\text{C-KRnet}, \boldsymbol{\theta}}(\mathbf{x}, t | \mathbf{x}_0) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0). \end{aligned}$$

We apply weighting functions $w_1(t) = t^{\frac{d}{2}+2}$ to the first part, and $w_2(t) = t^{d+2}$ to the last two parts. The resulting weighted loss function is then defined as

$$\begin{aligned} \mathcal{L}_w^k(\boldsymbol{\theta}) = & \gamma_1 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 t^{(\frac{d}{2}+2)} d\mu(\mathbf{x}) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0) \\ & + \gamma_2 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 t^{(d+2)} d\rho_{k-1}(\mathbf{x} | \mathbf{x}_0, t) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0) \\ & + \gamma_3 \int_{\Omega \times \Omega_0 \times (0, t_f]} |r(\mathbf{x}, \mathbf{x}_0, t; p_{\boldsymbol{\theta}})|^2 t^{(d+2)} d\rho_{\text{C-KRnet}}(\mathbf{x} | \mathbf{x}_0, t) d\rho(t | \mathbf{x}_0) d\rho(\mathbf{x}_0). \end{aligned}$$

We emphasize that the usage of the linearized SDE respects the causality, and the weighting functions introduced here not only help prevent blow-up but also assign greater weight to more challenging training regions, making our method well-suited for the transition FPE.

We introduce an indicator variable to distinguish samples drawn from different distributions. Specifically, let $N = N_1 + N_2 + N_3$ with proportions $N_1 : N_2 : N_3 = \gamma_1 : \gamma_2 : \gamma_3$. Define $S_1 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, \eta^i = 0)\}_{i=1}^{N_1}$, where $(\mathbf{x}^i, \mathbf{x}_0^i, t^i)$ are sampled from $\mu(\mathbf{x})\rho(t | \mathbf{x}_0)\rho(\mathbf{x}_0)$. Similarly, define $S_2 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, \eta^i = 1)\}_{i=N_1+1}^{N_1+N_2}$ where $\{(\mathbf{x}^i, \mathbf{x}_0^i, t^i)\}$ are drawn from $\rho_{k-1}(\mathbf{x} | \mathbf{x}_0, t)\rho(t | \mathbf{x}_0)\rho(\mathbf{x}_0)$. Define $S_3 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, \eta^i = 1)\}_{i=N_1+N_2+1}^N$ where $\{(\mathbf{x}^i, \mathbf{x}_0^i, t^i)\}$ are sampled from $\rho_{\text{C-KRnet}}(\mathbf{x} | \mathbf{x}_0, t)\rho(t | \mathbf{x}_0)\rho(\mathbf{x}_0)$. Let $S = S_1 \cup S_2 \cup S_3$. The empirical time-weighted loss function is then defined as:

$$(4.5) \quad \widehat{\mathcal{L}}_w^k(\boldsymbol{\theta}; S) := \frac{1}{N} \sum_{i=1}^N |r(\mathbf{x}^i, \mathbf{x}_0^i, t^i; p_{\boldsymbol{\theta}})|^2 \left(\eta^i t^{d+2} + (1 - \eta^i) t^{\frac{d}{2}+2} \right).$$

The optimal parameter $\boldsymbol{\theta}^*$ can be obtained via solving the optimization problem:

$$(4.6) \quad \boldsymbol{\theta}^* \in \arg \min_{\boldsymbol{\theta}} \widehat{\mathcal{L}}_w(\boldsymbol{\theta}; \tilde{S}).$$

The adaptive training process is summarized in Algorithm 4.1. The training set is divided into several mini-batches, and the Adam optimizer is used to update the parameters θ . The training set S is updated according to the strategy (4.2).

Algorithm 4.1 Solving the transition FPE with various initial states

Input: maximum epoch number N_e , maximum iteration number N_{adaptive} , rate $\gamma_1, \gamma_2, \gamma_3$, initial learning rate l_r , decay rate η , step size n_s , the number of the training samples N .

for $k = 0, \dots, N_{\text{adaptive}}$ **do**

if $k = 0$ **then**

$N_1 = \lfloor \frac{\gamma_1}{\gamma_1 + \gamma_3} * N \rfloor$, $N_3 = N - N_1$,

 Generate $S_1 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, 0)\}_{i=1}^{N_1}$ with $(\mathbf{x}^i, \mathbf{x}_0^i, t^i) \sim \mu(\mathbf{x})\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$,

$S_3 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, 1)\}_{i=N_1+1}^N$ with $(\mathbf{x}^i, \mathbf{x}_0^i, t^i) \sim \rho_{\text{C-KRnet}, \theta}(\mathbf{x}, t|\mathbf{x}_0)\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$.

 Initialize the training set $S = S_1 \cup S_3$.

else

$N_1 = \lfloor \gamma_1 * N \rfloor$, $N_2 = \lfloor \gamma_2 * N \rfloor$, $N_3 = N - N_1 - N_2$,

 Generate $S_1 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, 0)\}_{i=1}^{N_1}$ with $(\mathbf{x}^i, \mathbf{x}_0^i, t^i) \sim \mu(\mathbf{x})\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$,

S_2 consists of N_2 random samples drawn from S .

$S_3 = \{(\mathbf{x}^i, \mathbf{x}_0^i, t^i, 1)\}_{i=N_1+N_2+1}^N$, $(\mathbf{x}^i, \mathbf{x}_0^i, t^i) \sim \rho_{\text{C-KRnet}, \theta}(\mathbf{x}, t|\mathbf{x}_0)\rho(t|\mathbf{x}_0)\rho(\mathbf{x}_0)$.

 Update the training set: $S = S_1 \cup S_2 \cup S_3$.

end if

for $j = 1, \dots, N_e$ **do**

if $((k-1)N_e + j) \% n_s == 0$ **then**

$l_r = \eta * l_r$.

end if

 Divide S into n mini-batches $\{S^{ib}\}_{ib=1}^n$ randomly.

for $ib = 1, \dots, n$ **do**

 Compute the loss function (4.5) $\hat{\mathcal{L}}_w(\theta; S^{ib})$.

 Update θ by using the Adam optimizer.

end for

end for

end for

Output: The predicted solution $p_{\text{C-KRnet}, \theta}(\mathbf{x}, t|\mathbf{x}_0)$.

5. Importance sampling-based approximation of $p(\cdot, t)$. Once the transition PDF is available, a natural way to approximate the solution of equation (2.5) is by approximating the integral (2.7) through the Monte Carlo method,

$$(5.1) \quad p(\mathbf{x}, t) \approx \frac{1}{M} \sum_{i=1}^M p(\mathbf{x}, t|\mathbf{x}_0^i), \quad \mathbf{x}_0^i \sim p_0(\mathbf{x}),$$

where $\{\mathbf{x}_0^i\}_{i=1}^M$ are samples from the initial distribution p_0 . However, the transition PDF approximates the Dirac delta function as $t \rightarrow 0$. Consequently, a direct Monte Carlo method may incur large errors due to high variance. Alternatively, we apply importance sampling, leveraging the knowledge of the transition PDF.

$$(5.2) \quad \begin{aligned} p(\mathbf{x}, t) &= \int_{\Omega_0} \frac{p(\mathbf{x}, t|\mathbf{x}_0)p_0(\mathbf{x}_0)}{q(\mathbf{x}_0|\mathbf{x}, t)} q(\mathbf{x}_0|\mathbf{x}, t) d\mathbf{x}_0 \\ &\approx \hat{p}_\theta(\mathbf{x}, t) := \frac{1}{M} \sum_{i=1}^M \frac{p(\mathbf{x}, t|\mathbf{x}_0^i)p_0(\mathbf{x}_0^i)}{q(\mathbf{x}_0^i|\mathbf{x}, t)}, \quad \mathbf{x}_0^i \sim q(\cdot|\mathbf{x}, t). \end{aligned}$$

Recall that given an initial state $\mathbf{x}_0$ and time t , $\mathbf{x}$ concentrates around the initial state $\mathbf{x}_0$ for small time. The corresponding base distribution used in our conditional normalizing flow is

$$p_{\widehat{\mathbf{X}}_t}(\mathbf{x}, t | \mathbf{x}_0) = \mathcal{N}\left(\mathbf{x}; \mathbf{x}_0 + \int_0^t e^{\nabla \mathbf{f}(\mathbf{x}_0)(t-s)} \mathbf{f}(\mathbf{x}_0) ds, \Sigma[\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0]\right).$$

Let

$$(5.3) \quad q_1(\mathbf{x}_0 | \mathbf{x}, t) = \mathcal{N}(\mathbf{x}_0; m(\mathbf{x}, t), \tilde{\Sigma}(\mathbf{x}, t)),$$

where $m(\mathbf{x}, t) = \mathbf{x} - \int_0^t e^{\nabla \mathbf{f}(\mathbf{x})s} \mathbf{f}(\mathbf{x}) ds$, $\tilde{\Sigma}(\mathbf{x}, t) = \int_0^t e^{\nabla \mathbf{f}(\mathbf{x})s} \mathbf{g}(\mathbf{x}) \mathbf{g}(\mathbf{x})^\top e^{\nabla \mathbf{f}(\mathbf{x})^\top s} ds$. The distribution q_1 is primarily designed to capture the local concentration of the transition PDF and to mitigate the effect of large variance when t is small. However, the approximation accuracy of q_1 deteriorates as t increases. As t grows, the transition PDF becomes increasingly smooth and therefore does not hinder the approximation of $p(\mathbf{x}, t)$. This motivates incorporating the initial distribution into the proposal. Accordingly, we propose to use a mixture of $q_1(\cdot | \mathbf{x}, t)$ and $p_0(\cdot)$ as the proposal distribution in the importance sampling scheme. Define

$$(5.4) \quad q_2(\mathbf{x}_0 | \mathbf{x}, t) = \alpha(t) q_1(\mathbf{x}_0 | \mathbf{x}, t) + (1 - \alpha(t)) p_0(\mathbf{x}_0),$$

where $\alpha(t) \in (0, 1)$ is decreasing with respect to t , e.g., $\alpha(t) = \exp(-at)$, $a > 0$. The set of samples $\{\mathbf{x}_0^i\}_{i=1}^M \sim q_2(\mathbf{x}_0 | \mathbf{x}, t)$ can be generated by drawing $\alpha(t)M$ samples from q_1 and $(1 - \alpha(t))M$ samples from p_0 .

6. Numerical experiments. In this section, we present three numerical experiments to demonstrate the effectiveness of the proposed approach. The conditional normalizing flow is applied to solve the transition FPE under various initial states. Once obtained, the importance sampling method is used to evaluate $p(\mathbf{x}, t)$ given an initial distribution. We first consider a benchmark problem where the analytical solution of the transition PDF is available. We next apply our method to an SDE with nonlinear drift and constant diffusion. To further demonstrate the robustness, we investigate the case of state-dependent diffusion.

The experiments are implemented in PyTorch with an initial learning rate of 0.001. The number of quadrature points in the equation (3.6) is set to 10. To quantify the effectiveness of the proposed method, we consider several metrics. For some time t , we compute the relative L^2 error of the approximation for the transition PDF $p(\mathbf{x}, t | \mathbf{x}_0)$ if the exact transition PDF is known via

$$(6.1) \quad \text{Rel}(p_{\text{C-KRnet}}(\cdot, t | \cdot)) = \sqrt{\frac{\sum_{i=1}^{N_v} (p(\mathbf{x}^i, t | \mathbf{x}_0^i) - p_{\text{C-KRnet}}(\mathbf{x}^i, t | \mathbf{x}_0^i))^2}{\sum_{i=1}^{N_v} p(\mathbf{x}^i, t | \mathbf{x}_0^i)^2}},$$

where $\{(\mathbf{x}_0^i, \mathbf{x}^i)\}_{i=1}^{N_v} \subset \Omega_0 \times \Omega$ denotes the validation dataset. When approximating the solution of the FPE given an initial distribution $\rho_0(\mathbf{x})$, we consider the relative L^2 error of the approximation $\hat{p}_\theta(\mathbf{x}, t)$ of $p(\mathbf{x}, t)$, as well as the maximum mean discrepancy (MMD) [17]. The relative L^2 error reads

$$(6.2) \quad \text{Rel}(\hat{p}_\theta(\cdot, t)) = \sqrt{\frac{\sum_{i=1}^{N_v} (p(\mathbf{x}^i, t) - \hat{p}_\theta(\mathbf{x}^i, t))^2}{\sum_{i=1}^{N_v} p(\mathbf{x}^i, t)^2}},$$

where $\{\mathbf{x}^i\}_{i=1}^{N_v} \subset \Omega$ denotes the validation dataset. When the analytical form of the transition PDF is available, we employ Gauss–Kronrod quadrature rule to obtain the reference solution $p(\mathbf{x}, t)$ for a given initial distribution. In the absence of an analytical solution, the alternating-direction implicit (ADI) finite difference scheme [32] with a time step of 0.001 and a spatial discretization size of 0.01 is used to compute the reference solution for cases where the initial distributions are continuous over the entire domain $\mathbb{R}^2$. An MMD is introduced in particular to assess the effectiveness of the proposed method for discontinuous initial distributions.

$$(6.3) \quad \text{MMD}^2(t) = \frac{1}{M_1^2} \sum_{i,j=1}^{M_1} K(\mathbf{x}_t^i, \mathbf{x}_t^j) - \frac{2}{M_1 M_2} \sum_{i,j=1}^{M_1, M_2} K(\mathbf{x}_t^i, \mathbf{y}_t^j) + \frac{1}{M_2^2} \sum_{i,j=1}^{M_2} K(\mathbf{y}_t^i, \mathbf{y}_t^j),$$

$$K(\mathbf{x}, \mathbf{y}) = \frac{1}{3} \exp\left(-\frac{4\|\mathbf{x} - \mathbf{y}\|^2}{\hat{\sigma}^2}\right) + \frac{1}{3} \exp\left(-\frac{\|\mathbf{x} - \mathbf{y}\|^2}{\hat{\sigma}^2}\right) + \frac{1}{3} \exp\left(-\frac{\|\mathbf{x} - \mathbf{y}\|^2}{4\hat{\sigma}^2}\right).$$

$\{\mathbf{x}_t^i\}$ and $\{\mathbf{y}_t^i\}$ are samples at time t obtained using the Euler-Maruyama scheme with a time step of 0.001 and generated by the normalizing flow, respectively, with initial distribution $p_0(\mathbf{x})$. $\hat{\sigma}$ is the median distance between $\{\mathbf{x}_t^i\}$ and $\{\mathbf{y}_t^i\}$.

6.1. Beneš SDE. We start with demonstrating the performance of our method on the Beneš SDE,

$$(6.4) \quad \begin{cases} d\mathbf{X}_t = \tanh \mathbf{X}_t dt + \mathbf{I}_d d\mathbf{W}_t, \\ \mathbf{X}_0 = \mathbf{x}_0, \quad \mathbf{x}_0 \in [-1, 1]^d, \end{cases}$$

for which the corresponding transition PDF admits an analytical solution,

$$p(\mathbf{x}, t | \mathbf{x}_0) = \frac{1}{(2\pi t)^{d/2}} \exp\left(-\frac{d}{2}t\right) \exp\left(-\frac{1}{2t} \|\mathbf{x} - \mathbf{x}_0\|^2\right) \prod_{i=1}^d \frac{\cosh x_i}{\cosh x_{0,i}}.$$

6.1.1. Two-dimensional case. We first consider the two-dimensional case. The conditional normalizing flow is composed of eight coupling layers. The parameters of each coupling layer are generated by a neural network with 32 random Fourier features and two hidden layers with 32 neurons. Training is performed for 10 adaptive iterations, each with 300 epochs. The Adam optimizer is used with an initial learning rate of 0.001, which is halved every 2000 epochs. 2×10^5 training points are employed with a batch size of 10^4 each epoch. A validation dataset of 10^6 samples is used to compute relative errors throughout the training process.

We first investigate the approximation performance of the transition PDF for various initial states, using the following five types of loss functions:

1. Vanilla PINN, $t \in (0, 1.5)$, $\gamma_1 = 0.2, \gamma_2 = 0.6, \gamma_3 = 0.2$;
2. Vanilla PINN, $t \in (0.1, 1.5)$, $\gamma_1 = 0.2, \gamma_2 = 0.6, \gamma_3 = 0.2$;
3. Time-weighted loss function, $t \in (0, 1.5)$, $\gamma_1 = 0.2, \gamma_2 = 0.6, \gamma_3 = 0.2$.

We present the $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t | \cdot))$ (6.1) in Figure 1a, where the validation dataset $\{\mathbf{x}_0, \mathbf{x}\}$ consists of 10^5 samples drawn uniformly from $[-1, 1]^2 \times [-5, 5]^2$. The blue curve, which overlaps with the orange one, represents the relative error of the base distribution. While it achieves high accuracy as $t \rightarrow 0$, its approximation quality deteriorates as t increases. A similar trend is observed when applying vanilla PINN. Although the base distribution stays close to the ground truth and the conditional normalizing flow remains close to the identity at small times, the corresponding residual at small times may still dominate, so the approximation accuracy is nearly the

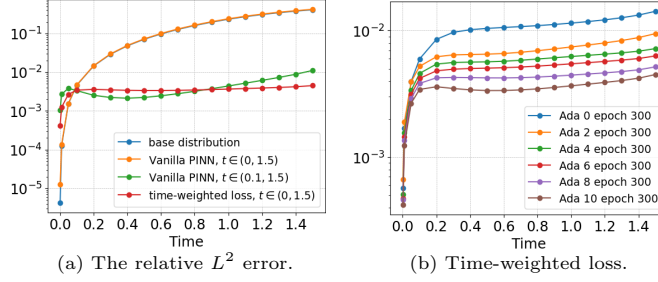


Fig. 1: Left: $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t|\cdot))$ for different loss functions. Right: the decay of $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t|\cdot))$ against the adaptive iterations for time-weighted loss.

same as that of the base distribution. If we exclude small times by setting the minimum training time to $t = 0.1$, the PDE constraints are no longer imposed near the initial condition. In this case, the approximation relies primarily on the continuity of the flow, and the accuracy improves for $t > 0.1$. By contrast, our method exploits the information from the base distribution and provides accurate approximations even as t increases, owing to the expressiveness of the conditional normalizing flow and the more balanced PINN's loss. Notably, we do not observe the typical error growth with time that often occurs in time-dependent problems. This highlights the effectiveness of the time-weighted loss function in mitigating error accumulation over time. Figure 1b further illustrates the decay of the $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t|\cdot))$ across adaptive iterations for the time-weighted loss function. Leveraging samples from the solution itself allows the model to capture the local behavior of the solution and enhances training efficiency.

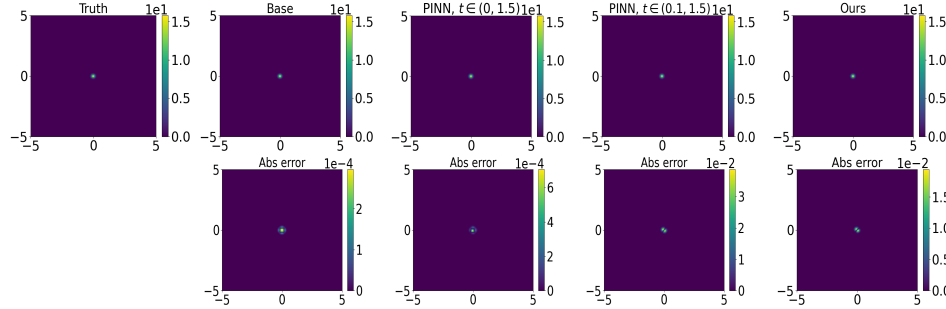
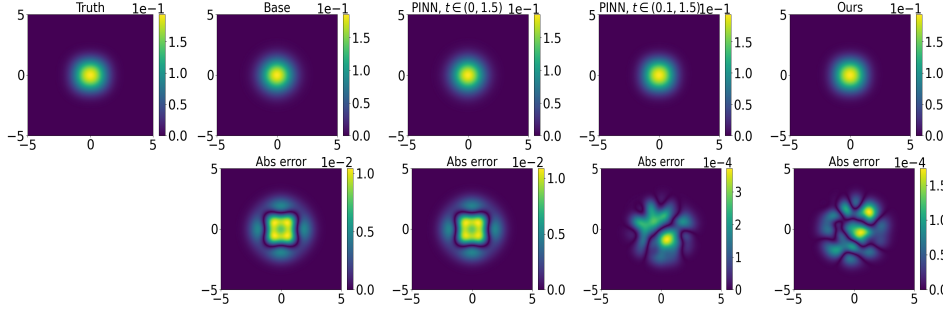
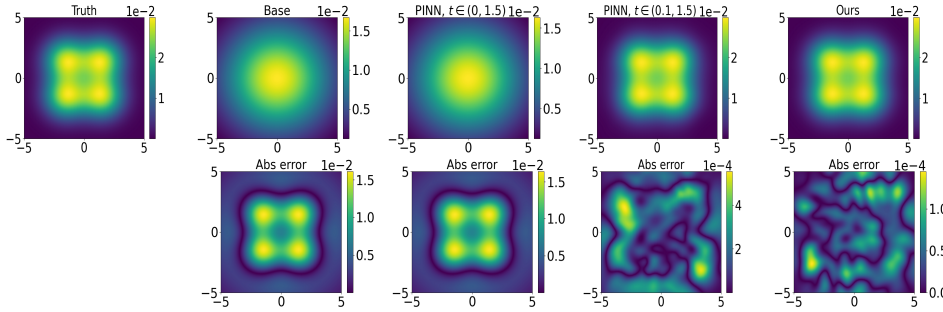


Fig. 2: Solutions (top row) and the absolute errors (bottom row) at $t = 0.01$. The first column shows the ground-truth solution. Columns 2–5 display numerical approximations obtained by, from left to right, the base distribution, Vanilla PINN with $t \in (0, 1.5)$, Vanilla PINN with $t \in (0.1, 1.5)$, time-weighted loss function with $t \in (0, 1.5)$.

Figure 2, Figure 3, and Figure 4 present the solutions at different times with the initial state $\mathbf{x}_0 = (0, 0)$ computed with the different loss functions. They all align closely to the ground truth except for the vanilla PINN's loss for $t \in (0, 1.5)$. While the base distribution and vanilla PINNs yield much more accurate approximations for small t , their performance degrades progressively as t increases. Especially, Figure 4 demonstrates that the conditional normalizing flow is capable of capturing multi-modal solutions. These results confirm the effectiveness of the proposed approach.

We next assess how well the solution of the FPE $p(\mathbf{x}, t)$ can be approximated under various initial conditions. To this end, we employ a variety of Beta distributions on $[-1, 1]^d$ as initial conditions, thereby evaluating the effectiveness of the conditional


 Fig. 3: Same as Figure 2 for $t = 0.5$.

 Fig. 4: Same as Figure 2 for $t = 1.5$.

normalizing flow for operator learning problems. More precisely,

$$\mathbf{Beta}(\mathbf{x}; \boldsymbol{\zeta}, \boldsymbol{\eta}) = \prod_{i=1}^d \frac{1}{B(\zeta_i, \eta_i)} \left(\frac{1+x_i}{2} \right)^{\zeta_i-1} \left(\frac{1-x_i}{2} \right)^{\eta_i-1}, \mathbf{x} \in [-1, 1]^d,$$

where $B(\zeta_i, \eta_i)$ denotes a beta function. We first explore the uniform case where $\boldsymbol{\zeta} = \boldsymbol{\eta} = (1, 1)$. We use the conditional normalizing flow trained with the time-weighted loss as an approximation of the transition PDF. Figure 5 reports the $\text{Rel}(p_{\boldsymbol{\theta}}(\cdot, t))$ (6.2) with a validation data set consisting 10^4 grid points over $[-5, 5]^2$ for different numbers of Monte Carlo samples used in evaluating $p_{\boldsymbol{\theta}}(\mathbf{x}, t)$ (5.2). Using q_1 as the proposal distribution yields the most accurate approximation for $t < 0.3$, whereas drawing samples directly from the initial distribution p_0 provides more reliable approximations for $t \geq 0.4$. By contrast, adopting q_2 , a mixture of q_1 and p_0 with $\alpha(t) = \exp(-6t)$, offers a balanced compromise, delivering acceptable accuracy across both small and large times. Figure 6 further illustrates the error decay with respect to the number of Monte Carlo samples. The relative error follows the Monte Carlo error scaling for $M \leq 10^5$. For a larger Monte Carlo sample size, however, the error in approximating the transition PDF becomes dominant. We also present the variance of the importance-sampling Monte Carlo estimator $\hat{p}_{\boldsymbol{\theta}}(\mathbf{x}, t)$ in (5.2) as a function of t in Figure 7 using 10^4 samples. The proposal q_1 substantially reduces the variance as $t \rightarrow 0$, but its variance increases as t grows. However, q_2 with different $\alpha(t)$, which combines q_1 with the initial distribution p_0 , yields a more balanced variance behavior over time. We note that traditional methods lose accuracy when the initial distribution is uniform, owing to the discontinuity of the initial distribution over the entire domain, if no special treatment of the discontinuity is taken. In contrast, our proposed methods remain unaffected by this discontinuity, thereby demonstrating their robustness in handling discontinuous initial data.

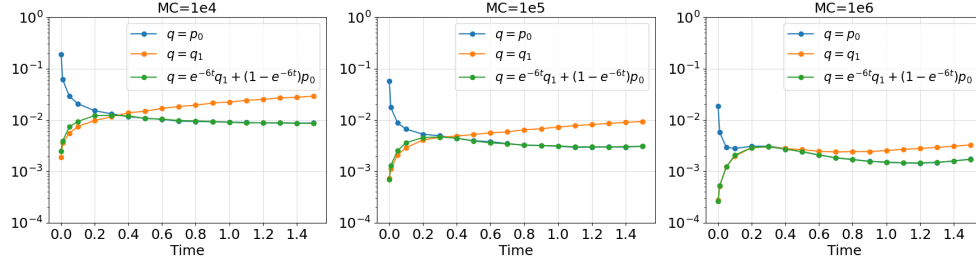


Fig. 5: Uniform case: Evolution of $\text{Rel}(\hat{p}_\theta(\cdot, t))$ over time. The subplots are arranged from left to right with $M = 10^4, 10^5, 10^6$.

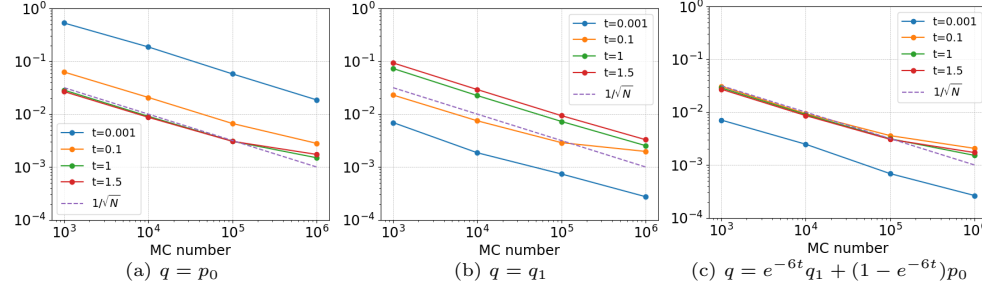


Fig. 6: Uniform case: The decay of $\text{Rel}(\hat{p}_\theta(\cdot, t))$ with respect to the number of Monte Carlo samples at different times. Left: Draw samples from p_0 ; Middle: Draw samples from q_1 ; Right: Draw samples from the mixture of p_0 and q_1 .

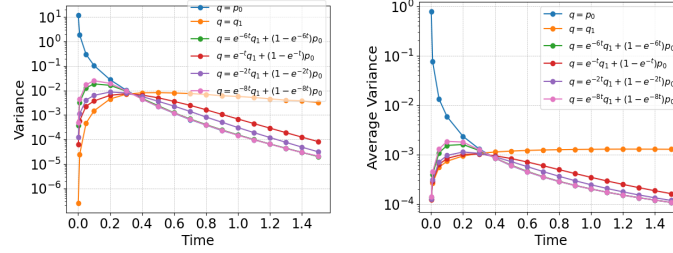


Fig. 7: The variance of $\hat{p}_\theta(\mathbf{x}, t)$ over time. Left: $\mathbf{x} = (0, 0)$. Right: Average variance of the set $\{\mathbf{x}^i\}$ consisting 10^4 grid points in $[-5, 5]^2$.

Moreover, we report in Figure 8 the time evolution of the $\text{Rel}(p_\theta(\cdot, t))$ for other Beta distributions chosen as initial conditions using q_2 with $\alpha(t) = \exp(-6t)$ as the proposal distribution. The numbers of Monte Carlo samples used in evaluating $\hat{p}_\theta(\mathbf{x}, t)$ are set to be 10^4 and 10^6 . The results confirm that the proposed method yields reliable approximations.

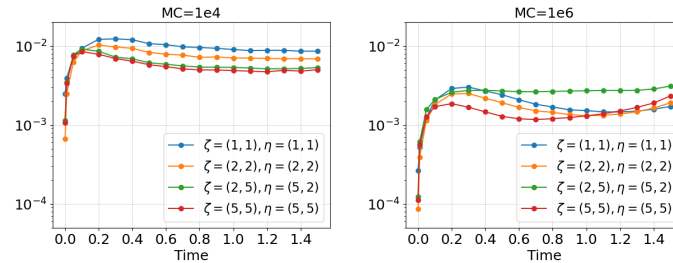


Fig. 8: Time evolution of the $\text{Rel}(\hat{p}_\theta(\cdot, t))$ for various initial distributions. Left: $M = 10^4$; Right: $M = 10^6$.

6.1.2. Four-dimensional case. We then solve the four-dimensional FPE with various initial distributions using the proposed time-weighted loss function. Moreover, when evaluating the interior residual, we scale the solution by a constant factor, set to 1000 in the four-dimensional case, to avoid numerical underflow, since the solution of the FPE becomes very small in higher dimensions. The conditional normalizing flow follows the KRnet structure with descending active transformation dimensions. It consists of three stages with $l_1 = 10, l_2 = 8$, and $l_3 = 6$. The parameters of each coupling layer are determined by a neural network with the same structure as in the two-dimensional case. We apply 20 adaptive iterations, each comprising 300 epochs. The updating rates are set to $\gamma_1 = 0.2$, $\gamma_2 = 0.4$, and $\gamma_3 = 0.4$. The initial learning rate is set to 0.001 and is halved every 4000 epochs. 2×10^5 training samples are employed with a batch size of 2×10^4 per epoch.

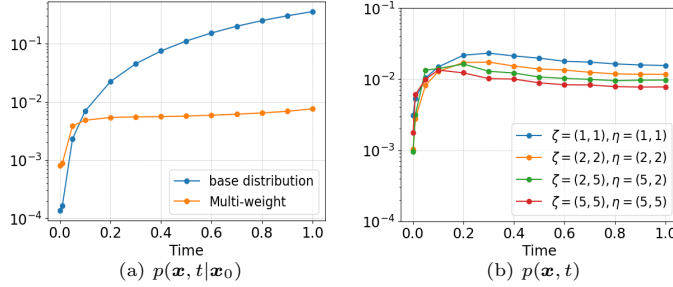


Fig. 9: The relative L^2 errors in the four-dimensional case. Left: $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t|\cdot))$; Right: $\text{Rel}(\hat{p}_{\theta}(\cdot, t))$ for various initial distributions with $M = 10^4$.

Figure 9a reports the $\text{Rel}(p_{\text{C-KRnet}}(\cdot, t|\cdot))$, where the validation dataset $\{\mathbf{x}_0, \mathbf{x}\}$ is uniformly sampled from $[-1, 1]^4 \times [-5, 5]^4$ with a total of 2×10^6 samples. It shows that the conditional normalizing flow achieves close agreement with the ground truth. Figure 9b presents $\text{Rel}(\hat{p}_{\theta}(\cdot, t))$ for different Beta distributions as initial conditions. For each initial condition, the numerical solution is obtained using the equation (5.2), where $q = q_2$ is defined in the equation (5.4) with $\alpha(t) = e^{-6t}$ and $M = 10^4$. A total of 10^5 test points are uniformly drawn from $[-5, 5]^4$. Together, these results demonstrate the effectiveness of the proposed method for moderately high-dimensional FPEs.

6.2. SDE with nonlinear drift and constant diffusion. In this section, we consider the following SDE with various initial distributions over $[-1, 1]^2$,

$$(6.5) \quad d \begin{pmatrix} X_1 \\ X_2 \end{pmatrix} = \begin{pmatrix} 2X_2 \\ 2X_1 - 0.8X_2 - 0.2X_1^3 \end{pmatrix} dt + \begin{pmatrix} \sqrt{0.4} \\ \sqrt{0.8} \end{pmatrix} d\mathbf{W}_t.$$

Both the structure of the conditional normalizing flow and the training setting are the same as in the two-dimensional Beneš equation. After obtaining the transition PDF, we approximate the PDF corresponding to various Beta initial distributions using the importance sampling (5.2) where $q = q_2$ (5.4) with $\alpha(t) = \exp(-6t)$ and $M = 10^4$.

Figure 10a presents the $\text{Rel}(\hat{p}_{\theta}(\cdot, t))$ (6.2) for different initial conditions, which is evaluated using a validation dataset consisting of 10^4 uniformly distributed grid points over $[-6, 6]^2$. We also report the MMD (6.3) with a sample size of 5×10^4 in Figure 10b. These results demonstrate that the conditional normalizing flow yields highly accurate approximations for small t , as the base distribution is quite close to the ground truth. It also provides reliable approximations as t increases.

6.3. SDE with nonlinear drift and state-dependent diffusion. For our final example, we validate the robustness of the proposed method in approximating

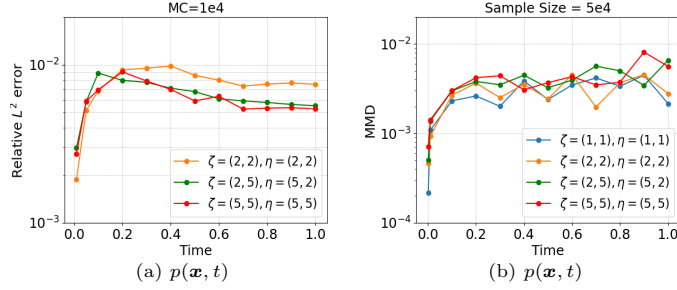


Fig. 10: Errors for various initial distributions in the nonlinear-drift, constant-diffusion case. Left: $\text{Rel}(\hat{p}_\theta(\cdot, t))$. Right: MMD.

the PDF of the following SDE with a state-dependent diffusion,

$$\begin{cases} d \begin{pmatrix} X_1 \\ X_2 \end{pmatrix} = \begin{pmatrix} X_2 \\ -0.5X_1 - 0.3X_1^3 + \sin X_2 \end{pmatrix} dt + \begin{pmatrix} 0.5 + 0.3X_1 & \\ & 0.4 + 0.1 \sin X_2 \end{pmatrix} d\mathbf{W}_t, \\ \mathbf{X}_0 = \mathbf{x}_0, \quad \mathbf{x}_0 \in [-1, 1]^2. \end{cases}$$

The conditional normalizing flow adopts the same structure as in the two-dimensional Beneš equation. Training is conducted over 4 adaptive iterations, each consisting of 1000 epochs. 4×10^5 training points are used, with a batch size of 5×10^4 per epoch.

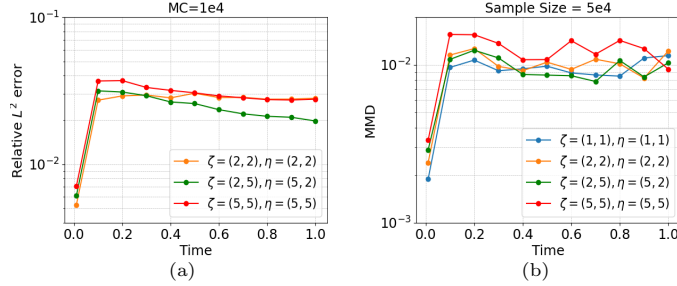


Fig. 11: Errors for various initial distributions in the nonlinear-drift, state-dependent-diffusion case. Left: $\text{Rel}(\hat{p}_\theta(\cdot, t))$. Right: MMD.

Figure 11a shows the $\text{Rel}(\hat{p}_\theta(\cdot, t))$ (6.2) for some initial conditions. The PDF $\hat{p}_\theta(\cdot, t)$ is evaluated using the importance sampling estimator (5.2) where $q = q_2$ with $\alpha(t) = \exp(-6t)$ and 10^4 Monte Carlo samples, while the $\text{Rel}(\hat{p}_\theta(\cdot, t))$ is computed on a validation dataset consisting of 10^4 uniformly distributed grid points over $[-5, 5]^2$, with the solution obtained by the ADI scheme used as the reference. We report the MMD (6.3) computed using 5×10^4 samples, as a function of time in the Figure 11b.

7. Conclusion. In this work, we proposed a transition-density-based operator learning method for solving the FPE under various initial conditions. The conditional normalizing flow is employed to approximate the transition PDF given any initial state, effectively constructing a mapping from the target stochastic process to a simple base process. The solution of the linearized SDE serves as a suitable candidate for the base stochastic process. We quantified the discrepancy between the PDFs of the linearized and original SDEs, which is characterized by differences in their drift and diffusion terms. By employing the solution of the linearized SDE and enforcing the identity mapping at $t = 0$, the initial condition is naturally satisfied by our conditional normalizing flow, thereby bypassing the explicit handling of the

Dirac delta initial condition. In addition, we designed a time-weighted loss function to alleviate numerical instabilities, ensuring that the proposed method respects causality in evolution PDEs while accounting for the increasing training difficulty as time progresses. To evaluate $p(\cdot, t)$ for arbitrary initial conditions over bounded domains, we used an importance sampling strategy to enhance accuracy. A series of numerical experiments validated the effectiveness and robustness. The four-dimensional Beneš SDE test demonstrated that the proposed method has the potential for solving moderately high-dimensional FPEs. This research contributes to the numerical solution of complex FPEs by combining the approximation capability of the base model (the linearized SDE) with the expressive power of the conditional normalizing flow. In future work, we will further refine this framework by improving the accuracy of the base stochastic process and exploring its potential applications. Another promising direction is to generalize the proposed method to a broader class of evolutionary partial differential equations.

Appendix A. Proof of Proposition 3.1.

Proof. Notice that the solution to the linearized SDE at time t is a Gaussian process, with PDF $p_{\sigma_t}(\mathbf{x}) = \mathcal{N}(\mathbf{x}; \mathbb{E}_{\sigma_t}[\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0], \Sigma_t[\widehat{\mathbf{X}}_t | \widehat{\mathbf{X}}_0 = \mathbf{x}_0])$ given the initial state $\mathbf{x}_0$. For notation simplicity, we omit the dependence on the initial state $\mathbf{x}_0$ and write $p_{\sigma_t}(\mathbf{x}) = \mathcal{N}(\mathbf{x}; \mathbb{E}_{\sigma_t}[\widehat{\mathbf{X}}_t], \Sigma_t[\widehat{\mathbf{X}}_t])$. Moreover, the linearized SDE and the original SDE share the same constant diffusion matrix, and $\nabla \mathbf{f}_\nu$ is bounded by assumption. In particular, if $C_f = 0$, then $\mathbf{f}_\sigma(\mathbf{x}) = \mathbf{f}_\nu(\mathbf{x})$ and consequently $\sigma_t = \nu_t$, which completes the proof. It remains to consider the case $C_f > 0$. It is easy to check that $(1 + |\mathbf{x}|)^{-2} |G^{ij}|, (1 + |\mathbf{x}|)^{-1} |\mathbf{f}_\nu|, (1 + |\mathbf{x}|)^{-1} |\mathbf{f}_\nu - \mathbf{f}_\sigma| \in L^1(\mathbb{R}^d \times [0, t_f], \{\sigma_t\}_{0 < t < t_f})$ for some finite t_f where $\|f\|_{L^1(\mathbb{R}^d \times [0, t_f], \{\sigma_t\}_{0 < t < t_f})} = \int_0^{t_f} \int_{\mathbb{R}^d} |f(\mathbf{x})| \sigma_t(d\mathbf{x}) dt$. Thus, according to Theorem 1.1 in [4], we have

$$(A.1) \quad H(\sigma_t | \nu_t) \leq \frac{1}{2} \int_0^t \int_{\mathbb{R}^d} \left| \mathbf{G}^{-1/2}(\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})) \right|^2 d\sigma_s ds.$$

Combining this result with the Pinsker–Csiszár–Kullback inequality [16] yields

$$(A.2) \quad \|\nu_t - \sigma_t\|_{\text{TV}}^2 \leq \frac{1}{4} \int_0^t \int_{\mathbb{R}^d} \left| \mathbf{G}^{-1/2}(\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})) \right|^2 d\sigma_s ds.$$

Notice that $\mathbf{G}_\sigma = \mathbf{G}_\nu = \mathbf{G}$ is a constant, non-degenerate matrix,

$$(A.3) \quad \left| \mathbf{G}^{-1/2}(\mathbf{f}_\sigma - \mathbf{f}_\nu) \right|^2 = (\mathbf{f}_\sigma - \mathbf{f}_\nu)^T \mathbf{G}^{-1}(\mathbf{f}_\sigma - \mathbf{f}_\nu) \leq \frac{1}{\lambda_{\min}(\mathbf{G})} |\mathbf{f}_\sigma - \mathbf{f}_\nu|^2,$$

where $\lambda_{\min}(\mathbf{G})$ is the smallest eigenvalue of $\mathbf{G}$. Meanwhile, for any $s > 0$, we have

$$\begin{aligned} \int_{\mathbb{R}^d} |\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})|^2 d\sigma_s &\leq 2 \int_{\mathbb{R}^d} |\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\sigma(\mathbf{x}_0)|^2 + |\mathbf{f}_\nu(\mathbf{x}_0) - \mathbf{f}_\nu(\mathbf{x})|^2 d\sigma_s \\ &= 2 \int_{\mathbb{R}^d} |\nabla \mathbf{f}_\nu(\mathbf{x}_0)(\mathbf{x} - \mathbf{x}_0)|^2 + \left| \int_0^1 \nabla \mathbf{f}_\nu(\mathbf{x}_0 + \tau(\mathbf{x} - \mathbf{x}_0))(\mathbf{x} - \mathbf{x}_0) d\tau \right|^2 d\sigma_s \\ &\leq 4C_f^2 \int_{\mathbb{R}^d} |\mathbf{x} - \mathbf{x}_0|^2 d\sigma_s = 4C_f^2 \text{tr}(\Sigma_s[\widehat{\mathbf{X}}_s]) + 4C_f^2 \left| \mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s] - \mathbf{x}_0 \right|^2, \end{aligned}$$

where the second inequality holds because $\|\nabla \mathbf{f}_\nu(\mathbf{x})\| \leq C_f, \forall \mathbf{x} \in \mathbb{R}^d$ and the last equality holds because σ_s is a Gaussian measure with mean $\mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s]$ and covariance

$\Sigma_s[\widehat{\mathbf{X}}_s]$. We give upper bounds for the two terms $\text{tr}(\Sigma_s[\widehat{\mathbf{X}}_s])$ and $|\mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s] - \mathbf{x}_0|^2$.

$$\begin{aligned} \text{tr}(\Sigma_s[\widehat{\mathbf{X}}_s]) &= \int_0^s \text{tr}(e^{\nabla \mathbf{f}(\mathbf{x}_0)(s-\tau)} \mathbf{g} \mathbf{g}^\top e^{(\nabla \mathbf{f}(\mathbf{x}_0))^\top (s-\tau)}) d\tau \\ &\leq 2\text{tr}(\mathbf{G}) \int_0^s \sigma_{\max}^2(e^{\nabla \mathbf{f}(\mathbf{x}_0)\tau}) d\tau \\ &\leq 2\text{tr}(\mathbf{G}) \int_0^s \left(e^{\sigma_{\max}(\nabla \mathbf{f}(\mathbf{x}_0)\tau)}\right)^2 d\tau \leq 2\text{tr}(\mathbf{G}) \int_0^s e^{2C_f\tau} d\tau = \text{tr}(\mathbf{G}) \frac{e^{2C_f s} - 1}{C_f}, \end{aligned}$$

where $\sigma_{\max}(\mathbf{A})$ is the spectral norm of the matrix $\mathbf{A}$. Meanwhile, equation (3.5) gives

$$|\mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s] - \mathbf{x}_0|^2 = \left| \int_0^s e^{\nabla \mathbf{f}(\mathbf{x}_0)\tau} \mathbf{f}(\mathbf{x}_0) d\tau \right|^2 \leq |\mathbf{f}(\mathbf{x}_0)|^2 \left(\frac{e^{C_f s} - 1}{C_f} \right)^2.$$

Therefore, given $\mathbf{x}_0$, $\text{tr}(\Sigma_s[\widehat{\mathbf{X}}_s]) = \mathcal{O}(s)$, $|\mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s] - \mathbf{x}_0|^2 = \mathcal{O}(s^2)$. Moreover,

$$\begin{aligned} H(\sigma_t | \nu_t) &\leq \frac{1}{2\lambda_{\min}(\mathbf{G})} \int_0^t \int_{\mathbb{R}^d} |\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})|^2 d\sigma_s ds, \\ &\leq \frac{2C_f^2}{\lambda_{\min}(\mathbf{G})} \int_0^t \left(\text{tr}(\mathbf{G}) \frac{e^{2C_f s} - 1}{C_f} + |\mathbf{f}(\mathbf{x}_0)|^2 \left(\frac{e^{C_f s} - 1}{C_f} \right)^2 \right) ds \\ &= \frac{2C_f^2}{\lambda_{\min}(\mathbf{G})} \left(\text{tr}(\mathbf{G}) (t^2 + \mathcal{O}(t^3)) + |\mathbf{f}(\mathbf{x}_0)|^2 \left(\frac{1}{3}t^3 + \mathcal{O}(t^4) \right) \right) = \mathcal{O}(t^2) \text{ as } t \rightarrow 0. \end{aligned}$$

Hence, $\|(\nu_t - \sigma_t)\|_{TV}^2 \leq \frac{1}{2} H(\sigma_t | \nu_t) = \mathcal{O}(t^2)$ as $t \rightarrow 0$. Furthermore, if the Hessian of the drift $\mathbf{f}_\nu$ is uniformly bounded by $H_{\mathbf{f}}$, then

$$\begin{aligned} &\int_{\mathbb{R}^d} |\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})|^2 d\sigma_s \\ &= \int_{\mathbb{R}^d} \left| \int_0^1 (1-h)(\mathbf{x} - \mathbf{x}_0)^\top \text{Hess}(\mathbf{x}_0 + h(\mathbf{x} - \mathbf{x}_0))(\mathbf{x} - \mathbf{x}_0) dh \right|^2 d\sigma_s \leq \frac{H_{\mathbf{f}}^2}{4} \int_{\mathbb{R}^d} |\mathbf{x} - \mathbf{x}_0|^4 d\sigma_s. \end{aligned}$$

$$\begin{aligned} \int_{\mathbb{R}^d} |\mathbf{x} - \mathbf{x}_0|^4 d\sigma_s &\leq 8 \int_{\mathbb{R}^d} |\mathbf{x} - \mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s]|^4 d\sigma_s + 8 \int_{\mathbb{R}^d} |\mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s] - \mathbf{x}_0|^4 d\sigma_s \\ &\leq 24\text{tr}^2(\Sigma_s[\widehat{\mathbf{X}}_s]) + 8 \int_{\mathbb{R}^d} \left| \int_0^s e^{\nabla \mathbf{f}(\mathbf{x}_0)\tau} \mathbf{f}(\mathbf{x}_0) d\tau \right|^4 d\sigma_s \\ &\leq 24\text{tr}^2(\mathbf{G}) \frac{(e^{2C_f s} - 1)^2}{C_f^2} + 8|\mathbf{f}(\mathbf{x}_0)|^4 \left(\frac{e^{C_f s} - 1}{C_f} \right)^4 = 96\text{tr}^2(\mathbf{G}) s^2 + \mathcal{O}(s^3), \quad s \rightarrow 0, \end{aligned}$$

having used that $\int_{\mathbb{R}^d} |\mathbf{x} - \mathbb{E}_{\sigma_s}[\widehat{\mathbf{X}}_s]|^4 d\sigma_s = \text{tr}^2(\Sigma_s[\widehat{\mathbf{X}}_s]) + 2\text{tr}(\Sigma_s(\widehat{\mathbf{X}}_s)^2)$ and $\text{tr}(\Sigma_s^2[\widehat{\mathbf{X}}_s]) \leq \text{tr}^2(\Sigma_s[\widehat{\mathbf{X}}_s])$. Hence, $\int_0^t \int_{\mathbb{R}^d} |\mathbf{f}_\sigma(\mathbf{x}) - \mathbf{f}_\nu(\mathbf{x})|^2 d\sigma_s ds \leq \frac{H_{\mathbf{f}}^2}{4} \int_0^t \int_{\mathbb{R}^d} |\mathbf{x} - \mathbf{x}_0|^4 d\sigma_s ds = 8H_{\mathbf{f}}^2 \text{tr}^2(\mathbf{G}) t^3 + \mathcal{O}(t^4)$, as $t \rightarrow 0$. Therefore, $H(\sigma_t | \nu_t) = \mathcal{O}(t^3)$, $\|\nu_t - \sigma_t\|_{TV}^2 = \mathcal{O}(t^3)$ as $t \rightarrow 0$. $\square$

Appendix B. Proof of Proposition 3.2.

Proof. Let $\rho = \nu + \sigma$, then $\nu \ll \rho$, $\sigma \ll \rho$. Define $p_\nu := \frac{d\nu}{d\rho}$, $p_\sigma := \frac{d\sigma}{d\rho}$, thus $\nu = p_\nu \rho$, $\sigma = p_\sigma \rho$, and $p_\nu + p_\sigma = 1$. For $\mathbf{f} = (f^1, \dots, f^m) \in L^2(\mathbb{R}^d, \nu; \mathbb{R}^m) \cap L^2(\mathbb{R}^d, \sigma; \mathbb{R}^m)$,

$$\begin{aligned} & \left| \int_{\mathbb{R}^d} \mathbf{f}(\mathbf{x}) d\nu(\mathbf{x}) - \int_{\mathbb{R}^d} \mathbf{f}(\mathbf{x}) d\sigma(\mathbf{x}) \right| = \left| \int_{\mathbb{R}^d} \mathbf{f}(\mathbf{x}) (p_\nu(\mathbf{x}) - p_\sigma(\mathbf{x})) d\rho(\mathbf{x}) \right| \\ & \leq \int_{\mathbb{R}^d} \left| \mathbf{f}(\mathbf{x}) (p_\nu(\mathbf{x}) + p_\sigma(\mathbf{x}))^{\frac{1}{2}} \frac{p_\nu(\mathbf{x}) - p_\sigma(\mathbf{x})}{(p_\nu(\mathbf{x}) + p_\sigma(\mathbf{x}))^{\frac{1}{2}}} \right| d\rho(\mathbf{x}) \\ & \leq \left(\int_{\mathbb{R}^d} |\mathbf{f}(\mathbf{x})|^2 (p_\nu(\mathbf{x}) + p_\sigma(\mathbf{x})) d\rho(\mathbf{x}) \right)^{\frac{1}{2}} \left(\int_{\mathbb{R}^d} \frac{(p_\nu(\mathbf{x}) - p_\sigma(\mathbf{x}))^2}{p_\nu(\mathbf{x}) + p_\sigma(\mathbf{x})} d\rho(\mathbf{x}) \right)^{\frac{1}{2}} \\ & \leq (\|\mathbf{f}\|_{L^2(\mathbb{R}^d, \nu; \mathbb{R}^m)} + \|\mathbf{f}\|_{L^2(\mathbb{R}^d, \sigma; \mathbb{R}^m)}) \sqrt{2\|\nu - \sigma\|_{\text{TV}}}. \quad \square \end{aligned}$$

Appendix C. Proof of Proposition 3.3.

Proof. Let $\Pi_k : \mathbb{R}^d \rightarrow \mathbb{R}^k$ denote the canonical projection onto the first k coordinates, that is, $\Pi_k(x_1, x_2, \dots, x_d) = (x_1, x_2, \dots, x_k)$. The KR map T_i admits $T_i(\mathbf{x}) = (T_i^1(x_1), T_i^2(x_1, x_2), \dots, T_i^d(x_1, \dots, x_d))$. For $1 \leq k \leq d$, denote its k -dimensional truncation by $T_i^{(k)}(\mathbf{x}^{(1:k)}) = \Pi_k \circ T_i(\mathbf{x})$ where $\mathbf{x}^{(1:k)} = (x_1, \dots, x_k)$. We denote by $\mu^{(k)} := \Pi_{k\#}\mu$ and $\nu_i^{(k)} := \Pi_{k\#}\nu_i$ the k -dimensional marginal distribution of μ and ν_i , respectively. Their corresponding PDFs are denoted by $p_\mu^{(k)}$ and $p_{\nu_i}^{(k)}$. By the definition of total variation distance, we have $\|\nu_i^{(k)} - \mu^{(k)}\|_{\text{TV}} \leq \|\nu_i - \mu\|_{\text{TV}} \rightarrow 0$.

We now prove $T_i(\mathbf{x}) \rightarrow \mathbf{x}$ in μ -probability as $i \rightarrow \infty$ by induction.

We first show that $T_i^1(x_1) \rightarrow x_1$ in μ -probability as $i \rightarrow \infty$. Let F_i^1 and F^1 be the cumulative distribution functions of $\nu_i^{(1)}$ and $\mu^{(1)}$, respectively. Let $\delta_i^1 = \sup_{x \in \mathbb{R}} |F_i^1(x) - F^1(x)|$. Since $\|\nu_i^{(1)} - \mu^{(1)}\|_{\text{TV}} \rightarrow 0$, we have $\lim_{i \rightarrow \infty} \delta_i^1 = 0$. Notice that $T_i^1(x_1) = (F_i^1)^{(-1)} \circ F^1(x_1)$, and $(F_i^1)^{(-1)}(s) = \inf\{y \in \mathbb{R} : F_i^1(y) \geq s\}$. For any $\epsilon > 0$, if $T_i^1(x_1) \geq x_1 + \epsilon$, then $F_i^1(x_1 + \epsilon) \leq F^1(x_1)$. Hence, $F^1(x_1 + \epsilon) - F^1(x_1) \leq F^1(x_1 + \epsilon) - F_i^1(x_1 + \epsilon) \leq \delta_i^1$. On the other hand, if $T_i^1(x_1) \leq x_1 - \epsilon$, then $F_i^1(x_1 - \epsilon) \geq F^1(x_1)$. Hence, $F^1(x_1) - F^1(x_1 - \epsilon) \leq F_i^1(x_1 - \epsilon) - F^1(x_1 - \epsilon) \leq \delta_i^1$. Therefore, we have $\{|T_i^1(x_1) - x_1| \geq \epsilon\} \subseteq \{F^1(x_1) - F^1(x_1 - \epsilon) \leq \delta_i^1\} \cup \{F^1(x_1 + \epsilon) - F^1(x_1) \leq \delta_i^1\}$. Let $\Omega^-(\delta) = \{x_1 \mid F(x_1) - F(x_1 - \epsilon) \leq \delta\}$, and $\Omega^+(\delta) = \{x_1 \mid F(x_1 + \epsilon) - F(x_1) \leq \delta\}$. Thus, $\mu^{(1)}(|T_i^1(x_1) - x_1| \geq \epsilon) \leq \mu^{(1)}(\Omega^-(\delta_i^1)) + \mu^{(1)}(\Omega^+(\delta_i^1))$, where $\mu^{(1)}(\Omega^-(\delta_i^1)) = \int 1_{\Omega^-(\delta_i^1)}(x_1) d\mu^{(1)}(x_1)$, and $\mu^{(1)}(\Omega^+(\delta_i^1)) = \int 1_{\Omega^+(\delta_i^1)}(x_1) d\mu^{(1)}(x_1)$. According to the dominated convergence theorem, we have

$$\lim_{i \rightarrow \infty} \mu^{(1)}(\Omega^-(\delta_i^1)) = \int 1_{\Omega^-(0)}(x_1) d\mu^{(1)}(x_1) = \mu^{(1)}(\Omega^-(0)) = \sum_{j \in \mathbb{Z}} \mu^{(1)}(\Omega^-(0) \cap ((j-1)\epsilon, j\epsilon]).$$

We note that $x_1 \in \Omega^-(0)$ implies $F(x_1) - F(x_1 - \epsilon) = 0$. For any $j \in \mathbb{Z}$, let $\bar{x}_j = \sup\{x_1 \in \Omega^-(0) \cap ((j-1)\epsilon, j\epsilon]\}$, then $\Omega^-(0) \cap ((j-1)\epsilon, j\epsilon] \subset ((j-1)\epsilon, \bar{x}_j]$ and

$$\mu^{(1)}(\Omega^-(0)) \leq \sum_{j \in \mathbb{Z}} \mu^{(1)}(((j-1)\epsilon, \bar{x}_j]) \leq \sum_{j \in \mathbb{Z}} \mu^{(1)}((\bar{x}_j - \epsilon, \bar{x}_j]) = 0.$$

Thus, $\lim_{i \rightarrow \infty} \mu^{(1)}(\Omega^-(\delta_i^1)) = 0$. Similarly, we have $\lim_{i \rightarrow \infty} \mu^{(1)}(\Omega^+(\delta_i^1)) = 0$. Hence, $\lim_{i \rightarrow \infty} \mu(\{|T_i^1(x_1) - x_1| \geq \epsilon\}) = 0$, and $T_i^1(x_1) \xrightarrow{i \rightarrow \infty} x_1$ in μ -probability.

Suppose that for any $1 \leq k \leq d-1$, $T_i^k(\mathbf{x}^{(1:k)}) \xrightarrow{i \rightarrow \infty} x_k$ in μ -probability. We next prove that $T_i^{k+1}(\mathbf{x}^{(1:k+1)}) \xrightarrow{i \rightarrow \infty} x_{k+1}$ in μ -probability. For notation simplicity,

we write $\zeta_{i,T_i^{(k)}}(dx_{k+1}) := \nu_i(dx_{k+1} | T_i^{(k)}(\mathbf{x}^{(1:k)}))$, $\zeta_i(dx_{k+1}) := \nu_i(dx_{k+1} | \mathbf{x}^{(1:k)})$, $\zeta_{T_i^{(k)}}(dx_{k+1}) := \mu(dx_{k+1} | T_i^{(k)}(\mathbf{x}^{(1:k)}))$, $\zeta(dx_{k+1}) := \mu(dx_{k+1} | \mathbf{x}^{(1:k)})$. We show that

$$(C.1) \quad \int \|\zeta_{i,T_i^{(k)}} - \zeta\|_{\text{TV}} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \rightarrow 0, \text{ as } i \rightarrow \infty.$$

Notice that $\|\zeta_{i,T_i^{(k)}} - \zeta\|_{\text{TV}} \leq \|\zeta_{i,T_i^{(k)}} - \zeta_{T_i^{(k)}}\|_{\text{TV}} + \|\zeta_{T_i^{(k)}} - \zeta\|_{\text{TV}}$, and

$$\begin{aligned} & \int \|\zeta_{i,T_i^{(k)}} - \zeta_{T_i^{(k)}}\|_{\text{TV}} d\mu^{(k)}(\mathbf{x}^{(1:k)}) = \int \|\zeta_i - \zeta\|_{\text{TV}} d\nu_i^{(k)}(\mathbf{x}^{(1:k)}) \\ &= \frac{1}{2} \int |p_{\nu_i}(x_{k+1} | \mathbf{x}^{(1:k)}) - p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)})| p_{\nu_i}^{(k)}(\mathbf{x}^{(1:k)}) dx_{k+1} d\mathbf{x}^{(1:k)} \\ &\leq \frac{1}{2} \int |p_{\nu_i}^{(k+1)} - p_{\mu}^{(k+1)}| d\mathbf{x}^{(1:k+1)} + \frac{1}{2} \int |p_{\mu}^{(k)} - p_{\nu_i}^{(k)}| d\mathbf{x}^{(1:k)} \\ &= \|\nu_i^{(k+1)} - \mu^{(k+1)}\|_{\text{TV}} + \|\nu_i^{(k)} - \mu^{(k)}\|_{\text{TV}} \leq 2\|\nu_i - \mu\|_{\text{TV}} \rightarrow 0, \text{ as } i \rightarrow \infty. \end{aligned}$$

Therefore, to prove (C.1), we only need to show $\int \|\zeta_{T_i^{(k)}} - \zeta\|_{\text{TV}} d\mu^{(k)}$ goes to 0 as $i \rightarrow \infty$. Since μ is absolutely continuous and $\int p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)}) dx_{k+1} = 1$ for $\mu^{(k)}$ -a.e. $\mathbf{x}^{(1:k)}$, $p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)})$ can be regarded as a $L^1(\mathbb{R})$ function for $\mu^{(k)}$ -a.e. $\mathbf{x}^{(1:k)}$. According to Lemma 1.2.31 in [20], for any $\epsilon > 0$, there exists a bounded continuous $L^1(\mathbb{R})$ function $\phi(x_{k+1}, \mathbf{x}^{(1:k)})$ with $\|\phi\|_{\infty} := \sup_{\mathbf{z} \in \mathbb{R}^k} \|\phi(\cdot, \mathbf{z})\|_{L^1(\mathbb{R})}$ such that

$$H_{\phi} := \int \|p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)}) - \phi(x_{k+1}, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) < \epsilon.$$

Thus,

$$\begin{aligned} & \int \|\zeta_{T_i^{(k)}} - \zeta\|_{\text{TV}} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \\ (C.2) \quad &= \frac{1}{2} \int \|p_{\mu}(x_{k+1} | T_i^{(k)}(\mathbf{x}^{(1:k)})) - p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \\ &\leq \frac{1}{2} \int \|p_{\mu}(x_{k+1} | T_i^{(k)}(\mathbf{x}^{(1:k)})) - \phi(x_{k+1}, T_i^{(k)}(\mathbf{x}^{(1:k)}))\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \\ &\quad + \frac{1}{2} \int \|\phi(x_{k+1}, T_i^{(k)}(\mathbf{x}^{(1:k)})) - \phi(x_{k+1}, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) + \frac{1}{2} H_{\phi}. \end{aligned}$$

Since $T_i^{(k)} \# \mu^{(k)} = \nu_i^{(k)}$ and $\|\nu_i^{(k)} - \mu^{(k)}\|_{\text{TV}} \rightarrow 0$, there exists $i_0 \in \mathbb{N}$ such that $\|\nu_i^{(k)} - \mu^{(k)}\|_{\text{TV}} < \frac{\epsilon}{1 + \|\phi\|_{\infty}}$ for any $i \geq i_0$. Then for $i \geq i_0$, we have

$$\begin{aligned} & \int \|p_{\mu}(x_{k+1} | T_i^{(k)}(\mathbf{x}^{(1:k)})) - \phi(x_{k+1}, T_i^{(k)}(\mathbf{x}^{(1:k)}))\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \\ (C.3) \quad &= \int \|p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)}) - \phi(x_{k+1}, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\nu_i^{(k)}(\mathbf{x}^{(1:k)}) \\ &= H_{\phi} + \int \|p_{\mu}(x_{k+1} | \mathbf{x}^{(1:k)}) - \phi(x_{k+1}, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d(\nu_i^{(k)} - \mu^{(k)})(\mathbf{x}^{(1:k)}) \\ &< \epsilon + 2(1 + \|\phi\|_{\infty}) \|\nu_i^{(k)} - \mu^{(k)}\|_{\text{TV}} < 3\epsilon. \end{aligned}$$

Meanwhile, $T_i^{(k)}(\mathbf{x}^{(1:k)}) \rightarrow \mathbf{x}^{(1:k)}$ in μ -probability as $i \rightarrow \infty$ and $\phi(x_{k+1}, \mathbf{x}^{(1:k)})$ is a bounded continuous function. Hence, $\phi(x_{k+1}, T_i^{(k)}(\mathbf{x}^{(1:k)})) \rightarrow \phi(x_{k+1}, \mathbf{x}^{(1:k)})$ in $\mu^{(k)}$ -probability with respect to the L^1 norm as $i \rightarrow \infty$. Since ϕ is bounded, hence integrable, according to the Vitali convergence theorem [14], we have $\int \|\phi(\cdot, T_i^{(k)}(\mathbf{x}^{(1:k)})) - \phi(\cdot, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \rightarrow 0$ as $i \rightarrow \infty$. Thus, there exists $i_1 \in \mathbb{N}$ such that

$$(C.4) \quad \int \|\phi(\cdot, T_i^{(k)}(\mathbf{x}^{(1:k)})) - \phi(\cdot, \mathbf{x}^{(1:k)})\|_{L^1(\mathbb{R})} d\mu^{(k)}(\mathbf{x}^{(1:k)}) < \epsilon, \text{ for any } i \geq i_1.$$

Substituting equations (C.3)-(C.4) into equation (C.2) gives $\int \|\zeta_{i, T_i^{(k)}} - \zeta\|_{TV} d\mu^{(k)} < \frac{5}{2}\epsilon$ as $i > \max\{i_0, i_1\}$. Hence, $\lim_{i \rightarrow \infty} \int \|\zeta_{i, T_i^{(k)}} - \zeta\|_{TV} d\mu^{(k)} = 0$ and equation (C.1) holds. For $\mu^{(k)}$ -a.e. $\mathbf{x}^{(1:k)}$, $T_i^{k+1}(\mathbf{x}^{(1:k)}, x_{k+1})$ is a one-dimensional KR map from ζ to $\zeta_{i, T_i^{(k)}}$. Let F_i^{k+1} and F^{k+1} be the cumulative distribution functions of ζ and $\zeta_{i, T_i^{(k)}}$, respectively. Denote $\delta_i^{k+1}(\mathbf{x}^{(1:k)}) = \sup_{x_{k+1}} |F_i^{k+1}(x_{k+1}) - F^{k+1}(x_{k+1})|$. Then

$$\int \delta_i^{k+1}(\mathbf{x}^{(1:k)}) d\mu^{(k)}(\mathbf{x}^{(1:k)}) \leq 2 \int \|\zeta_{i, T_i^{(k)}} - \zeta\|_{TV} d\mu^{(k)}(\mathbf{x}^{(1:k)}) \xrightarrow{i \rightarrow \infty} 0.$$

We want to show that $\mu^{(k+1)}(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon) \xrightarrow{i \rightarrow \infty} 0$ or equivalently,

$$\forall \eta > 0, \exists i_\eta \in \mathbb{N} : \mu^{(k+1)}(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon) \leq \eta, \forall i \geq i_\eta.$$

Let $R_\delta(\mathbf{x}^{(1:k)}) = \zeta(F^{k+1}(x_{k+1}) - F^{k+1}(x_{k+1} - \epsilon) \leq \delta) + \zeta(F^{k+1}(x_{k+1} + \epsilon) - F^{k+1}(x_{k+1}) \leq \delta)$. Following the same argument as for T_i^1 and since $0 \leq R_\delta \leq 2$, we have for any $\epsilon > 0$, $\int R_\delta d\mu^{(k)} \rightarrow 0$ as $\delta \rightarrow 0$. Let here δ_η be s.t. $\int R_\delta d\mu^{(k)} \leq \eta/2$. Since $\lim_{i \rightarrow \infty} \int \delta_i^{k+1}(\mathbf{x}^{(1:k)}) d\mu^{(k)} = 0$, there exists i_η s.t. $\int \delta_i^{k+1}(\mathbf{x}^{(1:k)}) d\mu^{(k)} \leq \delta_\eta \cdot \eta/2$, $\forall i \geq i_\eta$. Using that $\{|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon\} \cap \{\delta_i^{k+1} \leq \delta_\eta\} \subset \{F^{k+1}(x_{k+1}) - F^{k+1}(x_{k+1} - \epsilon) \leq \delta_\eta\} \cup \{F^{k+1}(x_{k+1} + \epsilon) - F^{k+1}(x_{k+1}) \leq \delta_\eta\}$, we have

$$\begin{aligned} & \mu^{(k+1)}(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon) \\ &= \mu^{(k+1)}(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon, \delta_i^{k+1}(\mathbf{x}^{(1:k)}) \leq \delta_\eta) \\ & \quad + \mu^{(k+1)}(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon, \delta_i^{k+1}(\mathbf{x}^{(1:k)}) > \delta_\eta) \\ &\leq \int \zeta(|T_i^{k+1}(\mathbf{x}^{(1:k+1)}) - x_{k+1}| \geq \epsilon) 1_{\delta_i^{k+1}(\mathbf{x}^{(1:k)}) \leq \delta_\eta} d\mu^{(k)} + \mu^{(k+1)}(\delta_i^{k+1} > \delta_\eta) \\ &\leq \int R_{\delta_\eta}(\mathbf{x}^{(1:k)}) d\mu^{(k)}(\mathbf{x}^{(1:k)}) + \mu^{k+1}(\delta_i^{k+1} > \delta_\eta) \leq \frac{\eta}{2} + \frac{1}{\delta_\eta} \int \delta_i^{k+1}(\mathbf{x}^{(1:k)}) d\mu^{(k)} \leq \eta, \forall i \geq i_\eta, \end{aligned}$$

which shows that $T_i^{k+1}(\mathbf{x}^{(1:k+1)}) \rightarrow x_{k+1}$ in $\mu^{(k+1)}$ -probability.

We now prove that $\lim_{i \rightarrow \infty} \int_{\mathbb{R}^d} |T_i(\mathbf{x}) - \mathbf{x}|^2 d\mu(\mathbf{x}) = 0$. Since $T_{i\#} \mu = \nu_i$,

$$\int_{\mathbb{R}^d} |T_i(\mathbf{x})|^2 d\mu(\mathbf{x}) = \int_{\mathbb{R}^d} |\mathbf{x}|^2 d\nu_i(\mathbf{x}) \rightarrow \int_{\mathbb{R}^d} |\mathbf{x}|^2 d\mu(\mathbf{x}), \text{ as } i \rightarrow \infty.$$

That is to say, $\lim_{i \rightarrow \infty} \|T_i\|_{L^2(\mathbb{R}^d, \mu)} = \|id\|_{L^2(\mathbb{R}^d, \mu)}$. Meanwhile, $2|T_i(\mathbf{x}) \cdot \mathbf{x}| \leq |T_i(\mathbf{x})|^2 + |\mathbf{x}|^2$, and $\{T_i\}_{i \geq 1}$ are uniformly integrable in $L^1(\mu)$. Thus, according to the Vitali convergence theorem [14], we have $\lim_{i \rightarrow \infty} \int_{\mathbb{R}^d} T_i(\mathbf{x}) \cdot \mathbf{x} d\mu(\mathbf{x}) = \int_{\mathbb{R}^d} |\mathbf{x}|^2 d\mu(\mathbf{x})$. Hence,

$$\|T_i - id\|_{L^2(\mathbb{R}^d, \mu)}^2 = \|T_i\|_{L^2(\mathbb{R}^d, \mu)}^2 + \|id\|_{L^2(\mathbb{R}^d, \mu)}^2 - 2 \int_{\mathbb{R}^d} T_i(\mathbf{x}) \cdot \mathbf{x} d\mu(\mathbf{x}) \xrightarrow{i \rightarrow \infty} 0.$$

Therefore, $\|T_i - id\|_{L^2(\mu)}^2 \rightarrow 0$ as $i \rightarrow \infty$. $\square$

Appendix D. Proof of Proposition 4.1.

Proof. For any $\tilde{p}(\mathbf{x}, t | \mathbf{x}_0) = \mathcal{N}(\mathbf{x}; \mathbf{m}(t), \Sigma(t))$, $\Sigma(t) \sim t\mathbf{G}$, we have,
 $\frac{\partial \tilde{p}}{\partial t} = \tilde{p} \left(-\frac{1}{2} \text{tr} \left(\Sigma^{-1} \frac{d\Sigma}{dt} \right) + \frac{d\mathbf{m}^\top}{dt} \Sigma^{-1} (\mathbf{x} - \mathbf{m}) + \frac{1}{2} (\mathbf{x} - \mathbf{m})^\top \Sigma^{-1} \frac{d\Sigma}{dt} \Sigma^{-1} (\mathbf{x} - \mathbf{m}) \right),$
 $\nabla \tilde{p} = \tilde{p} (-\Sigma^{-1} (\mathbf{x} - \mathbf{m})), \quad \nabla \cdot \nabla \cdot (\mathbf{g} \mathbf{g}^\top \tilde{p}) = \tilde{p} (\mathbf{g}^\top \Sigma^{-1} (\mathbf{x} - \mathbf{m}))^2 - \tilde{p} \mathbf{g}^\top \Sigma^{-1} \mathbf{g}.$
Hence, $r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p}) = \tilde{p} \tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})$,

$$\begin{aligned} \tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p}) = & -\frac{1}{2} \text{tr} \left(\Sigma^{-1} \frac{d\Sigma}{dt} \right) + \frac{d\mathbf{m}^\top}{dt} \Sigma^{-1} (\mathbf{x} - \mathbf{m}) + \frac{1}{2} (\mathbf{x} - \mathbf{m})^\top \Sigma^{-1} \frac{d\Sigma}{dt} \Sigma^{-1} (\mathbf{x} - \mathbf{m}) \\ & + \nabla \cdot \mathbf{f} - \mathbf{f}^\top \Sigma^{-1} (\mathbf{x} - \mathbf{m}) - \frac{1}{2} (\mathbf{g}^\top \Sigma^{-1} (\mathbf{x} - \mathbf{m}))^2 + \frac{1}{2} \mathbf{g}^\top \Sigma^{-1} \mathbf{g}. \end{aligned}$$

Thus, $|\tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})| \leq \tilde{C}_0 + \tilde{C}_1 t^{-1} + \tilde{C}_2 |\mathbf{x} - \mathbf{m}| t^{-1} + \tilde{C}_3 |\mathbf{x} - \mathbf{m}|^2 t^{-1} + \tilde{C}_4 |\mathbf{x} - \mathbf{m}|^2 t^{-2}$.

$$\begin{aligned} |\tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2 \leq & (C_0 + C_1 t^{-1} + C_2 t^{-2}) + (C_3 t^{-1} + C_4 t^{-2}) |\mathbf{x} - \mathbf{m}| \\ & + (C_5 t^{-1} + C_6 t^{-2} + C_7 t^{-3}) |\mathbf{x} - \mathbf{m}|^2 + (C_8 t^{-2} + C_9 t^{-3}) |\mathbf{x} - \mathbf{m}|^3 \\ & + (C_{10} t^{-2} + C_{11} t^{-3} + C_{12} t^{-4}) |\mathbf{x} - \mathbf{m}|^4. \end{aligned}$$

Notice that $(\mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \Sigma)} [|\mathbf{X} - \mathbf{m}|])^2 \leq \mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \Sigma)} [|\mathbf{X} - \mathbf{m}|^2]$.
 $\mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \Sigma)} [|\mathbf{X} - \mathbf{m}|^2] = \text{tr}(\Sigma)$, $\mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \Sigma)} [|\mathbf{X} - \mathbf{m}|^4] = \text{tr}^2(\Sigma) + 2\text{tr}(\Sigma^2)$,
 $\mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \Sigma)} [|\mathbf{X} - \mathbf{m}|^6] = \text{tr}^3(\Sigma) + 6\text{tr}(\Sigma)\text{tr}(\Sigma^2) + 8\text{tr}(\Sigma^3)$. Thus,

$$\begin{aligned} \int_{\mathbb{R}^d} |r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2 d\mathbf{x} &= (4\pi)^{-\frac{d}{2}} (\det(\Sigma))^{-\frac{1}{2}} \mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \frac{\Sigma}{2})} [|\tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2] \\ &\leq (4\pi)^{-\frac{d}{2}} (\det(\Sigma))^{-\frac{1}{2}} \left(C_0 + C_1 t^{-1} + C_2 t^{-2} + (C_3 t^{-1} + C_4 t^{-2}) \sqrt{\frac{\text{tr}(\Sigma)}{2}} \right. \\ &\quad \left. + \frac{1}{2} (C_5 t^{-1} + C_6 t^{-2} + C_7 t^{-3}) \text{tr}(\Sigma) + (C_8 t^{-2} + C_9 t^{-3}) \sqrt{\frac{1}{8} \text{tr}^3(\Sigma) + \frac{3}{4} \text{tr}(\Sigma) \text{tr}(\Sigma^2) + \text{tr}(\Sigma^3)} \right. \\ &\quad \left. + (C_{10} t^{-2} + C_{11} t^{-3} + C_{12} t^{-4}) \left(\frac{1}{4} \text{tr}^2(\Sigma) + \frac{1}{2} \text{tr}(\Sigma^2) \right) \right) = \mathcal{O}(t^{-d/2-2}) \text{ as } t \rightarrow 0. \end{aligned}$$

Similarly,

$$\begin{aligned} \int_{\mathbb{R}^d} |r(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2 \tilde{p} d\mathbf{x} &= (2\sqrt{3}\pi)^{-d} (\det(\Sigma))^{-1} \mathbb{E}_{\mathbf{X} \sim \mathcal{N}(\mathbf{x}; \mathbf{m}, \frac{\Sigma}{3})} [|\tilde{r}(\mathbf{x}, \mathbf{x}_0, t; \tilde{p})|^2] \\ &\leq (2\sqrt{3}\pi)^{-d} (\det(\Sigma))^{-1} \left(C_0 + C_1 t^{-1} + C_2 t^{-2} + (C_3 t^{-1} + C_4 t^{-2}) \sqrt{\frac{\text{tr}(\Sigma)}{3}} \right. \\ &\quad \left. + \frac{1}{3} (C_5 t^{-1} + C_6 t^{-2} + C_7 t^{-3}) \text{tr}(\Sigma) + (C_8 t^{-2} + C_9 t^{-3}) \sqrt{\frac{1}{27} \text{tr}^3(\Sigma) + \frac{2}{9} \text{tr}(\Sigma) \text{tr}(\Sigma^2) + \frac{8}{27} \text{tr}(\Sigma^3)} \right. \\ &\quad \left. + (C_{10} t^{-2} + C_{11} t^{-3} + C_{12} t^{-4}) \left(\frac{1}{9} \text{tr}^2(\Sigma) + \frac{2}{9} \text{tr}(\Sigma^2) \right) \right) = \mathcal{O}(t^{-(d+2)}) \text{ as } t \rightarrow 0. \quad \square \end{aligned}$$

REFERENCES

- [1] Z. ALDIRANY, R. COTTEREAU, M. LAFOREST, AND S. PRUDHOMME, Operator approximation of the wave equation based on deep learning of Green's function, *Computers & Mathematics with Applications*, 159 (2024), pp. 21–30.

- [2] R. BAPTISTA, B. HOSSEINI, N. B. KOVACHKI, AND Y. M. MARZOUK, Conditional sampling with monotone GANs: From generative models to likelihood-free inference, SIAM/ASA Journal on Uncertainty Quantification, 12 (2024), pp. 868–900.
- [3] N. M. BOFFI AND E. VANDEN-EIJNDEN, Probability flow solution of the Fokker–Planck equation, Machine Learning: Science and Technology, 4 (2023), p. 035012.
- [4] V. I. BOGACHEV, M. RÖCKNER, AND S. V. SHAPOSHNIKOV, Distances between transition probabilities of diffusions and applications to nonlinear Fokker–Planck–Kolmogorov equations, Journal of Functional Analysis, 271 (2016), pp. 1262–1300.
- [5] N. BOULLÉ, S. KIM, T. SHI, AND A. TOWNSEND, Learning Green’s functions associated with time-dependent partial differential equations, Journal of Machine Learning Research, 23 (2022), pp. 1–34.
- [6] J. BRUNA, B. PEHERSTORFER, AND E. VANDEN-EIJNDEN, Neural Galerkin schemes with active learning for high-dimensional evolution equations, Journal of Computational Physics, 496 (2024), p. 112588.
- [7] X. CHEN, L. YANG, J. DUAN, AND G. E. KARNIADAKIS, Solving inverse stochastic problems from discrete particle observations using the Fokker–Planck equation and physics-informed neural networks, SIAM Journal on Scientific Computing, 43 (2021), pp. B811–B830.
- [8] A. DAW, J. BU, S. WANG, P. PERDIKARIS, AND A. KARPATNE, Mitigating propagation failures in physics-informed neural networks using retain-resample-release (r3) sampling, Proceedings of Machine Learning Research, 202 (2023), pp. 7264–7302.
- [9] W. DENG, Finite element method for the space and time fractional Fokker–Planck equation, SIAM journal on numerical analysis, 47 (2009), pp. 204–226.
- [10] L. DINH, J. SOHL-DICKSTEIN, AND S. BENGIO, Density estimation using Real NVP, in International Conference on Learning Representations, 2017.
- [11] W. E AND Y. B., The deep Ritz method: A deep learning-based numerical algorithm for solving variational problems, Communications in Mathematics and Statistics, 6 (2018), pp. 1–12.
- [12] N. EL BEKRI, L. DRUMETZ, AND F. VERMET, FlowKac: An efficient neural fokker-planck solver using temporal normalizing flows and the feynman-kac formula, Transactions on Machine Learning Research (TMLR), (2025).
- [13] X. FENG, L. ZENG, AND T. ZHOU, Solving time dependent Fokker-Planck equations via temporal normalizing flow, Communications in Computational Physics, 32 (2022), pp. 401–423.
- [14] G. B. FOLLAND, Real analysis: modern techniques and their applications, John Wiley & Sons, 1999.
- [15] N. GABY, X. YE, AND H. ZHOU, Neural control of parametric solutions for high-dimensional evolution PDEs, SIAM Journal on Scientific Computing, 46 (2024), pp. C155–C185.
- [16] G. L. GILARDONI, On Pinsker’s and Vajda’s type inequalities for Csiszár’s f -divergences, IEEE Transactions on Information Theory, 56 (2010), pp. 5377–5386.
- [17] A. GRETTON, K. M. BORGWARDT, M. J. RASCH, B. SCHÖLKOPF, AND A. SMOLA, A kernel two-sample test, The journal of machine learning research, 13 (2012), pp. 723–773.
- [18] J. HE, Q. LIAO, AND X. WAN, Adaptive deep density approximation for stochastic dynamical systems, Journal of Scientific Computing, 102 (2025), p. 57.
- [19] Z. HU, Z. ZHANG, G. E. KARNIADAKIS, AND K. KAWAGUCHI, Score-based physics-informed neural networks for high-dimensional Fokker–Planck equations, SIAM Journal on Scientific Computing, 47 (2025), pp. C680–C705.
- [20] T. HYTONEN, J. VAN NEERVEN, M. VERAAR, AND L. WEIS, Analysis in Banach spaces, volume i: Martingales and littlewood-paley theory, Ergebnisse der Mathematik und ihrer Grenzgebiete, 3 (2016).
- [21] R. JORDAN, D. KINDERLEHRER, AND F. OTTO, The variational formulation of the Fokker–Planck equation, SIAM journal on mathematical analysis, 29 (1998), pp. 1–17.
- [22] P. KUMAR AND S. NARAYANAN, Solution of Fokker-Planck equation by finite element and finite difference methods for nonlinear systems, Sadhana, 31 (2006), pp. 445–461.
- [23] W. LEE, L. WANG, AND W. LI, Deep JKO: Time-implicit particle methods for general nonlinear gradient flows, Journal of Computational Physics, 514 (2024), p. 113187.
- [24] M. LEUTBECHER AND T. N. PALMER, Ensemble forecasting, Journal of computational physics, 227 (2008), pp. 3515–3539.
- [25] Z. LI, N. B. KOVACHKI, K. AZIZZADENESHELI, B. LIU, K. BHATTACHARYA, A. STUART, AND A. ANANDKUMAR, Fourier neural operator for parametric partial differential equations, in International Conference on Learning Representations.
- [26] S. LIU, W. LI, H. ZHA, AND H. ZHOU, Neural parametric Fokker–Planck equation, SIAM Journal on Numerical Analysis, 60 (2022), pp. 1385–1449.
- [27] L. LU, P. JIN, G. PANG, Z. ZHANG, AND G. E. KARNIADAKIS, Learning nonlinear operators via

- DeepONet based on the universal approximation theorem of operators, *Nature machine intelligence*, 3 (2021), pp. 218–229.
- [28] Y. LU AND B. HUANG, Structured output learning with conditional generative flows, in *Proceedings of the AAAI conference on artificial intelligence*, vol. 34, 2020, pp. 5005–5012.
 - [29] Y. LU, R. MAULIK, T. GAO, F. DIETRICH, I. G. KEVREKIDIS, AND J. DUAN, Learning the temporal evolution of multivariate densities via normalizing flows, *Chaos: An Interdisciplinary Journal of Nonlinear Science*, 32 (2022).
 - [30] R. N. MILLER, E. F. CARTER JR, AND S. T. BLUE, Data assimilation into nonlinear stochastic models, *Tellus A*, 51 (1999), pp. 167–194.
 - [31] H. NOH, J. LEE, AND E. YOON, FPL-net: A deep learning framework for solving the nonlinear Fokker–Planck–Landau collision operator for anisotropic temperature relaxation, *Journal of Computational Physics*, 523 (2025), p. 113665.
 - [32] L. PICHLER, A. MASUD, AND L. A. BERGMAN, Numerical solution of the Fokker–Planck equation by finite difference and finite element methods—a comparative study, in *Computational Methods in Stochastic Dynamics*, Springer, 2013, pp. 69–85.
 - [33] M. RAISSI, P. PERDIKARIS, AND G. E. KARNIADAKIS, Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations, *Journal of Computational Physics*, 378 (2019), pp. 686–707.
 - [34] F. SANTAMBROGIO, Optimal transport for applied mathematicians, Birkhäuser, NY, 55 (2015), p. 94.
 - [35] R. SAPORITI AND F. NOBILE, Neural Galerkin normalizing flow for transition probability density functions of diffusion models, arXiv preprint arXiv:2603.18907, (2026).
 - [36] R. SAPORITI AND F. NOBILE, Neural Galerkin normalizing flows for bayesian inference of diffusions with inaccessible boundaries, arXiv preprint arXiv:2606.04324, (2026).
 - [37] J. SIRIGNANO AND K. SPILIOPOULOS, DGM: A deep learning algorithm for solving partial differential equations, *Journal of Computational Physics*, 375 (2018), pp. 1339–1364.
 - [38] K. TANG, X. WAN, AND Q. LIAO, Deep density estimation via invertible block-triangular mapping, *Theoretical and Applied Mechanics Letters*, 10 (2020), pp. 143–148.
 - [39] K. TANG, X. WAN, AND Q. LIAO, Adaptive deep density approximation for Fokker–Planck equations, *Journal of Computational Physics*, 457 (2022), p. 111080.
 - [40] Y. TENG, X. ZHANG, Z. WANG, AND L. JU, Learning Green’s functions of linear reaction-diffusion equations with application to fast numerical solver, in *Mathematical and Scientific Machine Learning*, PMLR, 2022, pp. 1–16.
 - [41] S. WANG, H. WANG, AND P. PERDIKARIS, On the eigenvector bias of Fourier feature networks: From regression to solving multi-scale PDEs with physics-informed neural networks, *Computer Methods in Applied Mechanics and Engineering*, 384 (2021), p. 113938.
 - [42] T. WANG, Z. HU, K. KAWAGUCHI, Z. ZHANG, AND G. E. KARNIADAKIS, Tensor neural networks for high-dimensional Fokker–Planck equations, *Neural Networks*, 185 (2025), p. 107165.
 - [43] X. WANG, J. FENG, Q. LIU, C. TAN, Y. LIU, AND Y. XU, A deep learning framework for jointly solving transient fokker-planck equations with arbitrary parameters and initial distributions, arXiv preprint arXiv:2604.06001, (2026).
 - [44] Z. O. WANG, R. BAPTISTA, Y. MARZOUK, L. RUTHOTTO, AND D. VERMA, Efficient neural network approaches for conditional optimal transport with applications in Bayesian inference, *SIAM Journal on Scientific Computing*, 47 (2025), pp. C979–C1005.
 - [45] M. F. WEHNER AND W. WOLFER, Numerical evaluation of path-integral solutions to Fokker–Planck equations, *Physical Review A*, 27 (1983), p. 2663.
 - [46] Y. WEN, E. VANDEN-ELJNDEN, AND B. PEHERSTORFER, Coupling parameter and particle dynamics for adaptive sampling in neural Galerkin schemes, *Physica D: Nonlinear Phenomena*, 462 (2024), p. 134129.
 - [47] Y. XU, H. ZHANG, Y. LI, K. ZHOU, Q. LIU, AND J. KURTHS, Solving Fokker–Planck equation using deep learning, *Chaos: An Interdisciplinary Journal of Nonlinear Science*, 30 (2020).
 - [48] M. YANG, P. WANG, D. DEL CASTILLO-NEGRETTE, Y. CAO, AND G. ZHANG, A pseudoreversible normalizing flow for stochastic dynamical systems with various initial distributions, *SIAM Journal on Scientific Computing*, 46 (2024), pp. C508–C533.
 - [49] L. ZENG, X. WAN, AND T. ZHOU, Adaptive deep density approximation for fractional Fokker–Planck equations, *Journal of Scientific Computing*, 97 (2023), p. 68.
 - [50] L. ZENG, X. WAN, AND T. ZHOU, Bounded KRnet and its applications to density estimation and approximation, *SIAM Journal on Scientific Computing*, 47 (2025), pp. C1294–C1318.
 - [51] J. ZHAI, M. DOBSON, AND Y. LI, A deep learning method for solving Fokker–Planck equations, in *Mathematical and Scientific Machine Learning*, PMLR, 2022, pp. 568–597.
 - [52] M. ZHOU, S. OSHER, AND W. LI, Simulating Fokker–Planck equations via mean field control of score-based normalizing flows, arXiv preprint arXiv:2506.05723, (2025).