

Journal Pre-proof

Deep JKO: time-implicit particle methods for general nonlinear gradient flows

Wonjun Lee, Li Wang and Wuchen Li

PII: S0021-9991(24)00436-4
DOI: <https://doi.org/10.1016/j.jcp.2024.113187>
Reference: YJCPH 113187

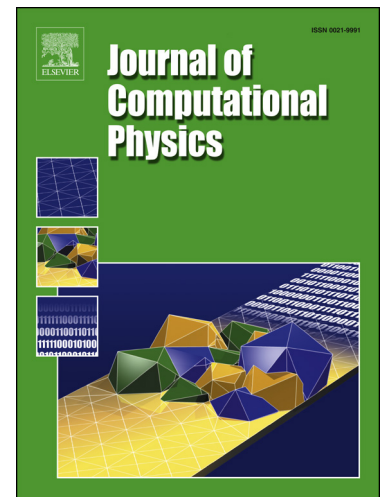
To appear in: *Journal of Computational Physics*

Received date: 16 November 2023
Revised date: 14 May 2024
Accepted date: 5 June 2024

Please cite this article as: W. Lee, L. Wang and W. Li, Deep JKO: time-implicit particle methods for general nonlinear gradient flows, *Journal of Computational Physics*, 113187, doi: <https://doi.org/10.1016/j.jcp.2024.113187>.

This is a PDF file of an article that has undergone enhancements after acceptance, such as the addition of a cover page and metadata, and formatting for readability, but it is not yet the definitive version of record. This version will undergo additional copyediting, typesetting and review before it is published in its final form, but we are providing this version to give early visibility of the article. Please note that, during the production process, errors may be discovered which could affect the content, and all legal disclaimers that apply to the journal pertain.

© 2024 Published by Elsevier.



Highlights

- A neural network empowered implicit particle method is proposed for computing high dimensional nonlinear gradient flows.
- Our method benefits from the structure preservation properties inherent in the JKO formulation and can handle high-dimensional problems due to its mesh-free discretization.
- We consider general nonlinear functionals and nonlinear mobility, which are currently beyond the scope of large-scale gradient flow solvers.
- Concepts from continuous normalizing flow are employed for efficient density computation, and an explicit recurrence relation for derivative computation is utilized to significantly streamline the backpropagation process.
- We have showcased the effectiveness and versatility of our method in both high dimensions and diverse equations, including one arising from an inverse problem—the Kalman-Wasserstein gradient flow.

Deep JKO: time-implicit particle methods for general nonlinear gradient flows

Wonjun Lee^a, Li Wang^b, Wuchen Li^c

^a*Institute for Mathematics and Its Applications, University of Minnesota, U.S.*

^b*School of Mathematics, University of Minnesota, U.S.*

^c*Department of Mathematics, University of South Carolina, U.S.*

Abstract

We develop novel neural network-based implicit particle methods to compute high-dimensional Wasserstein-type gradient flows with linear and nonlinear mobility functions. The main idea is to use the Lagrangian formulation in the Jordan–Kinderlehrer–Otto (JKO) framework, where the velocity field is approximated using a neural network. We leverage the formulations from the neural ordinary differential equation (neural ODE) in the context of continuous normalizing flow for efficient density computation. Additionally, we make use of an explicit recurrence relation for computing derivatives, which greatly streamlines the backpropagation process. Our methodology demonstrates versatility in handling a wide range of gradient flows, accommodating various potential functions and nonlinear mobility scenarios. Extensive experiments demonstrate the efficacy of our approach, including an illustrative example from Bayesian inverse problems. This underscores that our scheme provides a viable alternative solver for the Kalman–Wasserstein gradient flow.

Keywords: Nonlinear gradient flows, JKO scheme, Neural ODE, Particle methods, Density estimation, Nonlinear mobility.

1. Introduction

Deep learning has revolutionized many areas, fundamentally reshaping fields such as natural language processing, visual recognition, and beyond [27]. It has also emerged as the cornerstone of contemporary scientific computing. By harnessing the power of deep neural networks and their ability to approximate complex functions, it has become an invaluable tool for approximating solutions to partial differential equations (PDEs), particularly those of high dimensions. These equations often defy conventional computational techniques due to their complexity and dimensionality.

Several methodologies have employed neural network approximations for solving PDEs. One approach is the deep Ritz method, crafted upon the variational formulation of elliptic-type equations [44]. Another avenue is the deep

backward stochastic differential equation (BSDE) method, based on the probabilistic and control-based formulation of parabolic equations [23]. A third technique involves the deep Galerkin method [38] and the physics-informed neural network (PINN) [35]. These methodologies both hinge on the minimization of the (L^2) residual of the equation. Importantly, they possess the capability to be employed across a wide spectrum of equation types. In parallel, this paper aims to develop a versatile solver for a substantial class of nonlinear PDEs characterized by a gradient flow structure.

Gradient flows describe the evolution of a function or a probability measure over time, with the rate of change at any point being proportional to the negative gradient of the energy with respect to the function's value. They were originally used to study physical or biological systems that naturally obey some form of least action law. However, they have since emerged as powerful tools in various domains, including sampling, variational inference, the theory of neural networks, and more.

To be more specific, we consider the continuity equation of the form:

$$\partial_t \rho = -\nabla_{\mathbf{x}} \cdot (\rho \mathbf{v}) := \nabla_{\mathbf{x}} \cdot [\rho \nabla_{\mathbf{x}} (U'(\rho) + V + W * \rho)], \quad \rho(0, \cdot) = \rho_0. \quad (1)$$

Here, $\rho(t, \mathbf{x})$, with $\mathbf{x} \in \mathbb{R}^d$, represents the particle density function, ρ_0 is the initial density function, $U(\rho) \in \mathbb{R}$ represents the internal potential, $V(\mathbf{x}) \in \mathbb{R}$ is a drift potential, $W(\mathbf{x}, \mathbf{y}) = W(\mathbf{y}, \mathbf{x}) \in \mathbb{R}$ is an interaction potential involving $\mathbf{x}$ and $\mathbf{y}$ in $\mathbb{R}^d$, and $*$ denotes the convolution operator. Without loss of generality, we assume that U , V and W are sufficient regular in the sense that $V, W \in C(\mathbb{R})$, and $U \in C^1((0, \infty))$ with $\lim_{r \rightarrow \infty} \frac{U(r)}{r} = +\infty$ and $U(0) = 0$. See [10, Section 3.1] for more discussion.

This equation can be viewed as the gradient flow of the energy functional [40, 2]

$$\mathcal{E}(\rho) = \int_{\mathbb{R}^d} [U(\rho(\mathbf{x})) + V(\mathbf{x})\rho(\mathbf{x})] d\mathbf{x} + \frac{1}{2} \int_{\mathbb{R}^d \times \mathbb{R}^d} W(\mathbf{x} - \mathbf{y})\rho(\mathbf{x})\rho(\mathbf{y}) d\mathbf{x} d\mathbf{y} \quad (2)$$

in Wasserstein-2 space, which is the space of probability measures equipped with Wasserstein-2 distance, denoted as $\mathcal{W}_2$.

Among all the recent endeavors in devising structure-preserving methodologies for (1), our focus in this paper will be exclusively on the Jordan-Kinderlehrer-Otto (JKO) scheme [26]. This scheme centers on the task of determining $\rho^{k+1}(\cdot) \approx \rho(t^{k+1}, \cdot)$ with $t^{k+1} = (k+1)\Delta t$, given the approximation $\rho^k(\cdot) \approx \rho(t^k, \cdot)$:

$$\rho^{k+1} \in \arg \min_{\rho} \{ \mathcal{W}_2(\rho, \rho^k)^2 + 2\Delta t \mathcal{E}(\rho) \}. \quad (3)$$

Here, $\mathcal{W}_2(\rho, \rho^k)$ denotes the Wasserstein-2 distance between ρ and ρ^k . The scheme in (3) can be conceptualized as a time-implicit version of the fundamental gradient descent, with the expression on the right-hand side of (3) being recognized as the Wasserstein proximal operator associated with energy functional $\mathcal{E}$.

In general terms, there exist two methods for expressing the $\mathcal{W}_2$ distance, each leading to distinct implementations of (3). The first method is the Eulerian representation, which establishes a connection between two densities through a continuity equation involving a velocity field that minimizes kinetic energy during the transport. In contrast, the second method is grounded in the Lagrangian formulation [11]. This approach explores the diffeomorphism that maps one density onto another, a technique that has also gained prominence in modern machine learning applications.

In this paper, we adopt the Lagrangian viewpoint by employing the particle method, coupled with principles drawn from neural ordinary differential equations (neural ODEs). To elaborate, starting from a set of particles $\{\mathbf{x}_j^0\}$ that are i.i.d. samples from ρ^0 , we seek a sequence of transformations governing the update of all particles' positions as follows:

$$\{\mathbf{x}_j^0\} \xrightarrow{\mathbf{T}^1} \{\mathbf{x}_j^1\} \xrightarrow{\mathbf{T}^2} \{\mathbf{x}_j^2\} \cdots$$

Here, $\mathbf{x}_j^k = \mathbf{T}^k(\mathbf{x}_j^{k-1})$. As a consequence, the density evolves in accordance with $\rho^k = \mathbf{T}^k \# \rho^{k-1}$ for $k = 1, 2, \dots$, where $\#$ denotes the pushforward operator. We then approximate each map $\mathbf{T}^k$ by parameterizing the corresponding vector field using a deep neural network. Thanks to the Jacobi's formula, this results in a simple update in estimating the density ρ^k ; see (4c) below.

More precisely, our primary formulation is as follows. To transit from time t^k to t^{k+1} , we solve the following minimization problem:

$$\begin{aligned} \theta^* = & \arg \min_{\theta \in \Theta} \mathbb{E}_{\mathbf{x} \sim \rho} \left[\int_0^1 |\mathbf{v}_\theta(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2 d\tau + 2\Delta t \left(\frac{U(\rho(\mathbf{T}(1, \mathbf{x})))}{\rho(\mathbf{T}(1, \mathbf{x}))} + V(\mathbf{T}(1, \mathbf{x})) \right) \right] \\ & + \Delta t \mathbb{E}_{(\mathbf{x}, \mathbf{y}) \sim \rho \otimes \rho} [W(\mathbf{T}(1, \mathbf{x}), \mathbf{T}(1, \mathbf{y}))], \end{aligned} \quad (4a)$$

such that $\mathbf{T}(\tau, \mathbf{x})$ satisfies the dynamical constraint:

$$\frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}) = \mathbf{v}_\theta(\tau, \mathbf{T}(\tau, \mathbf{x})), \quad \mathbf{T}(0, \mathbf{x}) = \mathbf{x}, \quad (4b)$$

and the density $\rho(\mathbf{T}(1, \mathbf{x}))$ is computed as first evolving:

$$\frac{d}{d\tau} \log |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})| = \operatorname{div}(\mathbf{v}_\theta)(\tau, \mathbf{T}(\tau, \mathbf{x})), \quad \rho(\mathbf{T}(0, \mathbf{x})) = \rho^k(\mathbf{x}), \quad (4c)$$

followed by the pushforward relation:

$$\rho(\mathbf{T}(1, \mathbf{x})) = \rho^k(\mathbf{x}) / |\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x})|. \quad (4d)$$

In this problem, we employ Benamou-Brenier's dynamic formulation [1] to calculate the Wasserstein distance and reformulate the continuity equation using the flow map $\mathbf{T}(\tau, \mathbf{x})$, with τ serving as an artificial time. The associated

velocity field is denoted as $\mathbf{v}_\theta$, where $\theta \in \Theta$ represents the parameters of a neural network. The notation $\mathbb{E}_{\mathbf{x} \sim \rho}$ denotes the expectation value for $\mathbf{x}$ distributed according to the probability density ρ , and $\mathbb{E}(\mathbf{x}, \mathbf{y}) \sim \rho \otimes \rho$ represents the expectation over pairs of independent variables $(\mathbf{x}, \mathbf{y})$, where $\mathbf{x}$ and $\mathbf{y}$ are independent and satisfy the joint density $\rho \otimes \rho$. Equation (4c) is a significant outcome from the perspective of neural ODEs, providing a straightforward means to compute $|\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})|$. Consequently, this facilitates an efficient computation of density as outlined in (4d).

It's crucial to highlight that the proposed method goes beyond the current scope of merely approximating Wasserstein gradient flows. Similar approaches, as outlined in (4), can be employed for general gradient flows, including non-linear mobility Wasserstein and Kalman-Wasserstein gradient flows. In such cases, the identity (4c), utilized in neural ODE density estimations, takes on even greater significance. This is due to the fact that the objective function in (4) becomes dependent on the density, introducing additional nonlinearities into the optimization objective (loss function), see (24) in the subsequent context.

The exploration of computing (1) using the JKO scheme has experienced rapid growth. Previous endeavors employed finite difference or finite volume schemes [2, 10, 29, 41, 14, 19, 7], as well as Lagrangian methods [11, 13, 12]. More recently, efforts have been directed towards developing machine learning-based methods, particularly in tackling high-dimensional challenges. In this context, neural networks have been utilized to approximate maps or density functions [31, 18, 33, 25, 37, 4]. Furthermore, a variety of methods have arisen to approximate different formulations of one-step JKO schemes. These range from leveraging generative adversarial networks [30] to employing neural ODEs [36]. These approaches primarily target at sampling of high-dimensional distributions in machine learning applications [39, 42], manifesting as time-implicit updates of the projected dynamics within neural network spaces, specifically the natural gradient dynamics [30].

Contrasting previous methodologies, our approach employs neural ODE techniques to compute the time-implicit update of generalized Wasserstein-type gradient flows. This innovative strategy facilitates the streamlined computation of high-dimensional densities in Lagrangian coordinates. It is worth highlighting the similarities and disparities between our approach and its closest counterparts. One such sibling is the blob method [9]. While both approaches rely on particles and intrinsically exhibit energy dissipation, the dissipation in the blob method hinges on the chosen time step, whereas ours remains unconditionally stable. Furthermore, our method eliminates the need for kernel density estimation, a significant computational bottleneck present in the blob method. Another approach involves the direct use of a neural network to approximate the mapping between ρ^k and ρ in (3) [33, 25]. This approach can be seen as a special case of our method when the inner time step τ is set to 1. A critical distinction, however, lies in our utilization of the density evolution equation (4c) from the perspective of neural ODEs and the explicit formula (28) for computing derivatives, which significantly enhances the efficiency of density estimation. Furthermore, our approach can be readily extended to accommodate the non-

linear mobility case. A third comparable technique is the self-consistent velocity matching approach [4, 37]. In this method, neural networks are also employed to approximate the velocity field. However, instead of sequentially learning it within the JKO formulation as we do, this approach adopts an end-to-end learning strategy. While this method might lead to faster computation, it doesn't inherently preserve energy decay in general.

This paper is organized as follows. In section 2, we review the general formulation of Wasserstein gradient flows in Lagrangian coordinates and introduce the Lagrangian Wasserstein proximal operator and its particle version for time-implicit (i.e., JKO) updates of gradient flows. We expand our framework to include gradient flows characterized by metrics reliant on general nonlinear mobility functions. Section 3 outlines our proposition to approximate the push-forward map using the trajectory of neural ODEs. This approach empowers us to employ neural networks with diverse depths while simultaneously sidestepping the computationally intensive task of back-propagation. To illustrate the effectiveness and versatility of our approach, we present a series of numerical examples in Section 4, including Fokker-Planck equations, porous media equations, Kalman-Wasserstein gradient flow of Kullback-Leibler (KL) divergence, and nonlinear mobility gradient flows. Finally, the paper concludes in Section 5 with an overview of our findings and insights into future research directions.

2. Lagrangian JKO formulation

In this section, we present the main formulation of this paper. We first review the Lagrangian coordinates of Wasserstein gradient flows and then introduce our variational time-implicit schemes. One major component of our formulation is to use instantaneous change of variable formula (see (13)) in the continuous normalizing flow to achieve efficient density computation. We also consider the generalized nonlinear gradient flows, in which the Wasserstein-type metric is associated with nonlinear mobility functions.

2.1. Dynamic JKO schemes in Lagrangian coordinates

Recall the dynamic JKO formulation to (1) as follows. Given $\rho^n(\mathbf{x})$, then one can obtain $\rho^{n+1}(\mathbf{x})$ as $\rho(1, \mathbf{x})$ with $\rho(\tau, \mathbf{x})$ solving

$$\begin{cases} (\rho, \mathbf{v}) = \arg \inf_{\rho, \mathbf{v}} \int_0^1 \int_{\mathbb{R}^d} \rho |\mathbf{v}|^2 d\mathbf{x} d\tau + 2\Delta t \mathcal{E}(\rho(1, \cdot)) \\ \text{s.t. } \partial_\tau \rho + \nabla \cdot (\rho \mathbf{v}) = 0, \rho(0, \mathbf{x}) = \rho^n(\mathbf{x}). \end{cases} \quad (5)$$

In the above formulation, the minimization is over the probability density function $\rho(\tau, \mathbf{x})$, $\tau \in (0, 1)$, terminal time density $\rho(1, \mathbf{x})$, and vector fields $\mathbf{v}(\tau, \mathbf{x})$, such that the continuity equation holds.

Now we translate the variational problem (5) into the Lagrangian formulation. Assume the velocity field is sufficiently regular, then the solution $\rho(\tau, \mathbf{x})$ to the constrained continuity equation can be written as

$$\rho(\tau, \cdot) = \mathbf{T}(\tau, \cdot) \# \rho^n, \quad (6)$$

where $\mathbf{T}$ is the flow map that solves the following ODE:

$$\frac{d}{d\tau}\mathbf{T}(\tau, \mathbf{x}) = \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x})), \quad \mathbf{T}(0, \mathbf{x}) = \mathbf{x}. \quad (7)$$

Note that (6) implies that for any integrable test function ϕ , we have that

$$\int_{\mathbb{R}^d} \phi(\mathbf{x}) \rho(\tau, \mathbf{x}) d\mathbf{x} = \int_{\mathbb{R}^d} \phi(\mathbf{T}(\tau, \mathbf{x})) \rho^n(\mathbf{x}) d\mathbf{x}.$$

Then variational problem (5) rewrites as

$$\begin{cases} \min_{\mathbf{v}} \int_0^1 \int_{\mathbb{R}^d} \rho^n(\mathbf{x}) |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2 d\tau d\mathbf{x} + 2\Delta t \mathcal{E}(\mathbf{T}(1, \cdot) \# \rho^n) \\ \text{s.t.} \quad \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}) = \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x})), \quad \mathbf{T}(0, \mathbf{x}) = \mathbf{x}. \end{cases} \quad (8)$$

To proceed, it is necessary to derive an explicit formulation of $\mathcal{E}(\mathbf{T}(1, \cdot) \# \rho^n)$ that can be interpreted as an expectation of a functional, allowing for its representation using particles in the subsequent discussion. In the following, we present a comprehensive formulation for each component of equation (2). Specifically, the external potential and interaction potential have straightforward expectation formulations through the push forward relations, whereas the internal energy necessitates further reformulation.

External potential energy

$$\int_{\mathbb{R}^d} V(\mathbf{x}) \rho(1, \mathbf{x}) d\mathbf{x} = \int_{\mathbb{R}^d} V(\mathbf{T}(1, \mathbf{x})) \rho^n(\mathbf{x}) d\mathbf{x}. \quad (9)$$

Interaction energy

$$\begin{aligned} & \int_{\mathbb{R}^d \times \mathbb{R}^d} W(\mathbf{x}, \mathbf{y}) \rho(1, \mathbf{x}) \rho(1, \mathbf{y}) d\mathbf{x} d\mathbf{y} \\ &= \int_{\mathbb{R}^d \times \mathbb{R}^d} W(\mathbf{T}(1, \mathbf{x}), \mathbf{T}(1, \mathbf{y})) \rho^n(\mathbf{x}) \rho^n(\mathbf{y}) d\mathbf{x} d\mathbf{y}. \end{aligned} \quad (10)$$

Internal energy/entropy

$$\begin{aligned} & \int_{\mathbb{R}^d} U(\rho(1, \mathbf{x})) d\mathbf{x} \\ &= \int_{\mathbb{R}^d} U(\rho(1, \mathbf{T}(1, \mathbf{x}))) |\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x})| d\mathbf{x} \\ &= \int_{\mathbb{R}^d} U(\rho(1, \mathbf{T}(1, \mathbf{x}))) \frac{|\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x})|}{\rho^n(\mathbf{x})} \rho^n(\mathbf{x}) d\mathbf{x} \\ &= \int_{\mathbb{R}^d} U(\rho(1, \mathbf{T}(1, \mathbf{x}))) \frac{1}{\rho(1, \mathbf{T}(1, \mathbf{x}))} \rho^n(\mathbf{x}) d\mathbf{x}, \end{aligned} \quad (11)$$

where we have used the Monge-Amperè equation

$$\rho(1, \mathbf{T}(1, \mathbf{x})) |\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x})| = \rho^n(\mathbf{x}) \quad (12)$$

thanks to (6). In practice, computing the determinant of the Jacobian $\nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x})$ poses a significant bottleneck due to its cubic cost with respect to the dimension of $\mathbf{x}$. To address this issue, we draw inspiration from the concept of continuous normalizing flow and employ an efficient approach based on the instantaneous change of variable formula [15, 22]. The crucial element is the following formula.

Proposition 1.

$$\frac{d}{d\tau} \log |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})| = \text{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{x})). \quad (13)$$

Proof. Equation (13) is derived through a calculus that involves both ODE (7) and the Monge-Amperè equation (12). Assume that $\mathbf{T}(\tau, \mathbf{x})$ is invertible, i.e. $\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}) \neq 0$, then

$$\begin{aligned} \frac{d}{d\tau} \log |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})| &= \frac{\text{sgn}(\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}))}{|\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})|} \frac{d}{d\tau} (\det(\nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}))) \\ &= \text{tr} \left((\nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}))^{-1} \frac{d}{d\tau} \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}) \right) \\ &= \text{tr} \left((\nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}))^{-1} \nabla_{\mathbf{x}} \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}) \right) \\ &= \text{tr} \left((\nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}))^{-1} \nabla_{\mathbf{x}} \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x})) \right) \\ &= \text{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{x})), \end{aligned}$$

where the second equality uses the Jacobi's formula, and the last equality uses the chain rule, and $\text{sgn}(x) = \frac{x}{|x|}$. $\square$

As a result, the computational burden shifts towards calculating $\text{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{x}))$, which only involves a trace calculation and can be accomplished much more efficiently. In sum, (12) and (13) play a crucial role in enabling efficient density estimation, as elaborated in the next section and Section 3.1.

2.2. A Particle method

By interpreting the integral against $\rho(\mathbf{x})d\mathbf{x}$ as an expectation, (8) reveals a direct particle representation. More precisely, let $\{\mathbf{x}_j^n\}_{j=1}^N$ be particles sampled from $\rho^n(\mathbf{x})$, we discretize (8) as follows. Note here we omit the superscript n in $\mathbf{x}_j^n$.

Case 1 : (8) with (9) and/or (10)

$$\left\{ \begin{array}{l} \min_{\mathbf{v}} \frac{1}{N} \sum_{j=1}^N \left[\int_0^1 |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j))|^2 d\tau + 2\Delta t V(\mathbf{T}(1, \mathbf{x}_j)) \right. \\ \left. + \frac{2\Delta t}{N} \sum_{l=1}^N W(\mathbf{T}(1, \mathbf{x}_j), \mathbf{T}(1, \mathbf{x}_l)) \right] \\ \text{s.t. } \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}_j) = \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \mathbf{T}(0, \mathbf{x}_j) = \mathbf{x}_j. \end{array} \right. \quad (14)$$

Case 2 : (8) with (11)

$$\left\{ \begin{array}{l} \min_{\mathbf{v}} \frac{1}{N} \sum_{j=1}^N \left[\int_0^1 |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j))|^2 d\tau + 2\Delta t U \left(\frac{\rho^n(\mathbf{x}_j)}{|\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j)|} \right) \frac{|\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j)|}{\rho^n(\mathbf{x}_j)} \right] \\ \text{s.t. } \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}_j) = \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \mathbf{T}(0, \mathbf{x}_j) = \mathbf{x}_j \\ \frac{d}{d\tau} \log |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}_j)| = \text{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \log |\det \nabla_{\mathbf{x}} \mathbf{T}(0, \mathbf{x}_j)| = 0, \end{array} \right. \quad (15)$$

In cases where the energy solely comprises potential and interaction components, it suffices to monitor the particle positions. However, when internal energy is introduced, it becomes necessary also to track the density ρ , which is accomplished by monitoring the logarithmic determinant of the transport map.

More precisely, denote the minimizer of either problem $\mathbf{T}^{n+1}$, then

$$\mathbf{T}^{n+1}(1, \mathbf{x}_j^n) =: \mathbf{x}_j^{n+1},$$

which can be viewed as samples from $\rho^{n+1}(\mathbf{x})$. Therefore, starting from $\{\mathbf{x}_j^0\}$, we have a sequence of update

$$\{\mathbf{x}_j^0\} \xrightarrow{\mathbf{T}^1} \{\mathbf{x}_j^1\} \xrightarrow{\mathbf{T}^2} \{\mathbf{x}_j^2\} \xrightarrow{\mathbf{T}^3} \{\mathbf{x}_j^3\} \cdots, \quad (16)$$

where $\mathbf{x}_j^n = \mathbf{T}^n(1, \mathbf{x}_j^{n-1})$. Likewise, the density evolves as

$$\rho^0 \xrightarrow{\mathbf{T}^1} \rho^1 \xrightarrow{\mathbf{T}^2} \rho^2 \xrightarrow{\mathbf{T}^3} \rho^3 \cdots, \quad (17)$$

where $\rho^n = \mathbf{T}^n(1, \cdot) \# \rho^{n-1}$.

In practice, since we only evolve particles, a major difficulty in (17) lies in the *density estimation* of ρ^n , which can be very expensive and inaccurate in high dimensions. To resolve this issue, a key observation is that we don't need the full information of the density, but only the density evaluated along the trajectory of particles!

Suppose we are given a set of samples $\mathbf{x}_j^0$ drawn from the known analytical expression of ρ_0 . In the first JKO step, we can compute the updated density as follows:

$$\rho^1(\underbrace{\mathbf{T}^1(1, \mathbf{x}_j^0)}_{\mathbf{x}_j^1}) = \frac{\rho^0(\mathbf{x}_j^0)}{|\det \nabla_{\mathbf{x}} \mathbf{T}^1(1, \mathbf{x}_j^0)|}, \quad (18)$$

where $\mathbf{T}^1$ represents the map obtained in the first JKO step. Here, $\mathbf{x}_j^1$ corresponds to the transformed sample after applying $\mathbf{T}^1$ to $\mathbf{x}_j^0$. Following a similar procedure, we obtain the density in the subsequent JKO steps:

$$\rho^2(\underbrace{\mathbf{T}^2(1, \mathbf{x}_j^1)}_{\mathbf{x}_j^2}) = \frac{\rho^1(\mathbf{x}_j^1)}{|\det \nabla_{\mathbf{x}} \mathbf{T}^2(1, \mathbf{x}_j^1)|}, \quad (19)$$

and this iterative process can be continued for any subsequent JKO step.

Remark 1. We emphasize that this observation holds true for general non-linear internal potential functions, denoted as U . In the existing literature, U is commonly assumed to be the negative Boltzmann-Shannon entropy, given by $U(z) = z \log z$. In this particular case, the aforementioned difficulty can be circumvented. Indeed, since

$$U\left(\frac{\rho^n(\mathbf{x}_j)}{|\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j)|}\right) \frac{|\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j)|}{\rho^n(\mathbf{x}_j)} = \log \rho^n(\mathbf{x}_j) - \log |\det \nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j)|,$$

we can simplify the computation by avoiding the need to directly calculate $\rho^n(\mathbf{x}_j)$. This is because $\log \rho^n(\mathbf{x}_j)$ is separated from $\log |\det(\nabla_{\mathbf{x}} \mathbf{T}(1, \mathbf{x}_j))|$ in equation (15), enabling us to exclude it from the calculation. However, for general potential functions, the computation of $\rho^n(\mathbf{x}_j)$ cannot be circumvented.

2.3. Nonlinear mobility

In this section, we expand upon the previous method to encompass the generalized gradient flow, as delineated by the following equation:

$$\partial_t \rho = \nabla_{\mathbf{x}} \cdot \left(M(\rho) \nabla_{\mathbf{x}} \frac{\delta \mathcal{E}(\rho)}{\delta \rho} \right). \quad (20)$$

Here, the function $M(\rho) \geq 0$ represents a nonlinear mobility function. This type of equation appears in various contexts, such as thin films [3], phase separation in binary alloys described by the Cahn-Hilliard equation [6], and the transport phenomena of biological systems with overcrowding prevention [5], among others.

For general mobility function $M(\rho)$, equation (20) can also be interpreted as a gradient flow, especially when $M(\rho)$ is concave and satisfies

some other properties in [16, 8]. The generalized metric also admits a dynamic formulation, leading to the following extension of the dynamic JKO scheme (5):

$$\begin{cases} (\rho, \tilde{\mathbf{v}}) = \arg \inf_{\rho, \tilde{\mathbf{v}}} \int_0^1 \int_{\mathbb{R}^d} M(\rho) |\tilde{\mathbf{v}}|^2 d\mathbf{x} d\tau + 2\Delta t \mathcal{E}(\rho(1, \cdot)) \\ \text{s.t. } \partial_\tau \rho + \nabla \cdot (M(\rho) \tilde{\mathbf{v}}) = 0, \rho(0, \mathbf{x}) = \rho^n(\mathbf{x}). \end{cases} \quad (21)$$

In the above formulation, the minimization is over density function $\rho(\tau, \mathbf{x})$ for $\tau \in (0, 1)$, terminal time density $\rho(1, \mathbf{x})$, and vector fields $\tilde{\mathbf{v}}(\tau, \mathbf{x})$, such that the nonlinear mobility induced continuity equation holds.

Let $\mathbf{v} = \frac{M(\rho)}{\rho} \tilde{\mathbf{v}}$, equation (21) can be rewritten into

$$\begin{cases} (\rho, \mathbf{v}) = \arg \inf_{\rho, \mathbf{v}} \int_0^1 \int_{\mathbb{R}^d} \frac{\rho^2}{M(\rho)} |\mathbf{v}|^2 d\mathbf{x} d\tau + 2\Delta t \mathcal{E}(\rho(1, \cdot)) , \\ \text{s.t. } \partial_\tau \rho + \nabla \cdot (\rho \mathbf{v}) = 0, \rho(0, \mathbf{x}) = \rho^n(\mathbf{x}). \end{cases} \quad (22)$$

This reformulation states that the minimization is constrained by the classical continuity equation. It is worth noting that such reformulation does not change the optimizer of the problem. Indeed, let ϕ be the Lagrangian multiplier, then one can check that the critical point system of variational problems (21) and (22) lead to the same Hamilton-Jacobi equation for ϕ :

$$\partial_\tau \phi + \frac{1}{4} M'(\rho) |\nabla \phi|^2 = 0.$$

We note that variational problem (22) now admits a similar Lagrangian representation as in (8), with the only distinction being the first term, which now becomes:

$$\int_0^1 \int_{\mathbb{R}^d} \frac{\rho^2(\tau, \mathbf{x})}{M(\rho(\tau, \mathbf{x}))} |\mathbf{v}(\tau, \mathbf{x})|^2 d\tau d\mathbf{x}.$$

By employing the same technique as utilized in deriving (11), the integral with respect to $\mathbf{x}$ in the given expression rewrites as

$$\begin{aligned} & \int_{\mathbb{R}^d} \frac{\rho^2(\tau, \mathbf{x})}{M(\rho(\tau, \mathbf{x}))} |\mathbf{v}(\tau, \mathbf{x})|^2 d\mathbf{x} \\ &= \int_{\mathbb{R}^d} \frac{\rho^2(\tau, \mathbf{T}(\tau, \mathbf{x}))}{M(\rho(\tau, \mathbf{T}(\tau, \mathbf{x})))} |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2 |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})| d\mathbf{x} \\ &= \int_{\mathbb{R}^d} \frac{\rho^2(\tau, \mathbf{T}(\tau, \mathbf{x}))}{M(\rho(\tau, \mathbf{T}(\tau, \mathbf{x})))} |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2 \frac{|\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})|}{\rho^n(\mathbf{x})} \rho^n(\mathbf{x}) d\mathbf{x} \\ &= \int_{\mathbb{R}^d} \frac{\rho^2(\tau, \mathbf{T}(\tau, \mathbf{x}))}{M(\rho(\tau, \mathbf{T}(\tau, \mathbf{x})))} \frac{|\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2}{\rho(\tau, \mathbf{T}(\tau, \mathbf{x}))} \rho^n(\mathbf{x}) d\mathbf{x} \\ &= \int_{\mathbb{R}^d} \frac{\rho(\tau, \mathbf{T}(\tau, \mathbf{x}))}{M(\rho(\tau, \mathbf{T}(\tau, \mathbf{x})))} |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}))|^2 \rho^n(\mathbf{x}) d\mathbf{x}. \end{aligned} \quad (23)$$

As in (18), the density $\rho(\tau, \mathbf{T}(\tau, \mathbf{x}))$ is determined by the Monge-Amperè equation:

$$\rho(\tau, \mathbf{T}(\tau, \mathbf{x})) = \frac{\rho^n(\mathbf{x})}{|\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x})|}.$$

Similar to the JKO scheme presented in (14) or (15), given $\{\mathbf{x}_j^n\}_j$ and $\{\rho^n(\mathbf{x}_j^n)\}_j$ (we omit the superscript n in $\mathbf{x}_j^n$ in the following formula), the updated formulation can now be expressed as follows:

$$\left\{ \begin{array}{l} \min_{\mathbf{T}} \frac{1}{N} \sum_{j=1}^N \left[\int_0^1 \frac{\rho(\tau, \mathbf{T}(\tau, \mathbf{x}_j))}{M(\rho(\tau, \mathbf{T}(\tau, \mathbf{x}_j)))} |\mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j))|^2 d\tau + 2\Delta t \frac{U(\rho(1, \mathbf{T}(1, \mathbf{x}_j)))}{\rho(1, \mathbf{T}(1, \mathbf{x}_j))} \right] \\ \text{s.t.} \quad \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}_j) = \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \mathbf{T}(0, \mathbf{x}_j) = \mathbf{x}_j \\ \frac{d}{d\tau} \log |\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}_j)| = \operatorname{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \log |\det \nabla_{\mathbf{x}} \mathbf{T}(0, \mathbf{x}_j)| = 0 \\ \rho(\tau, \mathbf{T}(\tau, \mathbf{x}_j)) = \frac{\rho^n(\mathbf{x}_j)}{|\det \nabla_{\mathbf{x}} \mathbf{T}(\tau, \mathbf{x}_j)|}. \end{array} \right. \quad (24)$$

It is obvious that when $M(\rho) = \rho$, (24) reduces to (15).

3. Neural ODE empowered JKO schemes

In this section, we leverage neural network functions to approximate the vector field $\mathbf{v}$ as described in equations (14), (15), or, more broadly, equation (24). As a result, the optimization procedure aims at refining the parameters of the neural network. Building upon the efficiency gains achieved through the use of continuous normalizing flow for density estimation, as previously discussed, we incorporate a recursive relation formula proposed in a prior work [34] to compute derivatives, significantly improving the efficiency of backpropagation. Lastly, we conclude this section by presenting a concise summary of the algorithm.

3.1. Learning the potential function

In view of (14) and (15), the primary objective is to determine the optimal transport map, denoted as $\mathbf{T}$. However, this task can be computationally demanding, especially in high-dimensional scenarios. Fortunately, by applying the principles of optimal transport [20, 40], we can express the velocity field $\mathbf{v}(t, \mathbf{x})$ as the gradient of a potential function $\phi(t, \mathbf{x})$ such that $\mathbf{v}(t, \mathbf{x}) = -\nabla_{\mathbf{x}} \phi(t, \mathbf{x})$. Consequently, our focus shifts to finding the potential function ϕ . To achieve this, we utilize a neural network to approximate $\phi(t, \mathbf{x})$, denoted as $\phi_{\theta}(t, \mathbf{x})$. As a result, we obtain $\mathbf{v}_{\theta}(t, \mathbf{x}) = -\nabla_{\mathbf{x}} \phi_{\theta}(t, \mathbf{x})$ as the corresponding approximation of the velocity field.

To be more specific, let us consider solving equation (14). By substituting $\mathbf{v}$ with $\mathbf{v}_{\theta} = -\nabla_{\mathbf{x}} \phi_{\theta}$, the constraint in (14) becomes:

$$\frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{x}_j) = -\nabla_{\mathbf{x}} \phi_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{x}_j)), \quad \mathbf{T}(0, \mathbf{x}_j) = \mathbf{x}_j.$$

Similarly, in the case of the more complex equation (15), computing $\rho(\mathbf{T}(1, \mathbf{x}_j))$ calls for the following update:

$$\frac{d}{d\tau} \log |\det(\nabla_{\mathbf{x}}^2 \phi_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{x}_j)))| = \text{div}(\mathbf{v}_{\theta})(\tau, \mathbf{T}(\tau, \mathbf{x}_j)) = -\text{tr}(\nabla_{\mathbf{x}}^2 \phi_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{x}_j))).$$

Detailed computations of $\nabla^2 \phi_{\theta}$ will be given in the next subsection.

3.2. Deep JKO Algorithms

We hereby provide comprehensive details for computing the minimizer of (14), (15) or (24), when $\mathbf{v}$ is replaced by $\mathbf{v}_{\theta}$. Let N_{τ} denote the number of time discretization such that $[0, 1]$ is discretized into $0, \frac{1}{N_{\tau}}, \frac{2}{N_{\tau}}, \dots, \frac{N_{\tau}-1}{N_{\tau}},$ and θ is a vector of parameters for neural network functions. Given the density ρ^n at the n -th JKO step and sampled points $\{\mathbf{x}_j\}_{j=1}^N$, the loss function to compute the density at the $n+1$ -th JKO step is defined as follows:

$$\begin{aligned} L(\theta, \{\mathbf{x}_j\}_{j=1}^N) = \frac{1}{N} \sum_{j=1}^N \left[d\tau \sum_{k=1}^{N_{\tau}} |\nabla_{\mathbf{x}} \phi_{\theta}(kd\tau, \mathbf{z}_j^k)|^2 \right. \\ \left. + 2\Delta t \left(\frac{U(\rho(\mathbf{z}_j^{N_{\tau}}))}{\rho(\mathbf{z}_j^{N_{\tau}})} + V(\mathbf{z}_j^{N_{\tau}}) + \frac{1}{N} \sum_{l=1}^N W(\mathbf{z}_j^{N_{\tau}}, \mathbf{z}_l^{N_{\tau}}) \right) \right], \end{aligned} \quad (25)$$

where the inner time step is denoted as $d\tau = 1/N_{\tau}$ and the update equations are given by:

$$\begin{aligned} \mathbf{z}_j^{k+1} &= \mathbf{z}_j^k - d\tau \nabla_{\mathbf{x}} \phi_{\theta}(kd\tau, \mathbf{z}_j^k), \quad \mathbf{z}_j^0 = \mathbf{x}_j, \\ l_j^{k+1} &= l_j^k - d\tau \text{tr}(\nabla_{\mathbf{x}}^2 \phi_{\theta}(kd\tau, \mathbf{z}_j^k)), \quad l_j^0 = 0, \end{aligned}$$

for $k = 0, \dots, N_{\tau} - 1$ and the terminal density ρ at the location $\mathbf{z}_j^{N_{\tau}}$ is defined as

$$\rho(1, \mathbf{z}_j^{N_{\tau}}) = \frac{\rho^n(\mathbf{x}_j)}{\exp(l_j^{N_{\tau}})}. \quad (26)$$

The above equations for $\mathbf{z}_j^{k+1}$ and $\log \rho(\mathbf{z}_j^{k+1})$ utilize the forward Euler method to solve the ordinary differential equations (ODEs) in the constraints of (15). However, alternative ODE solvers such as Runge-Kutta can also be employed. For our numerical experiment, we utilize the RK4 as the ODE solver.

The loss function, denoted as $L(\theta, \{\mathbf{x}_j\}_{j=1}^N)$, takes two inputs: the neural network parameter θ and a set of sampled points $\{\mathbf{x}_j\}_{j=1}^N$. The objective is to minimize this loss function with respect to θ . This θ plays a vital role in approximating the potential function $\phi_{\theta}(\mathbf{x}, \tau)$, which relies on two variables: a state vector $\mathbf{x} \in \mathbb{R}^d$ and a time value $t \in [0, 1]$. The velocity is computed as:

$$\mathbf{v}_{\theta}(\tau, \mathbf{x}) = -\nabla_{\mathbf{x}} \phi_{\theta}(\tau, \mathbf{x}),$$

and its divergence is obtained as the trace of the Hessian:

$$\text{div}(\mathbf{v}_{\theta})(\tau, \mathbf{x}) = -\text{tr}(\nabla_{\mathbf{x}}^2 \phi_{\theta}(\tau, \mathbf{x})).$$

To simplify the calculation of gradients and Hessians in the formulas above, we utilize the explicit expression introduced in [34]. This approach capitalizes on the smoothness of the activation function within the neural network. Consequently, our optimization process experiences significant acceleration, as we no longer rely on automatic differentiation (AD) for backpropagation to compute gradients and Hessians in each iteration.

More precisely, we adopt a multi-layer ResNet that takes the input $\mathbf{s} = (\tau, \mathbf{x}) \in \mathbb{R}^{d+1}$, where $\mathbf{x} \in \mathbb{R}^d$ represents the initial state, and $\tau \in [0, 1]$ denotes time, yielding an output in $\mathbb{R}$. The neural network function consists of L hidden layers with recursive relations defined as follows:

$$\begin{aligned} \mathbf{u}_1 &= \sigma(W_0 \mathbf{s} + \mathbf{b}_0), \\ \mathbf{u}_{l+1} &= \mathbf{u}_l + \sigma(W_l \mathbf{u}_l + \mathbf{b}_l), \quad l = 1, \dots, L-1, \\ NN(\mathbf{s}; \theta) &= \mathbf{w}^\top \mathbf{u}_L. \end{aligned}$$

Here, $\theta = (W_0, W_l, b_l)$, with $W_0 \in \mathbb{R}^{m \times (d+1)}$ and $W_l \in \mathbb{R}^{m \times m}$ ($l = 1, \dots, L$) are dense matrices. Additionally, $b_k \in \mathbb{R}^m$ ($l = 0, \dots, L$) are biases, while m denotes the number of nodes in each hidden layer and $\mathbf{w} \in \mathbb{R}^m$ is a vector we perform dot product at the end of the ResNet. The element-wise activation function $\sigma(\mathbf{x}) = \log(\exp(\mathbf{x}) + \exp(-\mathbf{x}))$ is used, which is differentiable. The dual variable ϕ_θ is defined through

$$\phi_\theta(\tau, \mathbf{x}) = NN((\tau, \mathbf{x}); \theta). \quad (27)$$

Denote by $\nabla_{\mathbf{x}}$ a gradient operator with respect to the space variable $\mathbf{x}$ and $\nabla_{\mathbf{s}}$ be a gradient operator with respect to the state vector $\mathbf{s}$. Both the gradient $\nabla_{\mathbf{x}} \phi_\theta(\mathbf{x}, \tau)$ and the Hessian $\nabla_{\mathbf{x}}^2 \phi_\theta(\mathbf{x}, \tau)$ can be effortlessly computed using the explicit formulas presented in [34]. Specifically, for the case where $L = 2$, these formulas are as follows:

$$\nabla_{\mathbf{s}} \phi_\theta(\tau, \mathbf{x}) = W_0^\top \text{diag}(\sigma'(W_0 \mathbf{s} + b_0)) \mathbf{z}_1, \quad (28a)$$

$$\mathbf{z}_1 = \mathbf{w} + W_1^\top \text{diag}(\sigma'(W_1 \mathbf{u}_1 + b_1)) \mathbf{w}, \quad (28b)$$

and

$$\nabla_{\mathbf{s}}^2 \phi_\theta(\mathbf{x}, \tau) = t_0 + t_1, \quad (28c)$$

where

$$\begin{aligned} t_0 &= W_0^\top \text{diag}(\sigma''(W_0 \mathbf{s} + b_0) \odot \mathbf{z}_1) W_0, \\ t_1 &= J_0^\top W_1^\top \text{diag}(\sigma''(W_1 \mathbf{u}_1 + b_1) \odot \mathbf{w}) W_1 J_0, \\ J_0 &= W_0^\top \text{diag}(\sigma'(W_0 \mathbf{s} + b_0))^\top. \end{aligned}$$

Here $\odot$ represents the element-wise multiplication. Then, $\nabla_{\mathbf{x}} \phi_\theta(\mathbf{x}, \tau)$ is the first d elements of $\nabla_{\mathbf{s}} \phi_\theta(\tau, \mathbf{x})$, and $\nabla_{\mathbf{x}}^2 \phi_\theta(\tau, \mathbf{x}) = [\nabla_{\mathbf{s}}^2 \phi_\theta(\tau, \mathbf{x})]_{1:d, 1:d}$ where the subscript $1 : d, 1 : d$ denotes the submatrix of size $d \times d$. For comprehensive expressions of the gradient and Hessian operators for a general L , we recommend readers to see details in [34].

We now provide a summary of the proposed algorithms. Algorithm 2 serves as the primary algorithm, guiding the movement of particles along the gradient flow. Whenever sampling/resampling is necessary, Algorithm 1 will be invoked. Please note that the l update and density computation in Algorithm 1 are only required when working with general nonlinear internal energy and mobility. Additionally, in the algorithm, we utilize forward Euler, but it is adaptable for replacement with explicit Runge-Kutta type ODE solvers.

Algorithm 1 Generate samples and densities at n -th JKO step

Input:

- Outer JKO time step Δt .
- Inner dynamic formulation time step $d\tau$ and total number of inner steps N_τ .
- Initial distribution ρ^0 and initial samples $\{\mathbf{x}_j^0\}_{j=1}^N$ and densities $\{\rho(\mathbf{x}_j^0)\}_{j=1}^N$.
- Potential functions from previous steps $(\phi_\theta^0, \dots, \phi_\theta^{n-1})$.

Output:

- Sampled points $\{\mathbf{x}_j^n\}_{j=1}^N$ at n -th outer iteration.
- Density values $\{\rho(\mathbf{x}_j^n)\}_{j=1}^N$ at n -th outer iteration.

```

for  $i = 0, \dots, n-1$  do
  for  $j = 1, \dots, N$  do
     $\mathbf{z}_j^0 = \mathbf{x}_j^i, l_j^0 = 0$ 
    for  $k = 0, \dots, N_\tau - 1$  do
       $\mathbf{z}_j^{k+1} = \mathbf{z}_j^k - d\tau \nabla_{\mathbf{x}} \phi_\theta^i(kd\tau, \mathbf{z}_j^k)$ 
       $l_j^{k+1} = l_j^k - d\tau \text{tr}(\nabla_{\mathbf{x}}^2 \phi_\theta^i(kd\tau, \mathbf{z}_j^k))$ 
    end for
     $\mathbf{x}_j^{i+1} = \mathbf{z}_j^{N_\tau}, \rho(\mathbf{x}_j^{i+1}) = \frac{\rho^n(\mathbf{x}_j)}{\exp(l_j^{N_\tau})}$ 
  end for
end for

```

Algorithm 2 Deep JKO scheme

Input:

- Outer JKO time step Δt
- Inner dynamic formulation time step $d\tau$
- Initial distribution ρ^0
- Learning rate α for the Adam optimizer
- Total number of outer iterations K
- Resampling frequency C

Output:

- Neural network parameter θ and corresponding potential function $\phi_\theta^n(t, \mathbf{x})$, $n = 1, 2, \dots, K$
- Particle locations: $\{\mathbf{x}_j^n\}$, $n = 1, 2, \dots, K$, $j = 1, 2, \dots, N$.
- Density along particle trajectories: $\rho^n(\mathbf{x}_j^n)$, $n = 1, 2, \dots, K$, $j = 1, 2, \dots, N$.

Initialize θ from a standard normal distribution.

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

$c = 0$.

while not converged **do**

if $\text{mod}(c, C) = 0$ **then do**

 Sample $\{\mathbf{x}_j^0\}_{j=1}^N \stackrel{\text{i.i.d.}}{\sim} \rho^0$

 Compute $\{\mathbf{x}_j^n\}_{j=1}^N$ using Algorithm 1 and update the density value through (26).

end if

 Update $\theta \leftarrow \theta - \alpha \nabla_\theta L(\theta, \{\mathbf{x}_j^n\}_{j=1}^N)$

 where L is from (25) and the gradient is computed using automatic differentiation (AD).

$c \leftarrow c + 1$

end while

end for

4. Numerical examples

This section presents several numerical examples of the proposed deep JKO schemes for computing gradient flows with various choices of energy functionals and mobility functions.

4.1. Wasserstein gradient flow of the KL divergence

In this experiment, we present the computed solutions of (15) with the energy functional defined as

$$U(\rho) = D_{\text{KL}}(\rho||q) = \int \rho(\mathbf{x}) \log \frac{\rho(\mathbf{x})}{q(\mathbf{x})} d\mathbf{x}, \quad (29)$$

where q represents a reference density. The initial density ρ^0 is characterized by two Gaussian distributions centered at $(\pm 1.2, 0)$ with a standard deviation of $\sigma = 0.5$. The reference density q is defined as four Gaussian distributions centered at $(\pm 2, \pm 2)$ each with a covariance matrix represented by a diagonal matrix with every diagonal entry equal to 0.25.

For the numerical experiment, the inner timestep is set as $N_\tau = 1$ and the JKO time stepsize is set as $\Delta t = 0.025$. The model architecture includes 3 layers with $m = 64$ nodes in the hidden layers. The training is performed using a batch size of 1000, and the learning rate is chosen to be 10^{-5} . The total number of iterations is set to 10,000.

The two-dimensional results are depicted in Figure 1 and Figure 2. Figure 1 illustrates the evolution of the density from $t = 0$ to $t = 0.4$. In this figure, the color of the points represents the value of the density $\rho^n(\mathbf{x})$ at the corresponding location $\mathbf{x}$. Brighter colors indicate higher density values. The computation involves a total of 16 JKO steps. Figure 2 displays the trajectories of each particle with respect to time t and the energy value computed from the algorithm with respect to the iterations. As the iterations progress, the energy decreases and eventually converges to the stationary value.

We further extended our experiments to a 10-dimensional space while retaining the same energy functional (29). The initial density ρ^0 was set as a Gaussian distribution centered at the origin, and the reference density q was constructed using a combination of four Gaussian distributions centered at $(2, 2.5)$, $(-2.5, 2)$, $(-2, -2.5)$, and $(2.5, -2)$ each with a covariance matrix represented by a diagonal matrix with every diagonal entry equal to 0.25. In this experimental setup, we use a time step size of $\Delta \tau = 0.2$, a learning rate of 10^{-5} , and executed a total of 10,000 iterations. The results are illustrated in Figure 3, highlighting the initial 2D cross sections of the density's evolution from $t = 0$ to $t = 3.4$. These visual representations present kernel density plots derived from the point clouds generated by the algorithm at each time instance t .

4.2. Porous medium equation

The second experiment concerns the porous medium equation represented by

$$\partial_t \rho(t, \mathbf{x}) = \Delta \rho(t, \mathbf{x})^m, \quad m > 1, \mathbf{x} \in \mathbb{R}^d, \quad (30)$$

which can be interpreted as the Wasserstein gradient flow with the internal energy functional U defined as

$$U(\rho) = \int_{\mathbb{R}^d} \frac{1}{m-1} \rho(\mathbf{x})^m d\mathbf{x}.$$

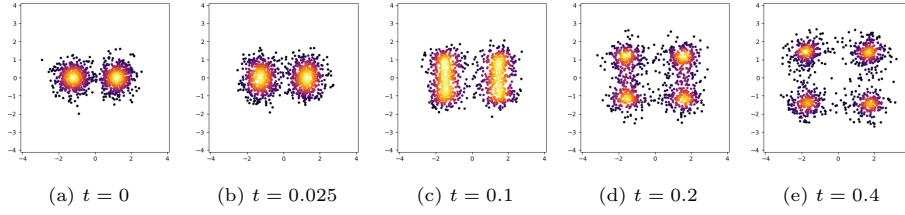


Figure 1: Computed solutions of 2D Fokker-Planck equation, i.e. the Wasserstein gradient flow of KL divergence (29). An initial density is a mixture of two Gaussians and the target density is a mixture of four Gaussians. The figures visually depict the evolution of these densities from $t = 0$ to $t = 0.4$.

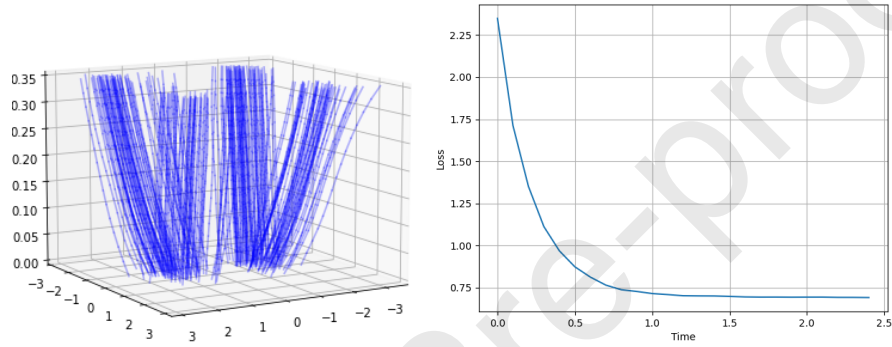


Figure 2: 3D plot (left) and the energy decay plot (right) from Figure 1. The left plot shows the trajectories of each particles where z -axis represents time from $t = 0$ to $t = 0.35$. The right plot illustrates the decay of the energy from $t = 0$ to $t = 2.4$.

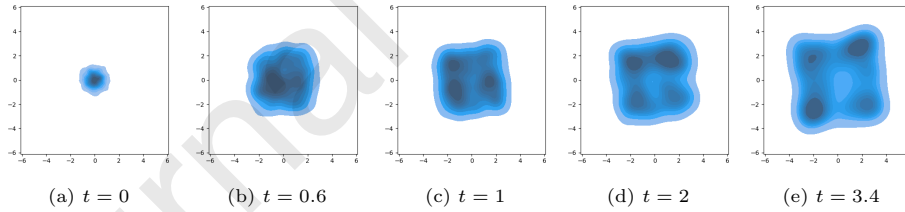


Figure 3: Evolution of densities from 10D Fokker-Planck equation from $t = 0$ to $t = 3.4$. An initial density is a Gaussian distribution centered at the origin and the target density is a mixture of four Gaussians.

This equation admits a close-form solution given by the Barenblatt formula:

$$\rho(t, \mathbf{x}) = (t + t_0)^{-\alpha} \left(C - \beta \|\mathbf{x}\|^2 (t + t_0)^{-\frac{2\alpha}{d}} \right)_+^{\frac{1}{m-1}}, \quad (31)$$

where $\alpha = \frac{d}{d(m-1)+2}$, $\beta = \frac{(m-1)\alpha}{2dm}$, C is a constant that makes $\int_{\mathbb{R}^d} \rho(t, \mathbf{x}) d\mathbf{x} = 1$. We choose $t_0 = 10^{-3}$. Formula (31) serves as a reference solution to evaluate the performance of our algorithm.

In the numerical experiment, we employ an inner timestep of $N_\tau = 2$ and set the JKO time step size to $\Delta t = 0.001$. The model architecture consists of 3 layers with 128 nodes in each hidden layer. During training, a batch size of 1000 is used, and the learning rate is set at 10^{-5} . The total number of iterations for the experiment is configured to be 10,000.

The experimental results are presented in Figure 4 and Figure 5. Figure 4 visualizes the density evolution from $t = 0$ to $t = 0.12$, where particles positions $\mathbf{x}$ and their corresponding density values, $\rho(t, \mathbf{x})$, are accurately computed at each time step. The color intensity at each particle's location represents its density, with brighter colors indicating higher density values. Notably, the computed solutions maintain radial symmetry, providing a precise representation. The green rings in the figures correspond to the support of the analytical solution (31) and align well with the computed solutions, affirming the algorithm's accuracy. Figure 5 displays particle trajectories from $t = 0$ to $t = 0.01$, exhibiting linear trajectories as expected from the analytical solution of the PDE.

We then utilize the algorithm to compute solutions for the porous medium equation in high-dimensional spaces, specifically in dimensions 6 and 10. For both cases, we choose an outer time step size of $\Delta t = 0.0005$, and set the inner timestep to $N_\tau = 2$. The model architecture comprises three layers with 128 nodes in the hidden layers. Training is carried out using a batch size of 1000, and the learning rate is fixed at 10^{-6} . The total number of iterations is set to be 20,000.

Figure 6 illustrates the evolution of densities in both 6D and 10D. To visualize the results and demonstrate the accuracy of the algorithm, 2D cross-section plots of the particles are presented, along with the support of the densities obtained from the analytical solution, represented as a green ring. The computed densities exhibit spherical evolution and closely match the analytical solutions.

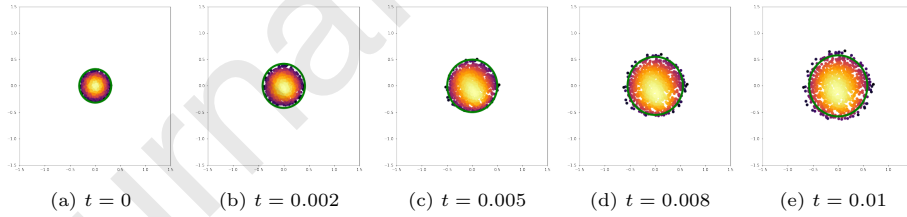


Figure 4: Computed solutions of 2D porous medium equation. The figures show the evolution of densities from $t = 0$ to $t = 0.12$. The algorithm calculates the positions of particles $\mathbf{x} \in \mathbb{R}^2$ and their corresponding density values $\rho(t, \mathbf{x})$ at each time t . The color of each particle represents its density value, with brighter colors indicating higher density values. The green ring indicates the support of the density from the analytical solution at the given time t .

Here, we discuss the complexity and accuracy of our scheme. For 2D and 6D cases, we use 4-layer fully connected neural network with 512 nodes per layer, and for 10D case, we use 4-layer fully connected neural network with 1024 nodes per layer. For all cases, we set the number of iterations to be 50,000 during the first outer iteration and 5,000 iterations thereafter. The cost per iteration is

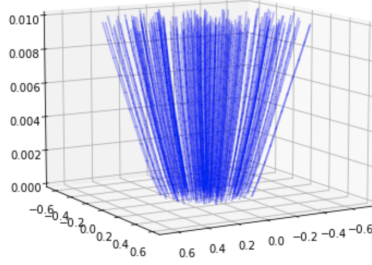


Figure 5: 3D plot from Figure 4. It shows the trajectories of each particles where z-axis represents time from $t = 0$ to $t = 0.01$.

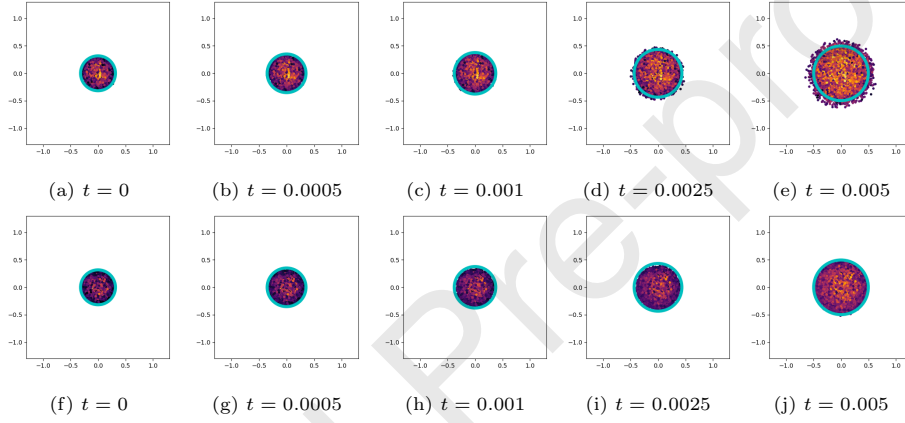


Figure 6: Computed solutions of 6D (first row) and 10D (second row) porous medium equation. The figures show the evolution of densities from $t = 0$ to $t = 0.005$ for 6D and 10D porous medium flows. The green ring indicates the support of the density from the analytical solution at the given time t .

however the same across different dimensions, approximately 1 to 2 minutes per 1000 iterations on 2 GPUs. We evaluate the accuracy by comparing the computed density using the density update formula with the exact solution. Figure 7 illustrates this comparison for a dataset $\{\mathbf{x}_j\}_{j=1}^N$ with $N = 50,000$. The L^2 error is calculated as:

$$L^2 \text{ error} = \left(\frac{1}{N} \sum_{j=1}^N |\rho_{\text{computed}}(\mathbf{x}_j^k) - \rho_{\text{exact}}(\mathbf{x}_j^k)|^2 \right)^{1/2}.$$

4.3. Nonlocal transport equation with nonlinear mobility

In this section, we explore an example that involves nonlinear mobility, a phenomenon frequently encountered in the study of transport in biological systems, where it serves to prevent overcrowding [5]. Particularly, we examine a

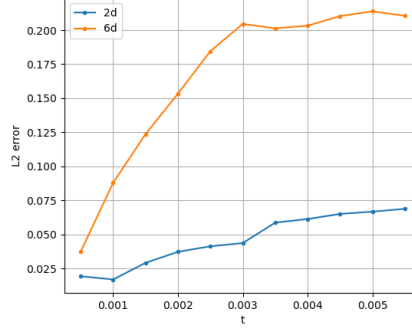


Figure 7: L^2 error from porous medium equation experiments for 2D and 6D results

specific mobility function, represented as $M(\rho) = \rho(1 - \rho)$, and a particular energy functional described by the equation:

$$\mathcal{E}(\rho) = \frac{1}{2} \int_{\mathbb{R}^d \times \mathbb{R}^d} W(\mathbf{x} - \mathbf{y}) \rho(\mathbf{x}) \rho(\mathbf{y}) d\mathbf{x} d\mathbf{y}, \quad W(\mathbf{x}) = |\mathbf{x}|^2. \quad (32)$$

The corresponding equation can be expressed as:

$$\partial_t \rho(t, \mathbf{x}) = \nabla \cdot \left(M(\rho(t, \mathbf{x})) \nabla \frac{\delta \mathcal{E}}{\delta \rho}(t, \mathbf{x}) \right). \quad (33)$$

This model is based on [17]. Beginning with an initial condition $\rho_0(\mathbf{x}) = \frac{3}{4}(1 - |\mathbf{x}|^2)_+$, it is anticipated that ρ_∞ will take the form of a characteristic function whose support is determined by the total mass. The experimental results are presented in Figure 8. The top row of figures illustrates particle evolution from a bird's-eye view, where each particle's position $\mathbf{x}$ is color-coded according to its associated density value, $\rho^n(\mathbf{x})$. The bottom row of figures presents a side view of the density values for each particle when their locations are projected onto the x -axis. These visualizations provide compelling evidence of the density converging toward a characteristic function.

4.4. Kalman-Wasserstein gradient flow

In the last example, we consider the Kalman-Wasserstein gradient flow [21]:

$$\partial_t \rho(t, \mathbf{u}) = \nabla_{\mathbf{u}} \cdot \left(\rho(t, \mathbf{u}) C(\rho) \nabla_{\mathbf{u}} \frac{\delta \mathcal{E}}{\delta \rho}(t, \mathbf{u}) \right), \quad (34)$$

where $\mathbf{u} \in \mathbb{R}^d$ is the parameter in the Bayesian inverse problem,

$$C(\rho) = \int_{\mathbb{R}^d} (\mathbf{u} - \mathbf{m}(\rho)) \otimes (\mathbf{u} - \mathbf{m}(\rho)) \rho(\mathbf{u}) d\mathbf{u}, \quad \mathbf{m}(\rho) = \int_{\mathbb{R}^d} \mathbf{u} \rho(\mathbf{u}) d\mathbf{u}, \quad (35)$$

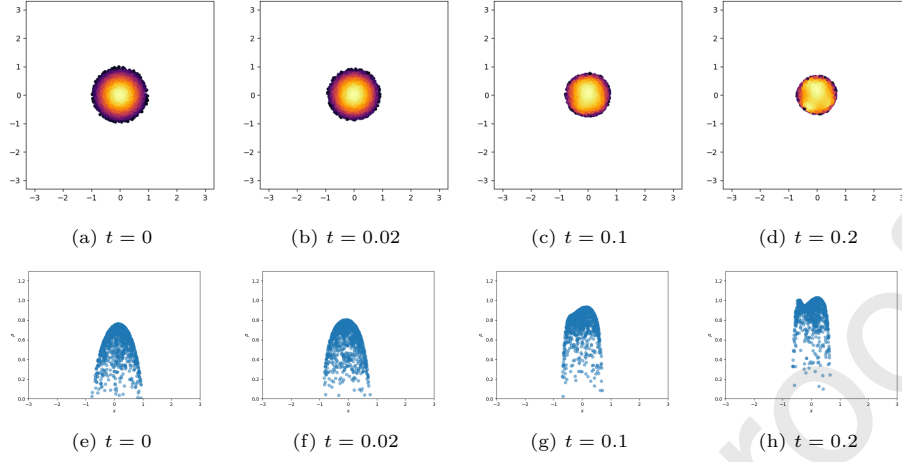


Figure 8: Computed solutions of nonlocal transport equation with nonlinear mobility (33). The figures show the evolution of densities from $t = 0$ to $t = 0.2$. Top row: bird eye view with color coded density. Bottom row: particle density (on the y -axis) versus its projected location on the x -axis.

and

$$\mathcal{E}(\rho) = D_{\text{KL}}(\rho \parallel \pi) = \int_{\mathbb{R}^d} \left(\Phi(\mathbf{u})\rho + \rho(\mathbf{u}) \ln \rho(\mathbf{u}) \right) d\mathbf{u} + \ln Z. \quad (36)$$

Here ρ is a probability density function of $\mathbf{u}$, $\pi = \frac{1}{Z}e^{-\Phi(\mathbf{u})}$ is a given target distribution and we assume $Z = \int e^{-\Phi(\mathbf{u})} d\mathbf{u} < +\infty$ is a normalization constant.

This equation is derived as a mean field limit of the Kalman Wasserstein sampler, aiming at identifying the unknown parameter $\mathbf{u}$ in the Bayesian setting. In particular, $\Phi(\mathbf{u})$ takes the form

$$\Phi(\mathbf{u}) = \frac{1}{2} \|\mathcal{G}(\mathbf{u}) - \mathbf{y}\|_{\Gamma}^2 + \frac{1}{2} \|\mathbf{u}\|_{\Gamma_0}^2, \quad (37)$$

where $\|\mathbf{u}\|_{\mathbf{A}} := \mathbf{u}^{\top} \mathbf{A}^{-1} \mathbf{u}$, $\mathcal{G}$ represents a forward map that relates the parameter $\mathbf{u}$ to the measurement $\mathbf{y}$, typically represented as a PDE operator. We assume that $\mathbf{u}$ follows a prior distribution $\mathcal{N}(0, \Gamma_0)$. The term $\frac{1}{2} \|\mathcal{G}(\mathbf{u}) - \mathbf{y}\|_{\Gamma}^2$ is commonly referred to as the likelihood function, and $\frac{1}{Z}e^{-\Phi(\mathbf{u})}$, with Z being the normalizing constant, represents the posterior distribution.

Then our goal of solving (34) is to sample from the posterior distribution. Instead of running a Lagenvin type dynamics as proposed in [21], we use our DeepJKO formulation. Note that $C(\rho)$ does not have an explicit $\mathbf{u}$ dependence,

(23) simplifies to

$$\begin{aligned}
& \int_{\mathbb{R}^d} \rho(\tau, \mathbf{u}) \mathbf{v}(\tau, \mathbf{u})^\top C(\rho(\tau, \mathbf{u}))^{-1} \mathbf{v}(\tau, \mathbf{u}) d\mathbf{u} \\
&= \int_{\mathbb{R}^d} \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{u}))^\top C(\rho(\tau, \mathbf{T}(\tau, \mathbf{u})))^{-1} \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{u})) \rho^n(\mathbf{u}) d\mathbf{u} \\
&= \int_{\mathbb{R}^d} \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{u}))^\top C(\tau)^{-1} \mathbf{v}(\tau, \mathbf{T}(\tau, \mathbf{u})) \rho^n(\mathbf{u}) d\mathbf{u}.
\end{aligned}$$

It's important to highlight that $C(\tau) = C(\rho(\tau, \mathbf{T}(\tau, \mathbf{u})))$ is reliant on the moments of ρ and thus remains unaffected by variations in u . Consequently, the JKO scheme (24) rewrites to:

$$\left\{ \begin{array}{l} \min_{\theta} \frac{1}{N} \sum_{j=1}^N \left\{ \int_0^1 \mathbf{v}_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{u}_j))^\top C(\tau)^{-1} \mathbf{v}_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{u}_j)) d\tau \right. \\ \quad \left. + 2\Delta t [\Phi(\mathbf{T}(1, \mathbf{u}_j)) + \log \rho^n(\mathbf{u}_j) - \log \det(\nabla_{\mathbf{u}} \mathbf{T}(1, \mathbf{u}_j))] \right\}, \\ \text{where } C(\tau) = \frac{1}{N} \sum_{j=1}^N (\mathbf{T}(\tau, \mathbf{u}_j) - \mathbf{m}) \otimes (\mathbf{T}(\tau, \mathbf{u}_j) - \mathbf{m}), \quad \mathbf{m} = \frac{1}{N} \sum_{j=1}^N \mathbf{T}(\tau, \mathbf{u}_j), \\ \text{s.t. } \frac{d}{d\tau} \mathbf{T}(\tau, \mathbf{u}_j) = \mathbf{v}_{\theta}(\tau, \mathbf{T}(\tau, \mathbf{u}_j)), \quad \mathbf{T}(0, \mathbf{u}_j) = \mathbf{u}_j. \\ \quad \frac{\partial}{\partial \tau} \log \det(\nabla_{\mathbf{u}} \mathbf{T}(\tau, \mathbf{u}_j)) = \operatorname{div}(\mathbf{v})(\tau, \mathbf{T}(\tau, \mathbf{u}_j)), \quad \log \det(\nabla_{\mathbf{u}} \mathbf{T}(0, \mathbf{u}_j)) = 0. \end{array} \right. \quad (38)$$

As a specific example, we consider a case where the forward map $\mathcal{G}$ is given by the one-dimensional elliptic boundary value problem:

$$-\frac{d}{dx} \left(e^{u_1} \frac{d}{dx} p_{\mathbf{u}}(x) \right) = 1, \quad x \in [0, 1],$$

with boundary conditions $p_{\mathbf{u}}(0) = 0$ and $p_{\mathbf{u}}(1) = u_2$. This problem has been considered in [24, 21, 28]. The explicit solution for this problem is

$$p_{\mathbf{u}}(x) = u_2 x + e^{-u_1} \left(-\frac{x^2}{2} + \frac{x}{2} \right).$$

Then the forward map $\mathcal{G}$ is defined by

$$\mathcal{G}(\mathbf{u}) = (p_{\mathbf{u}}(0.25), p_{\mathbf{u}}(0.75))^\top, \quad \text{where } \mathbf{u} = (u_1, u_2)^\top. \quad (39)$$

Then the Bayesian inverse problem is to find the distribution of the unknown $\mathbf{u}$ conditioned on the observation $\mathbf{y}$, assuming additive Gaussian noise. More precisely, we use $\Phi(\mathbf{u})$ defined in (37) and choose

$$\mathbf{y} = (27.5, 79.7)^\top, \quad \Gamma = 0.1^2 \mathbf{I}_2, \quad \Gamma_0 = 10^2 \mathbf{I}_2,$$

where $\mathbf{I}_2 \in \mathbb{R}^{2 \times 2}$ is the identity matrix. We run the dynamics by initially sampling according to

$$\rho_0(\mathbf{u}) = \mathcal{N}(0, 1) \times \mathcal{U}(90, 110), \quad (40)$$

That is, sample u_1 according to normal distribution $\mathcal{N}(0, 1)$, and u_2 according to uniform distribution $\mathcal{U}(90, 110)$.

In the computational approach, we set the matrix $C(\tau)^{-1}$ to be equal to $C(0)^{-1}$ to approximate the running costs. This approximation offers the advantage of reducing computational complexity by circumventing the need for backpropagation over the inverse cost function, thereby significantly increasing computation speed. Note that the approximation introduces an error of order Δt , which does not compromise the first order accuracy of the JKO scheme. The results of our experiment are presented in Figure 9, illustrating the density evolution from $t = 0$ to $t = 20$. The density converges to the stationary distribution, and the results are consistent with the numerical solution depicted in [21].

Figure 10 shows the loss plot and the contour plot of $\Phi(\mathbf{u})$. The loss plot displays the value of the loss function with the x-axis representing iterations and the y-axis representing the loss value. Here the loss function represents the relative entropy defined in (36). On the right, the figure illustrates the computed solution at $t = 20$, overlaid on the contour plot derived from $\Phi(\mathbf{u})$.

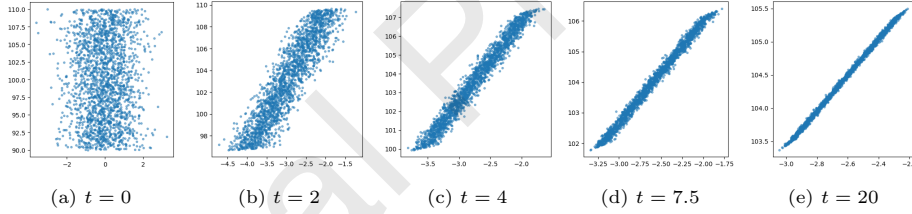


Figure 9: Computed solutions of a 2D Kalman-Wasserstein gradient flow (34).

5. Conclusion and discussion

In this paper, we introduce neural network-based implicit particle methods for computing high-dimensional Wasserstein-type gradient flows with linear and nonlinear mobility functions. Our approach centers around the Lagrangian formulation within the Jordan–Kinderlehrer–Otto (JKO) framework, where neural networks approximate the velocity field. Neural ODE techniques are harnessed to efficiently compute the time-implicit updates of both particle locations and the densities they carry. Additionally, we leverage an explicit recurrence relation to compute derivatives, significantly streamlining the backpropagation process. Our methodology exhibits versatility and can handle a broad spectrum of gradient flows while accommodating diverse potential functions and nonlinear mobility scenarios.

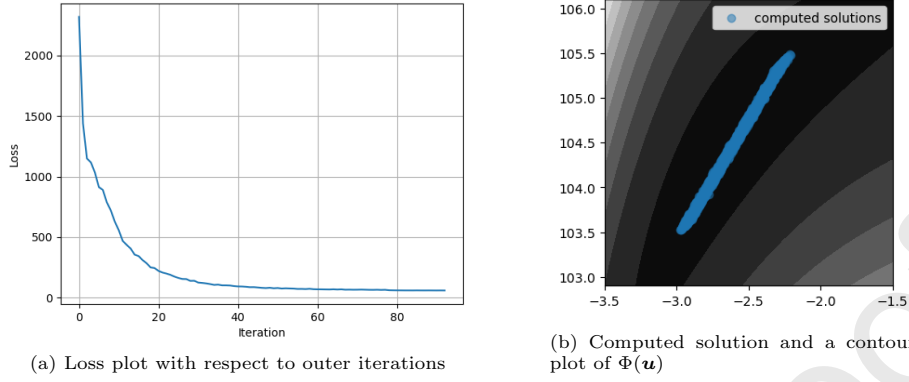


Figure 10: The loss function values across iterations are depicted on the x-axis, while the computed solution at $t = 20$ is superimposed on the contour plot generated from $\Phi(\mathbf{u})$ (on the right).

There are several promising directions for future work. Firstly, there is potential for accelerating the training process. This could involve leveraging the properties of the velocity field as suggested in [32], or exploring transfer learning techniques [43], and improving resampling methods. Another avenue for exploration is the convexity of the problem, investigating how it depends on parameters like the step size Δt and network architectures, including the choices of neural network activation functions. Understanding the convexity properties can contribute to faster training and provide insights into the convergence properties of the algorithm. Additionally, we aim to extend this approach to accommodate more complex energy functionals and scenarios where a gradient flow structure is not fully presented or is only part of the physical dynamics. This extension would broaden the applicability of our method to a wider range of problems in simulating general physics oriented partial differential equations.

Acknowledgements

W. Lee acknowledges support from the National Institute of Standards and Technology (NIST) under award number 70NANB22H021 and partially supported by AFOSR YIP award No. FA9550-23-1-0087. L. Wang acknowledges support from NSF grant DMS-1846854. W. Li's work is supported by AFOSR MURI FP 9550-18-1-502, AFOSR YIP award No. FA9550-23-1-0087, NSF RTG: 2038080, and NSF DMS-2245097. L. Wang and W. Li would like to express their gratitude for the hospitality extended by AIM, where part of the work was discussed.

References

- [1] J. BENAMOU AND Y. BRENIER, *A computational fluid mechanics solution*

- to the Monge-Kantorovich mass transfer problem, *Numer. Math.*, 84 (2000), pp. 375–393.
- [2] J. BENAMOU, G. CARLIER, AND M. LABORDE, *An augmented Lagrangian approach to Wasserstein gradient flows and applications*, *ESAIM: PROCEEDINGS AND SURVEYS*, 54 (2016), pp. 1–17.
 - [3] A. L. BERTOZZI, *The mathematics of moving contact lines in thin liquid films*, *Notices of the AMS*, 45 (1998), pp. 689–697.
 - [4] N. M. BOFFI AND E. VANDEN-EIJNDEN, *Probability flow solution of the fokker-planck equation*, *Machine Learning: Science and Technology*, 4 (2023), p. 035012.
 - [5] M. BURGER, M. DI FRANCESCO, AND Y. DOLAK-STRUSS, *The keller-segel model for chemotaxis with prevention of overcrowding: Linear vs. nonlinear diffusion*, *SIAM Journal on Mathematical Analysis*, 38 (2006), pp. 1288–1315.
 - [6] J. W. CAHN, *On spinodal decomposition*, *Acta metallurgica*, 9 (1961), pp. 795–801.
 - [7] C. CANCES, T. O. GALLOUËT, AND G. TODESCHI, *A variational finite volume scheme for wasserstein gradient flows*, *Numerische Mathematik*, 146 (2020), pp. 437–480.
 - [8] J. CARRILLO, S. LISINI, G. SAVARÉ, AND D. SLEPČEV, *Nonlinear mobility continuity equations and generalized displacement convexity*, *Journal of Functional Analysis*, 258 (2010), pp. 1273–1309.
 - [9] J. A. CARRILLO, K. CRAIG, AND F. S. PATACCHINI, *A blob method for diffusion*, *Calculus of Variations and Partial Differential Equations*, 58 (2019), pp. 1–53.
 - [10] J. A. CARRILLO, K. CRAIG, L. WANG, AND C. WEI, *Primal dual methods for Wasserstein gradient flows*, *Found. Comput. Math.*, 22 (2022), pp. 389–443.
 - [11] J. A. CARRILLO, D. MATTHES, AND M.-T. WOLFRAM, *Lagrangian schemes for Wasserstein gradient flows*, in *Handbook of Numerical Analysis*, vol. 22, Elsevier, 2021, pp. 271–311.
 - [12] J. A. CARRILLO AND J. S. MOLL, *Numerical simulation of diffusive and aggregation phenomena in nonlinear continuity equations by evolving diffeomorphisms*, *SIAM Journal on Scientific Computing*, 31 (2010), pp. 4305–4329.
 - [13] J. A. CARRILLO, H. RANETBAUER, AND M.-T. WOLFRAM, *Numerical simulation of nonlinear continuity equations by evolving diffeomorphisms*, *Journal of Computational Physics*, 327 (2016), pp. 186–202.

- [14] J. A. CARRILLO, L. WANG, AND C. WEI, *Structure preserving primal dual methods for gradient flows with nonlinear mobility transport distances*, (2023).
- [15] R. T. CHEN, Y. RUBANOVA, J. BETTENCOURT, AND D. K. DUVENAUD, *Neural ordinary differential equations*, Advances in neural information processing systems, 31 (2018).
- [16] J. DOLBEAULT, B. NAZARET, AND G. SAVARÉ, *A new class of transport distances between measures*, Calculus of Variations and Partial Differential Equations, 34 (2009), pp. 193–231.
- [17] S. FAGIOLI AND O. TSE, *On gradient flow and entropy solutions for non-local transport equations with nonlinear mobility*, Nonlinear Analysis, 221 (2022), p. 112904.
- [18] J. FAN, A. TAGHVAEI, AND Y. CHEN, *Variational wasserstein gradient flow*, arXiv preprint arXiv:2112.02424, (2021).
- [19] G. FU, S. OSHER, AND W. LI, *High order spatial discretization for variational time implicit schemes: Wasserstein gradient flows and reaction-diffusion systems*, arXiv:2303.08950 [math.NA], (2023).
- [20] W. GANGBO AND R. J. MCCANN, *The geometry of optimal transportation*, Acta Math., 177 (1996), pp. 113–161.
- [21] A. GARBUNO-INIGO, F. HOFFMANN, W. LI, AND A. M. STUART, *Interacting Langevin Diffusions: Gradient Structure and Ensemble Kalman Sampler*, SIAM Journal on Applied Dynamical Systems, 19 (2020), pp. 412–441.
- [22] W. GRATHWOHL, R. T. CHEN, J. BETTENCOURT, I. SUTSKEVER, AND D. DUVENAUD, *Fjord: Free-form continuous dynamics for scalable reversible generative models*, arXiv preprint arXiv:1810.01367, (2018).
- [23] J. HAN, A. JENTZEN, AND W. E, *Solving high-dimensional partial differential equations using deep learning*, Proceedings of the National Academy of Sciences, 115 (2018), pp. 8505–8510.
- [24] M. HERTY AND G. VISCONTI, *Kinetic methods for inverse problems.*, Kinetic & Related Models, 12 (2019).
- [25] Z. HU, C. LIU, Y. WANG, AND Z. XU, *Energetic Variational Neural Network Discretizations to Gradient Flows*, 2022.
- [26] 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.
- [27] Y. LECUN, Y. BENGIO, AND G. HINTON, *Deep learning*, nature, 521 (2015), pp. 436–444.

- [28] Q. LI, L. WANG, AND Y. YANG, *Differential-equation constrained optimization with stochasticity*, arXiv preprint arXiv:2305.04024, (2023).
- [29] W. LI, J. LU, AND L. WANG, *Fisher information regularization schemes for Wasserstein gradient flows*, J. Comput. Phys., 416 (2020), pp. 109449, 24.
- [30] A. T. LIN, W. LI, S. OSHER, AND G. MONTÚFAR, *Wasserstein proximal of gans*, in Geometric Science of Information, F. Nielsen and F. Barbaresco, eds., Cham, 2021, Springer International Publishing, pp. 524–533.
- [31] S. LIU, W. LI, H. ZHA, AND H. ZHOU, *Neural parametric fokker-planck equation*, SIAM Journal on Numerical Analysis, 60 (2022), pp. 1385–1449.
- [32] X. LIU, C. GONG, AND Q. LIU, *Flow straight and fast: Learning to generate and transfer data with rectified flow*, arXiv preprint arXiv:2209.03003, (2022).
- [33] P. MOKROV, A. KOROTIN, L. LI, A. GENEVAY, J. M. SOLOMON, AND E. BURNAEV, *Large-scale wasserstein gradient flows*, Advances in Neural Information Processing Systems, 34 (2021), pp. 15243–15256.
- [34] D. ONKEN, S. W. FUNG, X. LI, AND L. RUTHOTTO, *Ot-flow: Fast and accurate continuous normalizing flows via optimal transport*, in Proceedings of the AAAI Conference on Artificial Intelligence, vol. 35, 2021, pp. 9223–9232.
- [35] 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.
- [36] L. RUTHOTTO, S. J. OSHER, W. LI, L. NURBEKYAN, AND S. W. FUNG, *A machine learning framework for solving high-dimensional mean field game and mean field control problems*, Proceedings of the National Academy of Sciences, 117 (2020), pp. 9183–9193.
- [37] Z. SHEN, Z. WANG, S. KALE, A. RIBEIRO, A. KARBASI, AND H. HASSANI, *Self-consistency of the fokker planck equation*, in Conference on Learning Theory, PMLR, 2022, pp. 817–841.
- [38] J. SIRIGNANO AND K. SPILIOPOULOS, *Dgm: A deep learning algorithm for solving partial differential equations*, Journal of computational physics, 375 (2018), pp. 1339–1364.
- [39] A. VIDAL, S. WU FUNG, L. TENORIO, S. OSHER, AND L. NURBEKYAN, *Taming hyperparameter tuning in continuous normalizing flows using the JKO scheme*, Scientific Reports, 13 (2023), p. 4501.

- [40] C. VILLANI ET AL., *Optimal transport: old and new*, vol. 338, Springer, 2009.
- [41] L. WANG AND M. YAN, *Hessian informed mirror descent*, Journal of Scientific Computing, 92 (2022), p. 90.
- [42] C. XU, X. CHENG, AND Y. XIE, *Normalizing flow neural networks by JKO scheme*, in Thirty-seventh Conference on Neural Information Processing Systems, 2023.
- [43] W. XU, Y. LU, AND L. WANG, *Transfer learning enhanced deeponet for long-time prediction of evolution equations*, in Proceedings of the AAAI Conference on Artificial Intelligence, vol. 37, 2023, pp. 10629–10636.
- [44] B. YU AND W. E, *The deep ritz method: a deep learning-based numerical algorithm for solving variational problems*, Communications in Mathematics and Statistics, 6 (2018), pp. 1–12.

Declaration of interests

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

☐ The authors declare the following financial interests/personal relationships which may be considered as potential competing interests: