

Proceedings Track

A Variational Manifold Embedding Framework for Nonlinear Dimensionality Reduction

Editors: List of editors' names

Abstract

Dimensionality reduction algorithms like principal component analysis (PCA) are workhorses of machine learning and neuroscience, but each has well-known limitations. Variants of PCA are simple and interpretable, but not flexible enough to capture nonlinear data manifold structure. More flexible approaches have other problems: autoencoders are generally difficult to interpret, and graph-embedding-based methods can produce pathological distortions in manifold geometry. Motivated by these shortcomings, we propose a variational framework that casts dimensionality reduction algorithms as solutions to an optimal manifold embedding problem. By construction, this framework permits nonlinear embeddings, allowing its solutions to be more flexible than PCA. Moreover, the variational nature of the framework has useful consequences for interpretability: each solution satisfies a set of partial differential equations, and can be shown to reflect symmetries of the embedding objective. We discuss these features in detail and show that solutions can be analytically characterized in some cases. Interestingly, one special case exactly recovers PCA.

Keywords: dimensionality reduction, manifold embedding, principal component analysis

1. Introduction

Dimensionality reduction—the problem of identifying useful low-dimensional representations of high-dimensional data—is a central concern across machine learning, statistics, and neuroscience. But existing approaches have issues with accuracy, interpretability, or both. For example, principal component analysis (PCA) (Pearson, 1901; Hotelling, 1933; Bishop, 2006) is computationally cheap and can be interpreted in terms of a simple generative model (Tipping and Bishop, 1999), but cannot capture nonlinear structure in data.

Two complementary perspectives have motivated lines of work that go beyond PCA. The *geometric* perspective on dimensionality reduction casts the problem as identifying the low-dimensional structure (or *data manifold*) along which data tend to lie. Methods like Locally Linear Embedding (Roweis and Saul, 2000), Isomap (Tenenbaum et al., 2000), and diffusion maps (Coifman and Lafon, 2006) exploit this perspective to learn a low-dimensional data manifold embedded in the high-dimensional ambient space. Work on so-called neural manifolds, which seeks low-dimensional representations of neural data, also tends to take this perspective (Cunningham and Yu, 2014; Chung and Abbott, 2021; Perich et al., 2025).

The *probabilistic generative modeling* perspective on dimensionality reduction casts the problem as minimizing a divergence between the data distribution (or a suitably transformed version of it) and a distribution on a low-dimensional ‘latent’ space. Neural-network-based autoencoders (LeCun, 1987; Hinton and Zemel, 1993) and variational autoencoders (Kingma and Welling, 2013; Higgins et al., 2017) exploit this perspective in order to capture arbitrarily complex (smooth) structure in data, but may yield results that are difficult to interpret.

Proceedings Track

Methods which involve embedding samples into a graph, like Uniform Manifold Approximation and Projection (UMAP) (McInnes et al., 2018) and t-distributed stochastic neighbor embedding (t-SNE) (Maaten and Hinton, 2008), are nonlinear and somewhat easier to interpret, but have been observed to produce pathologies in downstream statistical and clustering analyses (Wattenberg et al., 2016; Chari and Pachter, 2023).

Motivated by these two perspectives, here we present a novel, nonlinear framework for dimensionality reduction that (i) involves minimizing a divergence between latent space distributions, like variational autoencoders; and (ii) constrains that optimization problem by making geometry-related *nonparametric* assumptions about the mapping between latent and ambient space. We show that this framework has a variety of reasonable features, is amenable to mathematical analysis, and includes PCA as a special case. Critically, the variational structure of the problem permits us to use tools from physics to reason about properties of optimal embeddings, including that symmetries of the data distribution constrain them, and that they satisfy a certain set of partial differential equations (PDEs).

2. A tractable variational framework for dimensionality reduction

When are two data manifolds the same? Let $\mathbf{x} \in \mathbb{R}^D$ denote an observation in a D -dimensional space, and $p_{\text{data}}(\mathbf{x})$ the associated data likelihood. When is p_{data} more or less identical to some other likelihood q ? Or, in more explicitly geometric language, when are the ‘data manifolds’ associated with these likelihoods equivalent? One reasonable notion of equivalence comes from demanding that q and p_{data} can be obtained from one another via a smooth bijection (i.e., a *diffeomorphism*) $\mathbf{g} : \mathbb{R}^D \rightarrow \mathbb{R}^D$. In particular, we ask that

$$q(\mathbf{x}) d\mathbf{x} = p_{\text{data}}(\mathbf{g}(\mathbf{x})) |\det \mathbf{J}_{\mathbf{g}}(\mathbf{x})| d\mathbf{x} \quad (1)$$

almost everywhere, where $\mathbf{J}_{\mathbf{g}}(\mathbf{x})$ is the $D \times D$ Jacobian matrix whose entries are $J_{ij}(\mathbf{x}) := \partial g_i(\mathbf{x}) / \partial x_j$. We will call the distribution $(\mathbf{g}^* p_{\text{data}})(\mathbf{x}) := p_{\text{data}}(\mathbf{g}(\mathbf{x})) |\det \mathbf{J}_{\mathbf{g}}(\mathbf{x})|$ the *pullback* of p_{data} by $\mathbf{g}$. Intuitively, the above equation says that one can obtain p_{data} from q by a smooth deformation of the state space $\mathbb{R}^D$.

Dimensionality reduction involves attempting to capture *most* of the structure of p_{data} by relating it to a distribution defined on a generally lower-dimensional space $\mathbb{R}^d$ for some $d \leq D$. Let $\mathbf{z} \in \mathbb{R}^d$ denote an element of this ‘latent’ space. Since it is usually impossible to capture *all* high-dimensional structure of p_{data} via a lower-dimensional approximation, we can only ask that $q(\mathbf{z})d\mathbf{z}$ is *approximately* the same as the pullback of p_{data} through a smooth embedding map $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$ (which we denote $\phi^* p_{\text{data}}$). More precisely, we want

$$q(\mathbf{z}) d\mathbf{z} \approx p_{\text{data}}(\phi(\mathbf{z})) \sqrt{\det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z}))} d\mathbf{z} \quad (2)$$

almost everywhere, where $\mathbf{J}(\mathbf{z})$ is the $D \times d$ Jacobian with entries $J_{ij}(\mathbf{z}) := \partial \phi_i / \partial z_j$. Notice that, since we may have $d < D$, the change of variables correction involves the Riemannian volume form $\sqrt{\det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z}))}$ (we cannot use $|\det \mathbf{J}(\mathbf{z})|$, since it may equal zero).

A variational objective for dimensionality reduction. Let $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$ be a smooth injection from a ‘latent’ space to an ‘ambient’ space (with $d \leq D$), and suppose that $\mathbf{z} \in \mathbb{R}^d$ is distributed according to a *prior* $q(\mathbf{z})$, and $\mathbf{x} \in \mathbb{R}^D$ is distributed according to a likelihood

Proceedings Track

p_{data} . Motivated by Eq. 2, we define an *optimal dimensionality reduction algorithm* (or *optimal embedding*) as a map ϕ which *maximizes*

$$J[\phi] = \int_{\mathbb{R}^d} dz \, q(z) \left[\frac{1}{2} \log \det(\mathbf{J}(z)^T \mathbf{J}(z)) + \log p_{\text{data}}(\phi(z)) - \log q(z) \right], \quad (3)$$

i.e., the negative of the Kullback-Leibler (KL) divergence between q and $\phi^* p_{\text{data}}$, while satisfying the finiteness conditions (see Appendix A)

$$\mathbb{E}_{z \sim q} \left\{ \sqrt{\det(\mathbf{J}(z)^T \mathbf{J}(z))} \right\} < \infty \quad \text{and} \quad \mathbb{E}_{z \sim q} \left\{ \|\phi(z)\|_2^2 \right\} < \infty. \quad (4)$$

Because $J = -D_{\text{KL}}(q \| \phi^* p_{\text{data}})$, we immediately know that J is upper bounded by zero¹, which corresponds to perfect performance. See Figure 1a for a schematic.

This objective has two terms, each of which perform slightly different roles. The log-likelihood term encourages ϕ to map into regions of high likelihood, while the log-determinant term encourages embeddings to have a nontrivial volume. (If only the log-likelihood term were present, an optimal embedding may map all z near a global maximum of p_{data} .) Interestingly, these terms also have clear physical analogues: the log-likelihood term functions like a *potential energy*, and the log-determinant term, which depends on derivatives of the embedding, functions like a *kinetic energy*.

As in physics, the fact that the objective depends not just on the embedding coordinates $\phi_i(z)$, but also on the derivatives $J_{ij} := \partial \phi_i / \partial z_j$, allows us to directly apply results from the calculus of variations (Goldstein, 1980; Elsgolc, 2007) that cannot be usefully leveraged in more generic (e.g., autoencoder) settings. For more discussion of this objective, including a Bayesian interpretation of it that casts optimizing J as identifying the maximum a posteriori estimate of ϕ , see Appendix A.

Examples of optimal one-dimensional embeddings. What do optimal embeddings look like? The naive way to estimate them is to parameterize ϕ and numerically optimize J via a gradient-descent-like algorithm. If we parameterize ϕ as a multilayer perceptron and choose p_{data} to be different Gaussian mixtures, optimizing J in the standard way tells us that optimal *one-dimensional* embeddings are generally nonlinear, and correspond to what one might expect: if there is an obvious path through the points, the optimal embedding follows it (Figure 1b).

3. Interpreting optimal one-dimensional embeddings

We can say more about optimal one-dimensional embeddings, and optimal embeddings more generally, by exploiting results from the calculus of variations. The most basic such result is that optimal embeddings satisfy the *Euler-Lagrange* (EL) *equations*, which we define next.

3.1. Optimal one-dimensional embeddings follow score vectors

If we approximate p_{data} with a one-dimensional manifold ($d = 1$), the objective becomes

$$J[\phi] = \int_{-\infty}^{\infty} dz \, q(z) \left[\frac{1}{2} \log \|\phi'(z)\|_2^2 + \log p_{\text{data}}(\phi(z)) - \log q(z) \right] \quad (5)$$

1. Note that this is only true if ϕ is an injection. See Appendix A.

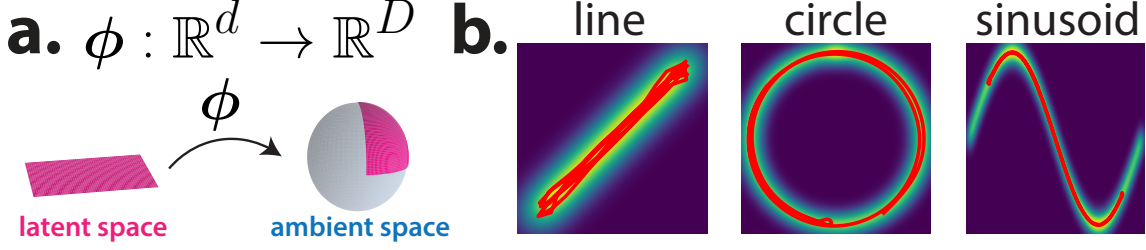


Figure 1: **a.** The function ϕ maps the latent space to the (higher-dimensional) ambient space. **b.** Optimal one-dimensional embeddings when p_{data} is a mixture of Gaussians centered at points along a line (left), circle (middle), or sinusoid (right).

where $\phi'(z) := (\phi'_1(z), \dots, \phi'_D(z))^T$ is the derivative of the embedding. This version of the objective looks extremely similar to a classical mechanics problem, with the main qualitative difference being that the kinetic term is $\frac{1}{2} \log \|\phi'(z)\|_2^2$ instead of $\frac{1}{2} \|\phi'(z)\|_2^2$.

We can view ϕ as tracing out a one-dimensional trajectory through $\mathbb{R}^D$. According to a standard result in the calculus of variations, these trajectories satisfy the EL equations

$$\frac{d}{dz} \left(\frac{\partial L}{\partial \phi'_i} \right) = \frac{\partial L}{\partial \phi_i} \quad , \text{ i.e., } \quad \frac{d}{dz} \left[q(z) \frac{\phi'_i(z)}{\|\phi'(z)\|_2^2} \right] = q(z) \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z)) \quad (6)$$

for each $i = 1, \dots, D$ (Goldstein, 1980; Elsgolc, 2007), where the *Lagrangian* L is defined as

$$L(\phi, \phi', z) := q(z) \left[\frac{1}{2} \log \|\phi'(z)\|_2^2 + \log p_{\text{data}}(\phi(z)) - \log q(z) \right]. \quad (7)$$

These equations place nontrivial constraints on the form of optimal embeddings; satisfying them is necessary, but not sufficient, for ϕ to be a global maximizer of J . If we define momentum-like variables $p_i(z) := \phi'_i(z) / \|\phi'(z)\|_2^2$, we find that the ‘dynamics’ defined by the EL equations looks like that of a massive particle moving through the landscape determined by the potential $-\log p_{\text{data}}$:

$$\phi'_i(z) = \frac{p_i(z)}{\|\mathbf{p}(z)\|_2^2} \quad p'_i(z) = s_i(\phi(z)) - \left[\frac{d}{dz} \log q(z) \right] p_i(z) \quad (8)$$

where $s_i(\phi(z)) := \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z))$ is the *score vector* associated with p_{data} at the point $\phi(z) \in \mathbb{R}^D$. These vectors, which play a foundational role in how diffusion models sample (Song et al., 2021) and generalize (Wang and Vastola, 2024; Vastola, 2025a), point in the direction in which the data likelihood increases most. Eq. 8 says that the trajectory $\phi(z)$ moves in the direction of the score $\mathbf{s}(\phi(z))$, and that the influence of the score vector is controlled by a q -dependent ‘friction’ factor and the norm of $\mathbf{p}(z)$. The norm $\|\mathbf{p}(z)\|_2^2$ plays the role of a ‘mass’. When $\mathbf{p}(z)$ is large (which can happen when the score vectors become large), the trajectory changes very slowly; when it is small, the trajectory changes quickly.

Proceedings Track

3.2. Connection to diffusion models

We can make the connection between optimal one-dimensional embeddings and diffusion models more concrete by observing that one significant difference between Eq. 8 and a probability flow ODE (Song et al., 2021) is that the former involves *inertia*: $\phi(z)$ does not *only* move in the direction of the score vector. Since the amount of inertia that appears in Eq. 8 depends on the prior q , we can modulate it by choosing a special prior. If we pick $q(z) = \gamma e^{\gamma z}$ for $z \in (-\infty, 0]$, in the large γ limit we obtain dynamics (see Appendix C)

$$\phi'_i(z) = \gamma \frac{s_i(\phi(z))}{\|\mathbf{s}(\phi(z))\|_2^2} \quad (9)$$

which follows score vectors, but moves faster or slower depending on their magnitude. If we reparameterize ‘time’, we recover an ODE that follows score vectors without the norm-dependent distortion, and hence looks like a probability flow ODE at a fixed noise scale.

4. Useful properties of the variational objective

What other properties of optimal embeddings can we reason about without actually optimizing the objective? In this section, we discuss two kinds: properties related to *flexibility*, and properties related to *interpretability*. Unlike in the previous section, here our comments apply for any choice of latent dimensionality (i.e., any $d \leq D$).

4.1. Flexibility-related properties

Objective is geometry-sensitive but not coordinate-sensitive. Our framework ought to reflect the geometry of the data manifold and prior, but not our incidental choice of coordinates²; in particular, if we relabeled each $\mathbf{z} \in \mathbb{R}^d$ according to a diffeomorphism $g : \mathbb{R}^d \rightarrow \mathbb{R}^d$, up to this relabeling, we would like J to have the same value and solutions. This reparameterization-invariance property trivially holds, since J equals a (negative) KL divergence, and the property holds for KL divergences. It is also straightforward to show that this property holds through direct computation (see Appendix D).

Objective has expected trivial solutions. Because J is bounded from above by zero, we know that any choice of ϕ that makes it achieve this bound is a global optimum. This property makes it clear that J has two types of reasonable trivial solutions. First, when $d = D$ and the prior q equals the data likelihood p_{data} —i.e., when there is no dimensionality reduction—the identity map $\phi(\mathbf{z}) = \mathbf{z}$ is optimal since it makes J equal zero. Second, if p_{data} is genuinely low-dimensional (say, d -dimensional), and there exists a smooth injection ψ for which the pullback $\psi^* p_{\text{data}}$ is normalized on $\mathbb{R}^d$, then $\phi = \psi$ is a global optimum. Importantly, this is true even if ψ is highly nonlinear.

2. In this respect, embedding is similar to other geometry-sensitive problems, like close-packing (Viazovska, 2017; Vastola, 2024).

Proceedings Track

4.2. Interpretability-related properties

Optimal embeddings satisfy EL equations. The Lagrangian density corresponding to our objective (i.e., its integrand) is

$$\mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) := q(\mathbf{z}) \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right]. \quad (10)$$

A useful consequence of our variational formulation is that it allows us to exploit mathematical results from the calculus of variations. In particular, optimal mappings ϕ satisfy the Euler-Lagrange (EL) equations (see Appendix B)

$$\sum_{j=1}^d \partial_j \left(\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \phi_i}, \text{ i.e., } \sum_j \left[\partial_j (J_{ji}^+) + J_{ji}^+ \partial_j [\log q(\mathbf{z})] \right] = \frac{\partial \log p_{\text{data}}(\phi(\mathbf{z}))}{\partial \phi_i}, \quad (11)$$

where $\mathbf{J}^+ := (\mathbf{J}^T \mathbf{J})^{-1} \mathbf{J}^T \in \mathbb{R}^{d \times D}$ is the Moore-Penrose pseudoinverse of $\mathbf{J}$ (which satisfies $\mathbf{J}^+ \mathbf{J} = \mathbf{I}_d$), and where we have relied on a well-known result on the derivative of a log-determinant (Petersen and Pedersen, 2008).

The EL equations define a system of D coupled nonlinear PDEs: there is one equation for each $\phi_i(\mathbf{z})$, and only up to second derivatives of the ϕ_k appear. (Unfortunately, due to the presence of $\mathbf{J}^+$, one may encounter derivatives of the ϕ_k in denominators.) Satisfying the EL equations is necessary but not sufficient for ϕ to maximize J , which allows one to use these PDEs to strongly constrain the solution space.

Continuous symmetries of the objective constrain solutions. We expect that symmetries of the data likelihood constrain how one ought to perform dimensionality reduction. For example, the dimensionality-reduced version of a radially symmetric likelihood is also probably radially symmetric. *Noether's theorem* (Noether, 1918; Neuenschwander, 2017; Vastola, 2025b), a powerful result frequently used in field theory settings, makes this intuition precise. It states that if a small perturbation $\phi(\mathbf{z}) \rightarrow \phi(\mathbf{z}) + \delta \phi(\mathbf{z})$ of the embedding causes the Lagrangian density to change by a total derivative, i.e., $\mathcal{L} \rightarrow \mathcal{L} + \nabla \cdot \mathbf{K}$ for some vector $\mathbf{K}$, then the quantity (the *Noether current*)

$$j_a := \sum_{b=1}^D \frac{\partial \mathcal{L}}{\partial (\partial_a \phi_b)} (\delta \phi_b) - K_a \quad (12)$$

is conserved (i.e., $\nabla \cdot \mathbf{j} = 0$) when the EL equations are satisfied (see Appendix E). Noether's theorem is useful because it allows us to make concrete statements about how continuous symmetries constrain optimal embeddings, even if and especially when the EL equations are too difficult to solve analytically.

5. Conserved quantities and consequences of continuous symmetries

In this section, we study explicit examples of how continuous symmetries constrain optimal embeddings. We consider three kinds of continuous symmetry: *reparameterization invariance*, which yields an analogue of energy conservation; *embedding translation invariance*, which yields an analogue of momentum conservation; and *embedding rotation invariance*, which yields an analogue of angular momentum conservation.

Proceedings Track

5.1. Embedding energy is conserved for optimal embeddings

If q is uniform, the Lagrangian density is locally reparameterization invariant, since an infinitesimal change of coordinates changes $\mathcal{L}$ by a total derivative. By Noether's theorem, this implies that the *stress-energy tensor*

$$T_{ij} = \sum_k \frac{\partial \mathcal{L}}{\partial (\partial_i \phi_k)} (\partial_j \phi_k(\mathbf{z})) - \delta_{ij} \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) = \sum_k J_{ik}^+ J_{kj} q - \delta_{ij} \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) = \delta_{ij} (q - \mathcal{L}) \quad (13)$$

is conserved ($\sum_i \partial_i T_{ij} = 0$). But since T_{ij} is only nonzero when $i = j$, conservation just means that $\mathcal{L} = \text{const}$. Hence, $\mathcal{L}$ being conserved is analogous to the conservation of energy.

Usefully, this result can be combined with the reparameterization invariance of the objective to establish a more general result that holds even for non-uniform q (one can also derive this result through a direct computation; see Appendix E): the **embedding energy**

$$\boxed{\mathcal{E} := -\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) - \log p_{\text{data}}(\phi(\mathbf{z})) + \log q(\mathbf{z})} \quad (14)$$

is constant (i.e., does not depend on $\mathbf{z}$) for all solutions to the EL equations (and in particular, for the global maximizers of J). Interestingly, since

$$0 \geq J[\phi] = - \int_{\mathbb{R}^d} d\mathbf{z} \, q(\mathbf{z}) \, \mathcal{E} = -\mathcal{E} \, , \quad (15)$$

embedding energy is always nonnegative, and global maximizers of J *minimize* it.

5.2. Translation invariance yields momentum conservation analogue

If p_{data} is uniform along some ambient space direction, so that it does not depend on ϕ_i , Noether's theorem implies that the *canonical momentum* along that direction is conserved:

$$\sum_j \partial_j \left(\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} q(\mathbf{z}) \right) = \sum_j \partial_j \left[J_{ji}^+ q(\mathbf{z}) \right] = 0 \, . \quad (16)$$

This conservation law almost fully characterizes the optimal embedding when p_{data} is uniform (on a ball of radius one-half centered at $\mathbf{x} = (1/2, \dots, 1/2)^T$, say) and $d = 1$. Together with embedding energy conservation, it implies $\phi'_i(\mathbf{z}) = \text{const.}/q(\mathbf{z})$, so

$$\phi_i(\mathbf{z}) = \int_{-\infty}^{\mathbf{z}} q(y) \, dy \quad (17)$$

for all $i = 1, \dots, D$. Interestingly, this mapping exactly matches the one used by efficient codes to convert a generally non-uniform stimulus distribution to a uniform distribution.

5.3. Rotation invariance yields angular momentum conservation analogue

If p_{data} is invariant to rotations in the ϕ_i - ϕ_j plane, Noether's theorem implies that

$$\sum_a \partial_a L_{ij}^a = 0 \quad \text{where} \quad L_{ij}^a := q(\mathbf{z}) \left[J_{ai}^+ \phi_j - J_{aj}^+ \phi_i \right] \, . \quad (18)$$

Proceedings Track

When $d = 1$, $\mathbf{J} \rightarrow \phi'$ and $\mathbf{J}^+ \rightarrow \phi' / \|\phi'\|_2^2$, so angular momentum conservation becomes

$$q(z) \frac{[\phi'_1(z)\phi_2(z) - \phi_1(z)\phi'_2(z)]}{\|\phi'(z)\|_2^2} = \text{const.} \quad (19)$$

When combined with embedding energy conservation, this conservation law can be used to show that optimal embeddings lie along circles when $D = 2$. Hence, the optimal embedding for a ring attractor structure, as expected, follows the ring.

6. A solvable case: PCA is optimal for Gaussian prior and likelihood

Optimal embeddings are difficult to determine analytically, in part because the EL equations are generally a system of highly nonlinear coupled PDEs. Do there exist any interesting solvable cases? Perhaps surprisingly, the answer is yes: the optimal embedding exactly corresponds to PCA when (i) the prior is an isotropic normal distribution $q(\mathbf{z}) = \mathcal{N}(\mathbf{z}; \mathbf{0}, \sigma_0^2 \mathbf{I})$, with $\sigma_0 > 0$; and (ii) the likelihood is a multivariate Gaussian $p_{\text{data}}(\mathbf{x}) = \mathcal{N}(\mathbf{x}; \mathbf{0}, \mathbf{\Sigma})$, with $\mathbf{\Sigma}$ a positive definite $D \times D$ matrix.

The PCA solution corresponds to the linear mapping $\phi(\mathbf{z}) = \sum_{k=1}^d \frac{\sqrt{\lambda_k}}{\sigma_0} (\mathbf{q}_k^T \mathbf{z}) \mathbf{v}_k$, where the $\mathbf{v}_k$ are the top d eigenvectors of the covariance matrix, and $\{\mathbf{q}_k\}_{k=1,\dots,d}$ is any set of orthonormal vectors (i.e., $\mathbf{q}_k^T \mathbf{q}_\ell = \delta_{k\ell}$). Although it is not mathematically obvious that the EL equations admit such a simple solution, the intuition is fairly clear: one can obtain an anisotropic Gaussian from an isotropic one by uniformly stretching it in different directions.

The one-dimensional case. Let us begin with the relatively simple $d = 1$ case, with

$$J[\phi] := \int_{-\infty}^{\infty} dz \frac{e^{-\frac{1}{2\sigma_0^2} z^2}}{\sqrt{2\pi\sigma_0^2}} \left\{ \frac{1}{2} \log \|\phi'(z)\|_2^2 - \frac{1}{2} \phi(z)^T \mathbf{\Sigma}^{-1} \phi(z) \right\} + \text{const.} \quad (20)$$

If the optimal ϕ *does* correspond to PCA, it should pick out the first principal component, i.e., an eigenvector of $\mathbf{\Sigma}$ whose eigenvalue is largest.³ The EL equations read

$$\frac{d}{dz} \left[q(z) \frac{\phi'(z)}{\|\phi'(z)\|_2^2} \right] = -q(z) \mathbf{\Sigma}^{-1} \phi(z). \quad (21)$$

Consider a linear ansatz $\phi(z) = \mathbf{v}z$ for some vector $\mathbf{v} \in \mathbb{R}^D$. Since $\phi'(z) = \mathbf{v}$, we require

$$\frac{d}{dz} \left[q(z) \frac{\mathbf{v}}{\|\mathbf{v}\|_2^2} \right] = -q(z) \mathbf{\Sigma}^{-1} \mathbf{v}z \quad \implies \quad \mathbf{\Sigma} \mathbf{v} = (\sigma_0^2 \|\mathbf{v}\|_2^2) \mathbf{v} \quad (22)$$

i.e., that $\mathbf{v}$ must be an eigenvector of $\mathbf{\Sigma}$ with eigenvalue $\lambda = (\sigma_0^2 \|\mathbf{v}\|_2^2)$. Hence, the EL equations force $\mathbf{v}$ to be *one* of the eigenvectors of $\mathbf{\Sigma}$, but not any specific eigenvector. The symmetry is broken by asking that ϕ maximize J (Eq. 20). Its ϕ -dependent terms equal

$$\mathbb{E}_{z \sim q} \left\{ -\frac{z^2}{2} \mathbf{v}^T \mathbf{\Sigma}^{-1} \mathbf{v} + \frac{1}{2} \log(\mathbf{v}^T \mathbf{v}) \right\} = -\frac{\sigma_0^2}{2} \mathbf{v}^T \mathbf{\Sigma}^{-1} \mathbf{v} + \frac{1}{2} \log(\mathbf{v}^T \mathbf{v}) = -\frac{1}{2} + \frac{1}{2} \log \left(\frac{\lambda}{\sigma_0^2} \right) \quad (23)$$

given our ansatz, which is maximized when λ is as large as possible. This happens precisely when $\mathbf{v}$ is the top eigenvector of $\mathbf{\Sigma}$, which (almost) finishes our argument. Do the EL equations admit nonlinear solutions? We can show the answer is no by exploiting embedding energy conservation (see Appendix F), which means $\phi(z) = z\mathbf{v}$ is indeed optimal.

3. If this is not unique, we expect multiple solutions.

Proceedings Track

The general case. For arbitrary $d \leq D$, the objective is

$$J[\phi] := \int_{\mathbb{R}^d} dz \frac{e^{-\frac{1}{2\sigma_0^2}\|z\|_2^2}}{(2\pi\sigma_0^2)^{d/2}} \left\{ \frac{1}{2} \log \det[\mathbf{J}(z)^T \mathbf{J}(z)] - \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) \right\} + \text{const.} \quad (24)$$

Although the EL equations are more complicated than in the one-dimensional case, if we assume a linear solution $\phi(z) = \mathbf{A}z$ for some $D \times d$ matrix $\mathbf{A}$, we recover an analogous constraint:

$$\Sigma \mathbf{A} = \sigma_0^2 \mathbf{A} \mathbf{A}^T \mathbf{A}. \quad (25)$$

If we exploit a singular value decomposition $\mathbf{A} = \sum_{k=1}^d s_k \mathbf{u}_k \mathbf{q}_k^T$ of $\mathbf{A}$, where $\mathbf{u}_k^T \mathbf{u}_\ell = \mathbf{q}_k^T \mathbf{q}_\ell = \delta_{k\ell}$ and the s_k are the singular values of $\mathbf{A}$, then this condition becomes

$$\Sigma \mathbf{u}_k = (\sigma_0^2 s_k^2) \mathbf{u}_k. \quad (26)$$

for all $k = 1, \dots, d$. As before, this implies each $\mathbf{u}_k$ is an eigenvector of Σ —although they must be *different* eigenvectors, since they are pairwise orthogonal. Also as before, the symmetry is broken by asking that ϕ maximize J . The ϕ -dependent part of J equals

$$\mathbb{E}_{z \sim q} \left\{ -\frac{1}{2} z^T \mathbf{A}^T \Sigma^{-1} \mathbf{A} z + \frac{1}{2} \log \det(\mathbf{A}^T \mathbf{A}) \right\} = -\frac{d}{2} + \frac{1}{2} \sum_{k=1}^d \log \left(\frac{\lambda_k}{\sigma_0^2} \right) \quad (27)$$

given our ansatz, which is maximized when each λ_k is chosen to be as large as possible. This happens precisely when the $\{\mathbf{u}_k\}$ are chosen to be the top d eigenvectors of Σ . Finally, we can argue that the solutions of the EL equations *must* be linear (see Appendix F).

7. Discussion

We introduced a variational framework for dimensionality reduction that is more flexible than PCA—and which includes it as a special case—and which has various features that aid interpretability, including that solutions satisfy PDEs, and that Noether’s theorem links data manifold symmetries to conservation laws. Also, since each dimensionality reduction algorithm associated with the framework optimizes the objective for specific choices of d , q and p_{data} , we can use these quantities to reason about the inductive biases of solutions.

Limitations. At least in the framework’s current form, it assumes that both q and p_{data} are distributions on continuous rather than discrete spaces, making it impossible to directly apply it to (for example) Poisson-like spiking data. On the other hand, since distributions of spikes are plausibly related by Bayes’ rule to the distributions of continuous latent variables (e.g., an animal’s allocentric position in space), an intuition made precise in encoding models like generalized linear models (Weber and Pillow, 2017) and probabilistic population codes (Ma et al., 2006; Vastola et al., 2023), there is still a sense in which the framework can be applied indirectly. Also, while the framework could in principle be used to approximate data manifolds of any dimensionality, it is unclear how scalable it is, and how useful certain analytic insights (e.g., related to the EL equations) are for high-dimensional data sets of interest. Symmetries, for example, strongly constrain low-dimensional embeddings, but not necessarily high-dimensional ones.

Proceedings Track

References

- Christopher M. Bishop. *Pattern Recognition and Machine Learning (Information Science and Statistics)*. Springer-Verlag, Berlin, Heidelberg, 2006. ISBN 0387310738.
- Tara Chari and Lior Pachter. The specious art of single-cell genomics. *PLOS Computational Biology*, 19(8):e1011288, August 2023. doi: 10.1371/journal.pcbi.1011288. URL <https://doi.org/10.1371/journal.pcbi.1011288>. Publisher: Public Library of Science.
- SueYeon Chung and L. F. Abbott. Neural population geometry: An approach for understanding biological and artificial neural networks. *Current Opinion in Neurobiology*, 70: 137–144, 2021. ISSN 0959-4388. doi: <https://doi.org/10.1016/j.conb.2021.10.010>. URL <https://www.sciencedirect.com/science/article/pii/S0959438821001227>.
- Ronald R. Coifman and Stéphane Lafon. Diffusion maps. *Applied and Computational Harmonic Analysis*, 21(1):5–30, 2006. doi: 10.1016/j.acha.2006.04.006.
- John P Cunningham and Byron M Yu. Dimensionality reduction for large-scale neural recordings. *Nature Neuroscience*, 17(11):1500–1509, November 2014. ISSN 1546-1726. doi: 10.1038/nn.3776. URL <https://doi.org/10.1038/nn.3776>.
- Lev D. Elsgolc. *Calculus of variations*. Dover, 2007.
- Herbert Goldstein. *Classical mechanics*. Addison-Wesley, 1980.
- Irina Higgins, Loic Matthey, Arka Pal, Christopher Burgess, Xavier Glorot, Matthew Botvinick, Shakir Mohamed, and Alexander Lerchner. beta-VAE: Learning basic visual concepts with a constrained variational framework. In *International Conference on Learning Representations*, 2017. URL <https://openreview.net/forum?id=Sy2fzU9gl>.
- Geoffrey E Hinton and Richard Zemel. Autoencoders, Minimum Description Length and Helmholtz Free Energy. In J. Cowan, G. Tesauro, and J. Alspector, editors, *Advances in Neural Information Processing Systems*, volume 6. Morgan-Kaufmann, 1993. URL https://proceedings.neurips.cc/paper_files/paper/1993/file/9e3cfc48eccf81a0d57663e129aef3cb-Paper.pdf.
- Harold Hotelling. Analysis of a complex of statistical variables into principal components. *Journal of educational psychology*, 24(6):417, 1933.
- Diederik P Kingma and Max Welling. Auto-Encoding Variational Bayes. *arXiv e-prints*, art. arXiv:1312.6114, December 2013. doi: 10.48550/arXiv.1312.6114.
- Yann LeCun. PhD thesis: Modeles connexionnistes de l’apprentissage (connectionist learning models). 1987.
- Wei Ji Ma, Jeffrey M Beck, Peter E Latham, and Alexandre Pouget. Bayesian inference with probabilistic population codes. *Nature Neuroscience*, 9(11):1432–1438, November 2006. ISSN 1546-1726. doi: 10.1038/nn1790. URL <https://doi.org/10.1038/nn1790>.

Proceedings Track

- Laurens van der Maaten and Geoffrey Hinton. Visualizing Data using t-SNE. *Journal of Machine Learning Research*, 9(86):2579–2605, 2008. URL <http://jmlr.org/papers/v9/vandermaaten08a.html>.
- Leland McInnes, John Healy, and James Melville. UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction. *arXiv e-prints*, art. arXiv:1802.03426, February 2018. doi: 10.48550/arXiv.1802.03426.
- Dwight E Neuenschwander. *Emmy Noether’s wonderful theorem*. JHU Press, 2017.
- E. Noether. Invariante variationsprobleme. *Nachrichten von der Gesellschaft der Wissenschaften zu Göttingen, Mathematisch-Physikalische Klasse*, 1918:235–257, 1918. URL <http://eudml.org/doc/59024>.
- Karl Pearson. Liii. on lines and planes of closest fit to systems of points in space. *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science*, 2(11): 559–572, 1901. doi: 10.1080/14786440109462720. URL <https://doi.org/10.1080/14786440109462720>.
- Matthew G. Perich, Devika Narain, and Juan A. Gallego. A neural manifold view of the brain. *Nature Neuroscience*, 28(8):1582–1597, August 2025. ISSN 1546-1726. doi: 10.1038/s41593-025-02031-z. URL <https://doi.org/10.1038/s41593-025-02031-z>.
- K. B. Petersen and M. S. Pedersen. The matrix cookbook, October 2008. URL <http://www2.imm.dtu.dk/pubdb/p.php?3274>. Version 20081110.
- Sam T. Roweis and Lawrence K. Saul. Nonlinear dimensionality reduction by locally linear embedding. *Science*, 290(5500):2323–2326, 2000. doi: 10.1126/science.290.5500.2323.
- Yang Song, Jascha Sohl-Dickstein, Diederik P Kingma, Abhishek Kumar, Stefano Ermon, and Ben Poole. Score-based generative modeling through stochastic differential equations. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=PxTIG12RRHS>.
- Joshua B. Tenenbaum, Vin de Silva, and John C. Langford. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290(5500):2319–2323, 2000. doi: 10.1126/science.290.5500.2319.
- Michael E. Tipping and Christopher M. Bishop. Probabilistic Principal Component Analysis. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 61(3): 611–622, September 1999. ISSN 1369-7412. doi: 10.1111/1467-9868.00196. URL <https://doi.org/10.1111/1467-9868.00196>.
- John Vastola. Optimal packing of attractor states in neural representations. In Sophia Sanborn, Christian Shewmake, Simone Azeglio, and Nina Miolane, editors, *Proceedings of the 2nd NeurIPS Workshop on Symmetry and Geometry in Neural Representations*, volume 228 of *Proceedings of Machine Learning Research*, pages 425–442. PMLR, 16 Dec 2024. URL <https://proceedings.mlr.press/v228/vastola24a.html>.

Proceedings Track

- John Vastola. Generalization through variance: how noise shapes inductive biases in diffusion models. In *The Thirteenth International Conference on Learning Representations*, 2025a. URL <https://openreview.net/forum?id=7lUdo8Vuqa>.
- John J. Vastola. Dynamical symmetries in the fluctuation-driven regime: an application of Noether’s theorem to noisy dynamical systems. In *Proceedings of the 3rd NeurIPS Workshop on Symmetry and Geometry in Neural Representations*, 2025b. URL <https://openreview.net/forum?id=1LiIJc7oCJ>.
- John J. Vastola, Zach Cohen, and Jan Drugowitsch. Is the information geometry of probabilistic population codes learnable? In Sophia Sanborn, Christian Shewmake, Simone Azeglio, Arianna Di Bernardo, and Nina Miolane, editors, *Proceedings of the 1st NeurIPS Workshop on Symmetry and Geometry in Neural Representations*, volume 197 of *Proceedings of Machine Learning Research*, pages 258–277. PMLR, 03 Dec 2023. URL <https://proceedings.mlr.press/v197/vastola23a.html>.
- Maryna S. Viazovska. The sphere packing problem in dimension 8. *Annals of Mathematics*, 185(3):991–1015, 2017. ISSN 0003486X. URL <http://www.jstor.org/stable/26395747>.
- Binxu Wang and John Vastola. The unreasonable effectiveness of gaussian score approximation for diffusion models and its applications. *Transactions on Machine Learning Research*, 2024. ISSN 2835-8856. URL <https://openreview.net/forum?id=IOuknSHM2j>.
- Martin Wattenberg, Fernanda Viégas, and Ian Johnson. How to use t-sne effectively. *Distill*, 2016. doi: 10.23915/distill.00002. URL <http://distill.pub/2016/misread-tsne>.
- Alison I. Weber and Jonathan W. Pillow. Capturing the Dynamical Repertoire of Single Neurons with Generalized Linear Models. *Neural Computation*, 29(12):3260–3289, December 2017. ISSN 0899-7667. doi: 10.1162/neco_a.01021. URL https://doi.org/10.1162/neco_a.01021. eprint: https://direct.mit.edu/neco/article-pdf/29/12/3260/1026921/neco_a.01021.pdf.

Proceedings Track

Appendix A. Motivating the objective function

In this appendix, we present an alternative Bayesian argument to motivate our objective function

$$J[\phi] = \int_{\mathbb{R}^d} d\mathbf{z} \, q(\mathbf{z}) \left[\frac{1}{2} \log \det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})) + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right], \quad (28)$$

and also briefly comment on technical subtleties related to this objective. As described in Sec. 2, D is the ambient space dimensionality, $d \leq D$ is the latent space dimensionality, p_{data} is a likelihood on $\mathbb{R}^D$, q is a distribution (the ‘prior’) on $\mathbb{R}^d$, and $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$ is a smooth injection chosen to maximize J , subject to certain finiteness conditions.

A.1. A Bayesian argument for the objective function

Assumed ground truth distribution. Let $\mathbf{x} \in \mathbb{R}^D$ denote an observation. Assume that observations genuinely have low-dimensional structure, in the sense that they can be viewed as generally lying on or near some d -dimensional manifold embedded in $\mathbb{R}^D$ (with $d \leq D$).

We can make the idea of a low-dimensional manifold more concrete by imagining that there exists some smooth injection $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$ which embeds any particular location $\mathbf{z} \in \mathbb{R}^d$ on this manifold (given in coordinates that may differ from those of the ambient space) into $\mathbb{R}^D$. We can make our demand that observations generally lie ‘near’ the manifold more concrete by assuming a noise model $p_{\text{noise}}(\mathbf{x}|\mathbf{z}, \phi) := \mathcal{N}(\mathbf{x}; \phi(\mathbf{z}), \epsilon \mathbf{I}_D)$, where $\epsilon > 0$ is small. Note that this noise model allows observations to be off-manifold, but that this only occurs with a low probability.

If we assume that latent states are sampled according to a distribution $q(\mathbf{z})$, then the joint distribution of observations and latent states is

$$p(\mathbf{x}, \mathbf{z}) := p_{\text{noise}}(\mathbf{x}|\mathbf{z}, \phi)q(\mathbf{z}) = \mathcal{N}(\mathbf{x}; \phi(\mathbf{z}), \epsilon \mathbf{I}_D)q(\mathbf{z}). \quad (29)$$

Assume that q and $\epsilon > 0$ are known, but the mapping ϕ is not. How can we infer it?

Assumed model. One sensible way to proceed is to exploit Bayes’ rule. If we model observations as distributed according to $p_{\text{data}}(\mathbf{x})$, our knowledge of the noise model $p_{\text{noise}}(\mathbf{x}|\mathbf{z}, \phi)$ allows us to infer latent states from observations via

$$p(\mathbf{z}|\phi) := \int p(\mathbf{z}|\mathbf{x}, \phi)p_{\text{data}}(\mathbf{x}) \, d\mathbf{x} = \int \frac{p_{\text{noise}}(\mathbf{x}|\mathbf{z}, \phi)}{\int p_{\text{noise}}(\mathbf{x}|\mathbf{z}', \phi) d\mathbf{z}'} p_{\text{data}}(\mathbf{x}) \, d\mathbf{x}. \quad (30)$$

Let $p(\phi)$ denote a prior over mappings that encodes our preferences about it (e.g., that it is continuous). By Bayes’ rule, the posterior distribution of ϕ given a latent state $\mathbf{z}$ is

$$p(\phi|\mathbf{z}) = \frac{p(\mathbf{z}|\phi)p(\phi)}{\int p(\mathbf{z}|\phi')p(\phi') \, d\phi'}, \quad (31)$$

where we note that, since the distributions over ϕ are defined over *function spaces*, there are technical subtleties related to normalization. Fortunately, these difficulties are irrelevant from the point of view of identifying the most likely mapping ϕ .

Proceedings Track

Using Bayes' rule to estimate the embedding map. Let $\{\mathbf{z}_i\}_{i=1}^N$ be $N \geq 1$ independent and identically distributed samples from q , the ground truth distribution of latent states. We can estimate ϕ by maximizing the scaled log-posterior probability

$$\begin{aligned} \frac{1}{N} \log p(\phi | \{\mathbf{z}_i\}) &= \frac{1}{N} \log p(\{\mathbf{z}_i\} | \phi) + \frac{1}{N} \log p(\phi) + \text{const.} \\ &= \frac{1}{N} \sum_{i=1}^N \log p(\mathbf{z}_i | \phi) + \frac{1}{N} \log p(\phi) + \text{const.} \end{aligned} \quad (32)$$

where ‘const.’ denotes a ϕ -independent term that comes from the normalization of the posterior. In the large N limit, the central limit theorem implies that J' approaches its average value, which is

$$J'[\phi] = \int_{\mathbb{R}^d} d\mathbf{z} \, q(\mathbf{z}) \log p(\mathbf{z} | \phi) + \frac{1}{N} \log p(\phi) , \quad (33)$$

where we have dropped the constant term since it does not affect our optimization. Our argument so far indicates that maximizing J' with respect to ϕ yields a Bayes-optimal estimate of it. The final remaining step, after which we will recover an objective equivalent to the one we study in the main text, is to simplify J' .

Taking the small-noise limit. We can simplify J' by using Laplace's method to approximate $p(\mathbf{z} | \phi)$ in the small noise ($\epsilon \rightarrow 0$) limit. The probability $p_{\text{noise}}(\mathbf{x} | \mathbf{z})$ is obviously maximized when $\mathbf{x} = \phi(\mathbf{z})$; similarly, if $\mathbf{x} \in \phi(\mathbf{R}^d)$, then $p(\mathbf{z} | \mathbf{x})$ is maximized at the unique $\mathbf{z}$ for which $\phi(\mathbf{z}) = \mathbf{x}$ (this is unique since ϕ is an injection).

Denote this unique point by $\mathbf{z}_*(\mathbf{x})$. At a small displacement $\Delta\mathbf{z}$ away from it,

$$\frac{1}{2} \|\mathbf{x} - \phi(\mathbf{z}_*(\mathbf{x}) + \Delta\mathbf{z})\|_2^2 \approx \frac{1}{2} (\Delta\mathbf{z})^T \mathbf{J}(\mathbf{z}_*(\mathbf{x}))^T \mathbf{J}(\mathbf{z}_*(\mathbf{x})) (\Delta\mathbf{z}) , \quad (34)$$

i.e., the Hessian near $\mathbf{z}_*(\mathbf{x})$ is $\mathbf{J}(\mathbf{z}_*(\mathbf{x}))^T \mathbf{J}(\mathbf{z}_*(\mathbf{x}))$. By Laplace's method, this implies

$$\int e^{-\frac{1}{2\epsilon} \|\mathbf{x} - \phi(\mathbf{z}')\|_2^2} d\mathbf{z}' \xrightarrow{\epsilon \rightarrow 0} \frac{(2\pi\epsilon)^{d/2}}{\sqrt{\det[\mathbf{J}(\mathbf{z}_*(\mathbf{x}))^T \mathbf{J}(\mathbf{z}_*(\mathbf{x}))]}} . \quad (35)$$

Furthermore,

$$\begin{aligned} p(\mathbf{z} | \phi) &\approx \frac{1}{(2\pi\epsilon)^{d/2}} \int e^{-\frac{1}{2\epsilon} \|\mathbf{x} - \phi(\mathbf{z})\|_2^2} p_{\text{data}}(\mathbf{x}) \sqrt{\det[\mathbf{J}(\mathbf{z}_*(\mathbf{x}))^T \mathbf{J}(\mathbf{z}_*(\mathbf{x}))]} d\mathbf{x} \\ &\xrightarrow{\epsilon \rightarrow 0} (2\pi\epsilon)^{(D-d)/2} p_{\text{data}}(\phi(\mathbf{z})) \sqrt{\det[\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})]} \end{aligned} \quad (36)$$

in the $\epsilon \rightarrow 0$ limit. Up to irrelevant constants, this gives us the result we want: maximizing the average log-posterior probability of ϕ yields our objective, possibly plus an extra $\log p(\phi)$ term if the prior over ϕ is non-uniform. More explicitly,

$$J'[\phi] \xrightarrow{\epsilon \rightarrow 0} \int_{\mathbb{R}^d} d\mathbf{z} \, q(\mathbf{z}) \left\{ \frac{1}{2} \log \det[\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})] + \log p_{\text{data}}(\phi(\mathbf{z})) \right\} + \frac{1}{N} \log p(\phi) + \frac{(D-d)}{2} \log(2\pi\epsilon) , \quad (37)$$

which only differs from Eq. 28 by a constant offset when $p(\phi)$ is uniform (or when we eliminate the prior term by taking $N \rightarrow \infty$).

Proceedings Track

A.2. Additional comments about the objective

The embedding map needs to be an injection. The latent states $\mathbf{z}$ index locations on the d -dimensional manifold embedded in the ambient space, and hence function as coordinates for it. This intuition suggests that ϕ ought to be an injection; any given point on the manifold should only correspond to a single value of $\mathbf{z}$.

There is a more important technical reason for us to demand that ϕ is an injection. If ϕ is an injection, the pullback of p_{data} through ϕ is not generally normalized, i.e.,

$$\int_{\mathbb{R}^d} p_{\text{data}}(\phi(\mathbf{z})) \sqrt{\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})} d\mathbf{z} \leq 1 . \quad (38)$$

But because it is less than or equal to one, we can still invoke Gibbs' inequality to argue

$$J[\phi] = \int_{\mathbb{R}^d} q(\mathbf{z}) \log \left[\frac{(\phi^* p_{\text{data}})(\mathbf{z})}{q(\mathbf{z})} \right] d\mathbf{z} \leq 0 . \quad (39)$$

(The argument is the same as the standard one used to prove that the KL divergence is nonnegative, so we omit it.) If ϕ were *not* an injection, J does not necessarily have an upper bound, and hence our optimization problem is not well-defined. For example, suppose

$$\begin{aligned} q(\mathbf{z}) &:= 1 \quad \text{for } \mathbf{z} \in [0, 1] \quad \text{and} \\ p_{\text{data}}(x, y) &:= \frac{1}{2\pi} \delta(x^2 + y^2 - 1) \quad \text{for } (x, y) \in \mathbb{R}^2, \end{aligned} \quad (40)$$

i.e., q is uniform on $[0, 1]$ and p_{data} is uniform on the unit circle. Each embedding

$$\phi_1(\mathbf{z}; k) := \cos(2\pi k \mathbf{z}) \quad \phi_2(\mathbf{z}; k) := \sin(2\pi k \mathbf{z}) \quad (41)$$

for $k \in \mathbb{N}$ yields the same $\log p_{\text{data}}(\phi(\mathbf{z}))$ term, but have different Jacobian terms. Note,

$$\frac{1}{2} \log \|\phi'(\mathbf{z})\|_2^2 = \frac{1}{2} \log [(2\pi k)^2] = \log(2\pi k) , \quad (42)$$

which implies that we can make this term arbitrarily large by choosing k to be larger and larger. This type of construction is only possible if we allow each (x, y) on the unit circle to correspond to multiple values of $\mathbf{z}$; such constructions are not possible if we insist that ϕ is an injection, as Gibbs' inequality indicates.

Boundary conditions. We are looking for smooth injections $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$ that maximize Eq. 28 while also satisfying the finiteness conditions

$$\mathbb{E}_{\mathbf{z} \sim q} \left\{ \sqrt{\det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z}))} \right\} < \infty \quad \text{and} \quad \mathbb{E}_{\mathbf{z} \sim q} \{ \|\phi(\mathbf{z})\|_2^2 \} < \infty . \quad (43)$$

Since a small region of $\mathbb{R}^d$ becomes a small region of $\mathbb{R}^D$ with surface area

$$dV = \sqrt{\det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z}))} d\mathbf{z} \quad (44)$$

when projected through ϕ , the first condition is equivalent to asking that the average surface area (or the total q -weighted ‘mass’ of the embedding) is finite. This condition prevents ‘needle’-like embeddings which project small regions of $\mathbb{R}^d$ to extremely large regions of $\mathbb{R}^D$. It also plays a crucial role in our argument that the PCA objective does not admit nonlinear solutions (see Appendix F).

The second condition prevents pathological embeddings in the tails of q (i.e., large $\|\mathbf{z}\|_2$) by controlling their norm.

Proceedings Track

Appendix B. Derivation of Euler-Lagrange equations

In this appendix, we derive the Euler-Lagrange (EL) equations for our objective

$$J[\phi] = \int_{\mathbb{R}^d} dz \, q(\mathbf{z}) \left[\frac{1}{2} \log \det(\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})) + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right], \quad (45)$$

where $J_{ij} := \partial \phi_i / \partial z_j$ are the elements of the $D \times d$ Jacobian matrix $\mathbf{J}$. We will assume that $\mathbf{J}$ has full column rank, so that the determinant that appears is nonzero. The Lagrangian density corresponding to this objective is

$$\mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) = q(\mathbf{z}) \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right]. \quad (46)$$

B.1. Euler-Lagrange equations for arbitrary-dimensional embeddings

The EL equations are

$$\sum_{j=1}^d \partial_j \left(\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \phi_i}, \quad (47)$$

so we only need to evaluate these derivatives explicitly and simplify the result. First, note

$$\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} = \frac{\partial \mathcal{L}}{\partial J_{ij}} = q(\mathbf{z}) \frac{\partial}{\partial J_{ij}} \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) \right] = q(\mathbf{z}) [\mathbf{J}(\mathbf{J}^T \mathbf{J})^{-1}]_{ij}, \quad (48)$$

where we have used a matrix cookbook result on the derivative of a log-determinant (Petersen and Pedersen, 2008). Note that $\mathbf{J}^+ := (\mathbf{J}^T \mathbf{J})^{-1} \mathbf{J}^T$ is the Moore-Penrose pseudoinverse of $\mathbf{J}$, i.e., the $d \times D$ matrix which satisfies $\mathbf{J}^+ \mathbf{J} = \mathbf{I}_d$. Using it allows us to write the above partial derivatives in vector form as

$$\frac{\partial \mathcal{L}}{\partial \mathbf{J}} = (\mathbf{J}^+)^T q(\mathbf{z}). \quad (49)$$

More straightforwardly, the partial derivative of $\mathcal{L}$ with respect to ϕ_i is the score of p_{data} :

$$\frac{\partial \mathcal{L}}{\partial \phi_i} = q(\mathbf{z}) \frac{\partial \log p_{\text{data}}(\phi(\mathbf{z}))}{\partial \phi_i}. \quad (50)$$

These results immediately imply that the EL equation for ϕ_i reads

$$\sum_j \partial_j \left(\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \phi_i} \implies \sum_j \partial_j [J_{ji}^+ q(\mathbf{z})] = q(\mathbf{z}) \frac{\partial \log p_{\text{data}}(\phi(\mathbf{z}))}{\partial \phi_i}. \quad (51)$$

Simplifying, we have

$$\sum_j \partial_j (J_{ji}^+) + \sum_j J_{ji}^+ \frac{\partial_j q}{q} = \frac{\partial \log p_{\text{data}}(\phi(\mathbf{z}))}{\partial \phi_i}, \quad (52)$$

or in vector form,

$$[\partial + \nabla \log q(\mathbf{z})]^T \mathbf{J}^+ = \mathbf{s}_\phi(\mathbf{z})^T, \quad (53)$$

where we have defined the score vector $\mathbf{s}_i(\mathbf{z}) := \partial \log p_{\text{data}}(\phi(\mathbf{z})) / \partial \phi_i$.

Proceedings Track

B.2. Euler-Lagrange equations for one-dimensional embeddings

The one-dimensional ($d = 1$) form of the objective reads

$$J[\phi] = \int_{-\infty}^{\infty} dz \, q(z) \left[\frac{1}{2} \log \|\phi'(z)\|_2^2 + \log p_{\text{data}}(\phi(z)) - \log q(z) \right] \quad (54)$$

and the corresponding Lagrangian is

$$L(\phi, \phi', z) := q(z) \left[\frac{1}{2} \log \|\phi'(z)\|_2^2 + \log p_{\text{data}}(\phi(z)) - \log q(z) \right] . \quad (55)$$

In this setting, the EL equations have the somewhat simpler form

$$\frac{d}{dz} \left(\frac{\partial L}{\partial \phi'_i} \right) = \frac{\partial L}{\partial \phi_i} \quad , \text{ i.e., } \quad \frac{d}{dz} \left[q(z) \frac{\phi'_i(z)}{\|\phi'(z)\|_2^2} \right] = q(z) \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z)) . \quad (56)$$

Simplifying as in the previous subsection, we end up with

$$\frac{d}{dz} \left[\frac{\phi'_i(z)}{\|\phi'(z)\|_2^2} \right] + \frac{q'(z)}{q(z)} \frac{\phi'_i(z)}{\|\phi'(z)\|_2^2} = \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z)) , \quad (57)$$

or if we prefer to be even more explicit,

$$\frac{\phi''_i(z)}{\|\phi'(z)\|_2^2} - \frac{2\phi'_i(z) \sum_k \phi'_k(z)}{\|\phi'(z)\|_2^4} + \frac{q'(z)}{q(z)} \frac{\phi'_i(z)}{\|\phi'(z)\|_2^2} = \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z)) . \quad (58)$$

Proceedings Track

Appendix C. Optimal 1D embeddings and diffusion models

In Sec. 3, we showed that optimal 1D embeddings $\phi : \mathbb{R} \rightarrow \mathbb{R}^D$ satisfy a system

$$\phi'_i(z) = \frac{p_i(z)}{\|\mathbf{p}(z)\|_2^2} \quad p'_i(z) = s_i(\phi(z)) - \left[\frac{d}{dz} \log q(z) \right] p_i(z) \quad (59)$$

of first-order ODEs, where the $p_i(z)$ are auxiliary momentum-like variables and $s_i(\phi(z)) := \frac{\partial}{\partial \phi_i} \log p_{\text{data}}(\phi(z))$ are the components of the *score vector* associated with p_{data} at the point $\phi(z) \in \mathbb{R}^D$. In this appendix, we argue that optimal embeddings exactly follow score vectors in a certain limit.

Consider a prior $q(z) = \gamma e^{\gamma z}$ which is nonzero only for $z \in (-\infty, 0]$. For this q ,

$$\phi'_i(z) = \frac{p_i(z)}{\|\mathbf{p}(z)\|_2^2} \quad p'_i(z) = s_i(\phi(z)) - \gamma p_i(z) . \quad (60)$$

Note that $\gamma > 0$ controls the ‘time scale’ on which $p_i(z)$ approaches its steady state value, at which $p_i(z) = s_i(\phi(z))/\gamma$. More precisely, Eq. 60 implies that $p_i(z)$ obeys the self-consistent equation

$$p_i(z) = \int_{-\infty}^z s_i(\phi(z')) e^{-\gamma(z-z')} dz' . \quad (61)$$

Since the exponential decay factor becomes a delta function in the large γ limit, i.e.,

$$\gamma e^{-\gamma(z-z')} \xrightarrow{\gamma \rightarrow \infty} \delta(z - z') , \quad (62)$$

in this limit we have

$$p_i(z) \approx \frac{1}{\gamma} \int_{-\infty}^z s_i(\phi(z')) \delta(z - z') dz' = \frac{1}{\gamma} s_i(\phi(z)) , \quad (63)$$

which means that the momentum vectors match the score vectors. Eq. 60 then implies

$$\phi'_i(z) \approx \gamma \frac{s_i(\phi(z))}{\|\mathbf{s}(z)\|_2^2} . \quad (64)$$

As we noted in Sec. 3, this implies that trajectories follow score vectors, albeit with a speed that changes depending on their norm. If we perform a reparameterization of ‘time’ with

$$\tau(z) := \int_{-\infty}^z \frac{1}{\|\mathbf{s}(z')\|_2^2} dz' \quad \tau'(z) = \frac{1}{\|\mathbf{s}(z)\|_2^2} , \quad (65)$$

then we have the equivalent dynamics

$$\phi'_i(\tau) = s_i(\phi(\tau)) \quad (66)$$

which exactly follows the score field.

Proceedings Track

Appendix D. The objective is reparameterization invariant

Recall that our objective is

$$J[\phi] = \int_{\mathbb{R}^d} dz \, q(\mathbf{z}) \left[\frac{1}{2} \log \det[\mathbf{J}(\mathbf{z})^T \mathbf{J}(\mathbf{z})] + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right]. \quad (67)$$

One of its key properties is that it is *invariant to diffeomorphisms*; said differently, it is not sensitive to our choice of latent space coordinates. This property is desirable, since one's choice of coordinates is arbitrary, and different coordinate systems can reflect the same underlying geometry. In this appendix, we show that this property holds via a direct computation.

Let $\mathbf{g} : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a diffeomorphism, and consider a change of variables $\mathbf{z} = \mathbf{g}(\mathbf{z}')$ from $\mathbf{z}$ to $\mathbf{z}'$. Note that

$$\phi(\mathbf{z}) = \phi(\mathbf{g}(\mathbf{z}')) , \quad (68)$$

so in the new coordinate system ϕ becomes $\tilde{\phi}(\mathbf{z}') := \phi(\mathbf{g}(\mathbf{z}'))$. The Jacobian of $\tilde{\phi}$ in this new coordinate system is

$$\tilde{J}_{ij} := \frac{\partial \tilde{\phi}_i}{\partial z'_j} = \sum_{k=1}^d \frac{\partial \tilde{\phi}_i}{\partial z_k} \frac{\partial z_k}{\partial z'_j} = \sum_{k=1}^d \frac{\partial \phi_i}{\partial z_k} \frac{\partial z_k}{\partial z'_j} , \quad (69)$$

which means that $\tilde{\mathbf{J}} = \mathbf{J} \mathbf{J}_g$, i.e., that the new Jacobian equals the old one times the Jacobian of $\mathbf{g}$. Finally, the measure $q(\mathbf{z})d\mathbf{z}$ becomes

$$q(\mathbf{g}(\mathbf{z}')) |\det \mathbf{J}_g| d\mathbf{z}' , \quad (70)$$

so in the new coordinate system we effectively go from $q(\mathbf{z})$ to $\tilde{q}(\mathbf{z}') := q(\mathbf{g}(\mathbf{z}')) |\det \mathbf{J}_g|$. Putting these results together, we find that $J[\phi]$ equals

$$\begin{aligned} J[\phi] &= \int_{\mathbb{R}^d} d\mathbf{z}' \, q(\mathbf{g}(\mathbf{z}')) |\det \mathbf{J}_g| \left[\frac{1}{2} \log \det[\mathbf{J}(\mathbf{g}(\mathbf{z}'))^T \mathbf{J}(\mathbf{g}(\mathbf{z}'))] + \log p_{\text{data}}(\phi(\mathbf{g}(\mathbf{z}'))) - \log q(\mathbf{g}(\mathbf{z}')) \right] \\ &= \int_{\mathbb{R}^d} d\mathbf{z}' \, \tilde{q}(\mathbf{z}') \left[\frac{1}{2} \log \det[\mathbf{J}(\mathbf{g}(\mathbf{z}'))^T \mathbf{J}(\mathbf{g}(\mathbf{z}'))] + \log p_{\text{data}}(\tilde{\phi}(\mathbf{z}')) - \log q(\mathbf{g}(\mathbf{z}')) \right] \\ &= \int_{\mathbb{R}^d} d\mathbf{z}' \, \tilde{q}(\mathbf{z}') \left[\frac{1}{2} \log \det[\tilde{\mathbf{J}}(\mathbf{z}')^T \tilde{\mathbf{J}}(\mathbf{z}')] + \log p_{\text{data}}(\tilde{\phi}(\mathbf{z}')) - \log \tilde{q}(\mathbf{z}') \right] , \end{aligned} \quad (71)$$

where two $\log \det |\mathbf{J}_g|$ terms cancelled out in the last step. Since the final line is just the objective in the new coordinate system determined by the diffeomorphism $\mathbf{g}$, J is indeed reparameterization invariant.

Note that this property also follows straightforwardly from the fact that J equals a (negative) KL divergence.

Appendix E. Review and consequences of Noether's theorem

Noether's theorem, which we first discuss in Sec. 4 and then explore in more detail in Sec. 5, links continuous quasi-symmetries of the objective function to conservation laws (Noether, 1918; Neuenschwander, 2017; Vastola, 2025b). In this appendix, we review its statement and provide a derivation of embedding energy conservation.

E.1. Statement of Noether's theorem

Let $\phi := (\phi_1, \dots, \phi_D)^T$ be a vector of D fields defined on $\mathbb{R}^d$, i.e., each field is a smooth function from $\mathbb{R}^d$ to $\mathbb{R}$. Let

$$J[\phi] := \int_{\mathbb{R}^d} dz \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) \quad (72)$$

be an objective function that depends on such a vector of fields through a Lagrangian density $\mathcal{L}(\phi, \mathbf{J}, \mathbf{z})$, where the $D \times d$ Jacobian matrix $\mathbf{J}$ has components $J_{ij} = \partial\phi_i/\partial z_j$.

We will call a field *extremal* if it makes J stationary, which happens if and only if it satisfies the EL equations

$$\sum_{j=1}^d \partial_j \left(\frac{\partial \mathcal{L}}{\partial (\partial_j \phi_i)} \right) = \frac{\partial \mathcal{L}}{\partial \phi_i} . \quad (73)$$

Next, define the *stress-energy tensor*

$$T_{ij}(\mathbf{z}) := \sum_{k=1}^D \frac{\partial \mathcal{L}}{\partial (\partial_i \phi_k)} (\partial_j \phi_k(\mathbf{z})) - \delta_{ij} \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) . \quad (74)$$

Following Neuenschwander (2017), Noether's theorem is the following result.

Theorem 1 (Noether's theorem) *Consider a transformation from coordinates $\mathbf{z}$ and fields ϕ to*

$$\begin{aligned} \mathbf{z}' &= \mathbf{z} + \epsilon(\delta\mathbf{z})(\mathbf{z}, \phi) \\ \phi' &= \phi + \epsilon(\delta\phi)(\mathbf{z}, \phi) \end{aligned} \quad (75)$$

where $\epsilon > 0$ is infinitesimally small. Note that the perturbations $\delta\mathbf{z}$ and $\delta\phi$ are allowed to depend on $\mathbf{z}$ and ϕ . If this change of coordinates and fields changes the Lagrangian density by a total derivative, i.e.,

$$\mathcal{L}(\phi', \mathbf{J}', \mathbf{z}') = \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) + \epsilon \nabla \cdot \mathbf{K}(\phi, \mathbf{J}, \mathbf{z}) + \mathcal{O}(\epsilon^2) \quad (76)$$

for some function $\mathbf{K}$, then we call such a transformation a (continuous) quasi-symmetry of the objective, and conclude that the **Noether current**

$$j_a := \sum_{b=1}^D \frac{\partial \mathcal{L}}{\partial (\partial_a \phi_b)} (\delta\phi_b) - \sum_{c=1}^d T_{ac}(\delta z_c) - K_a \quad (77)$$

is conserved (i.e., $\nabla \cdot \mathbf{j} = 0$) whenever ϕ satisfies the EL equations.

Proceedings Track

For our *particular* Lagrangian density

$$\mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) = q(\mathbf{z}) \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) + \log p_{\text{data}}(\phi(\mathbf{z})) - \log q(\mathbf{z}) \right], \quad (78)$$

since (see Appendix B)

$$\frac{\partial \mathcal{L}}{\partial J_{ij}} = q(\mathbf{z}) J_{ji}^+, \quad (79)$$

we have

$$T_{ij}(\mathbf{z}) := \sum_k q(\mathbf{z}) J_{ik}^+ J_{kj} - \delta_{ij} \mathcal{L}(\phi, \mathbf{J}, \mathbf{z}) = \delta_{ij} [q(\mathbf{z}) - \mathcal{L}(\phi, \mathbf{J}, \mathbf{z})] \quad (80)$$

and

$$\begin{aligned} j_a &= \sum_{b=1}^D q(\mathbf{z}) J_{ab}^+ (\delta \phi_b) - \sum_{c=1}^d \delta_{ac} (q - \mathcal{L}) (\delta z_b) - K_a \\ &= \sum_{b=1}^D q(\mathbf{z}) J_{ab}^+ (\delta \phi_b) - (q(\mathbf{z}) - \mathcal{L}) (\delta z_a) - K_a. \end{aligned} \quad (81)$$

E.2. Direct derivation of embedding energy conservation

We can derive embedding energy conservation through a direct computation, without going through Noether's theorem. We want to show that the quantity

$$\boxed{\mathcal{E} := -\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) - \log p_{\text{data}}(\phi(\mathbf{z})) + \log q(\mathbf{z})}, \quad (82)$$

which we call the **embedding energy**, is constant (i.e., does not depend on $\mathbf{z}$) for all solutions of the EL equations. To see this, multiply Eq. 53 on the right by $\mathbf{J}$:

$$[\partial + \nabla \log q(\mathbf{z})]^T \mathbf{J}^+ \mathbf{J} = \mathbf{s}_\phi(\mathbf{z})^T \mathbf{J}. \quad (83)$$

Since $\mathbf{J}^+ \mathbf{J} = \mathbf{I}$, one term is easy to simplify:

$$[\nabla \log q(\mathbf{z})]^T \mathbf{J}^+ \mathbf{J} = [\nabla \log q(\mathbf{z})]^T. \quad (84)$$

The derivative term is somewhat more delicate. In terms of components, we have

$$\sum_{j,i} \partial_j (J_{ji}^+) J_{ik} = \sum_{j,i} \partial_j (J_{ji}^+ J_{ik}) - J_{ji}^+ \partial_j (J_{ik}) = - \sum_{j,i} J_{ji}^+ \partial_j (J_{ik}). \quad (85)$$

But note that, since $(\mathbf{J}^+)^T$ is the derivative of the kinetic term with respect to $\mathbf{J}$, we have

$$\partial_k \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) \right] = \sum_{i,j} \frac{\partial}{\partial J_{ij}} \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) \right] \partial_k (J_{ij}) = \sum_{i,j} J_{ji}^+ \partial_k (J_{ij}) = \sum_{i,j} J_{ji}^+ \partial_j (J_{ik}),$$

where in the last step we used the fact that $\partial_k (J_{ij}) = \partial_{kj}^2 \phi_i = \partial_{jk}^2 \phi_i = \partial_j (J_{ik})$. Hence, we conclude that

$$\nabla \left[\frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) \right] + \nabla \log q(\mathbf{z}) = \mathbf{s}_\phi(\mathbf{z})^T \mathbf{J} = \nabla \log p_{\text{data}}(\phi(\mathbf{z})), \quad (86)$$

Proceedings Track

or equivalently that

$$\nabla \left\{ \frac{1}{2} \log \det(\mathbf{J}^T \mathbf{J}) + \log p_{\text{data}}(\boldsymbol{\phi}(\mathbf{z})) - \log q(\mathbf{z}) \right\} = 0 . \quad (87)$$

This implies that the bracketed quantity is constant, i.e., that $\mathcal{E}$ is conserved for solutions of the EL equations.

As an aside, note that

$$0 \geq J[\phi] = - \int_{\mathbb{R}^d} d\mathbf{z} \, q(\mathbf{z}) \, \mathcal{E} = -\mathcal{E} . \quad (88)$$

This implies that we ought to minimize $\mathcal{E}$, and that if J is perfectly maximized (such that the corresponding KL divergence is minimized), $\mathcal{E} = 0$.

Proceedings Track

Appendix F. There are no nonlinear solutions to the PCA objective

In the main text (Sec. 6), we argue that linear maps satisfy the EL equations associated with the PCA objective, and find that the linear maps that maximize J pick out the top eigenvectors of the data's covariance matrix. Here, we argue that it is impossible for these EL equations to have nonlinear solutions, and hence for J to have any nonlinear global maximizers. This point is important, since otherwise we have not established that linear maps are the global maximizers of the PCA objective.

For technical reasons, we specialize to the prior with $\sigma_0 = 1$ below, but this assumption can easily be relaxed.

F.1. The one-dimensional case

The objective is

$$J[\phi] := \int_{-\infty}^{\infty} dz \frac{e^{-\frac{1}{2}z^2}}{\sqrt{2\pi}} \left\{ \frac{1}{2} \log \|\phi'(z)\|_2^2 - \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) \right\} + \text{const.} \quad (89)$$

where we are optimizing over smooth injections $\phi : \mathbb{R} \rightarrow \mathbb{R}^D$. Recall from Sec. 2 that we also insist on two finiteness conditions, which here read

$$\mathbb{E}_{z \sim q} \{ \|\phi'(z)\|_2 \} < \infty \quad \text{and} \quad \mathbb{E}_{z \sim q} \{ \|\phi\|_2^2 \} < \infty. \quad (90)$$

Note that the second condition implies that

$$\mathbb{E}_{z \sim q} \{ \phi_i(z)^2 \} < \infty \quad (91)$$

for each $i = 1, \dots, D$. This, in turn, implies that each $\phi_i(z)$ can be written in terms of a (probabilist's) Hermite polynomial expansion

$$\phi_i(z) = \sum_{n=0}^{\infty} c_{ni} H_n(z) = c_{0i} + c_{1i}z + c_{2i}(z^2 - 1) + \dots \quad (92)$$

By the conservation of embedding energy (see Sec. 5), we have

$$-\frac{1}{2} \|\phi'(z)\|_2^2 + \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) - \frac{z^2}{2} = \text{const.} \quad (93)$$

which implies that

$$\begin{aligned} \|\phi'(z)\|_2 &= C \exp \left\{ \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) - \frac{z^2}{2} \right\} \\ &= C \exp \left\{ \frac{(\mathbf{c}_0^T \Sigma^{-1} \mathbf{c}_0)}{2} + (\mathbf{c}_0^T \Sigma^{-1} \mathbf{c}_1)z + \frac{(\mathbf{c}_1^T \Sigma^{-1} \mathbf{c}_1)}{2} z^2 + \mathcal{O}(z^3) - \frac{z^2}{2} \right\} \end{aligned} \quad (94)$$

for some constant $C > 0$, where $\mathbf{c}_n$ denotes the vector whose i -th component is c_{ni} . Unless the higher-order coefficients are zero, this norm will grow at least as quickly as $\exp(z^n)$ for some $n > 2$, and hence the first finiteness condition will be violated.

Proceedings Track

F.2. The general case

In the general case,

$$J[\phi] := \int_{\mathbb{R}^d} dz \frac{e^{-\frac{1}{2}\|z\|_2^2}}{(2\pi)^{d/2}} \left\{ \frac{1}{2} \log \det[\mathbf{J}(z)^T \mathbf{J}(z)] - \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) \right\} + \text{const.} \quad (95)$$

where we are optimizing over smooth injections $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^D$. The general finiteness conditions are

$$\mathbb{E}_{z \sim q} \left\{ \sqrt{\det(\mathbf{J}(z)^T \mathbf{J}(z))} \right\} < \infty \quad \text{and} \quad \mathbb{E}_{z \sim q} \left\{ \|\phi(z)\|_2^2 \right\} < \infty. \quad (96)$$

As above, the second finiteness condition implies that each $\phi_i(z)$ can be written in terms of a Hermite polynomial expansion. By the conservation of embedding energy,

$$-\log \sqrt{\det(\mathbf{J}(z)^T \mathbf{J}(z))} + \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) - \frac{\|z\|_2^2}{2} = \text{const.} \quad (97)$$

which implies

$$\sqrt{\det(\mathbf{J}(z)^T \mathbf{J}(z))} = C \exp \left\{ \frac{1}{2} \phi(z)^T \Sigma^{-1} \phi(z) - \frac{\|z\|_2^2}{2} \right\} \quad (98)$$

for some constant $C > 0$. By the same argument as before, the argument of this exponential will grow at least as quickly as a polynomial in z , and hence not satisfy the first finiteness condition, unless the coefficients of the higher-order Hermite expansion terms are zero. Since only the zeroth and first order terms of the expansion are permitted, we conclude that the global maximizers of J must be linear.