
Shapes analysis for time series

Thibaut Germain^{1*}

Centre Borelli, ENS Paris-Saclay
4 av. des sciences, 91190

Samuel Gruffaz¹

Centre Borelli, ENS Paris-Saclay
4 av. des sciences, 91190

Charles Truong¹

Centre Borelli, ENS Paris-Saclay
4 av. des sciences, 91190

Laurent Oudre¹

Centre Borelli, ENS Paris-Saclay
4 av. des sciences, 91190

Alain Durmus

CMAP, CNRS, Ecole polytechnique
Institut Polytechnique de Paris
91120 Palaiseau, France

Abstract

Analyzing inter-individual variability of physiological functions is particularly appealing in medical and biological contexts to describe or quantify health conditions. Such analysis can be done by comparing individuals to a reference one with time series as biomedical data. This paper introduces an unsupervised representation learning (URL) algorithm for time series tailored to inter-individual studies. The idea is to represent time series as deformations of a reference time series. The deformations are diffeomorphisms parameterized and learned by our method called TS-LDDMM. Once the deformations and the reference time series are learned, the vector representations of individual time series are given by the parametrization of their corresponding deformation. At the crossroads between URL for time series and shape analysis, the proposed algorithm handles irregularly sampled multivariate time series of variable lengths and provides shape-based representations of temporal data. In this work, we establish a representation theorem for the graph of a time series and derive its consequences on the LDDMM framework. We showcase the advantages of our representation compared to existing methods using synthetic data and real-world examples motivated by biomedical applications.

1 Introduction

Our goal is to analyze the inter-individual variability within a time series dataset, an approach of significant interest in physiological contexts [25, 58, 4, 21]. Specifically, we aim to develop an unsupervised feature representation method that encodes the specificities of individual time series in comparison to a reference time series. In physiology, examining the various "shapes" in a time series related to biological phenomena and their variations due to individual differences or pathological conditions is common. However, the term "shape" lacks a precise definition and is more intuitively understood as the silhouette of a pattern in a time series. In this paper, we refer to the shape of a time series as the graph of this signal.

¹Complete affiliation: Université Paris Saclay, Université Paris Cité, ENS Paris Saclay, CNRS, SSA, INSERM, Centre Borelli, F-91190, Gif-sur-Yvette, France.

*Corresponding author. Contact at thibaut.germain@ens-paris-saclay.fr

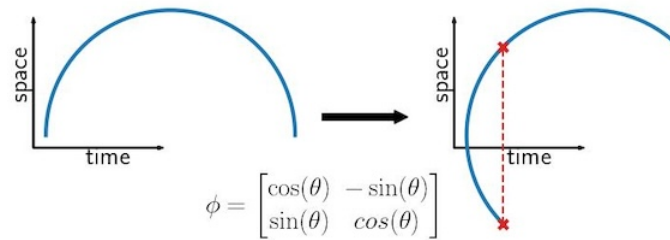


Figure 1: A time series' graph $G = \{(t, s(t)) : t \in I\}$ can lose its structure after applying a general diffeomorphism ϕ . G : a time value can be related to two values on the space axis.

Although community structures with representatives can be learned in an unsupervised manner [55, 39] using contrastive loss [20, 54, 39] or similarity measures [2, 21, 45, 62], the study of inter-individual variability of shapes within a cluster [42, 51] remains an open problem in unsupervised representation learning (URL), particularly for *irregularly sampled* time series with *variable lengths*.

Our work explicitly focuses on learning shape-based representation of time series. First, we propose to view the shape of a time series not merely as its curve $\{s_t : t \in I\}$, but as its graph $G(s) = \{(t, s(t)) : t \in I\}$. Then, building on the shape analysis literature [5, 57], we adopt the Large Deformation Diffeomorphic Metric Mapping (LDDMM) framework [5, 57] to analyze these graphs. The core idea is to represent each element $G(s^j)$ of a dataset $(s^j)_{j \in [N]}$ as the transformation of a reference graph $G(s_0)$ by a diffeomorphism ϕ_j , i.e. $G(s^j) \sim \phi_j.G(s_0)$. The diffeomorphism ϕ_j is learned by integrating an ordinary differential equation parameterized by a Reproducing Kernel Hilbert Space (RKHS). The parameters $(\alpha_j)_{j \in [N]}$ encoding the diffeomorphisms $(\phi_j)_{j \in [N]}$ yield the representation features of the graphs $(G(s^j))_{j \in [N]}$. Finally, these shape-encoding features can be used as inputs to any statistical or machine-learning model.

However, a time series graph transformation by a general diffeomorphism is not always a time series graph, see e.g. Figure 1, thus a time series graph is more than a simple curve [23]. Our contributions arise from this observation: we specify the class of diffeomorphisms to consider and show how to learn them. This change is fruitful in representing transformations of time series graphs as illustrated in Figure 2.

Our contributions can be summarized as follows:

- We propose an unsupervised method (TS-LDDMM) to analyze the inter-individual variability of shapes in a time series dataset (Section 4). In particular, the method can handle multivariate time series *irregularly sampled* and with *variable sizes*.
- We motivate our extension of LDDMM to time series by introducing a theoretical framework with a representation theorem for time series graph (Theorem 1) and kernels related to their structure (Lemma 1).
- We demonstrate the identifiability of the model by estimating the true generating parameter of synthetic data, and we highlight the sensitivity of our method concerning its hyperparameters (Appendix G.1), also providing guidelines for tuning (Appendix D).
- We highlight the *interpretability* of TS-LDDMM for studying the inter-individual variability in a clinical dataset (Section 5).
- We illustrate the quantitative interest of such representation on classification tasks on real shape-based datasets with regular and irregular sampling (Appendices H and I).

2 Notations

We denote by integer ranges by $[k : l] = \{k, \dots, l\} \subset \mathcal{P}(\mathbb{Z})$ and $[l] = [1 : l]$ with $k, l \in \mathbb{N}$, by $C^m(I, E)$ the set of m -times continuously differentiable function defined on an open set U to a normed vector space E , by $\|u\|_\infty = \sup_{x \in U} |u(x)|$ for any bounded function $u : U \rightarrow E$, and by $\mathbb{N}_{>0}$ is the set of positive integers.

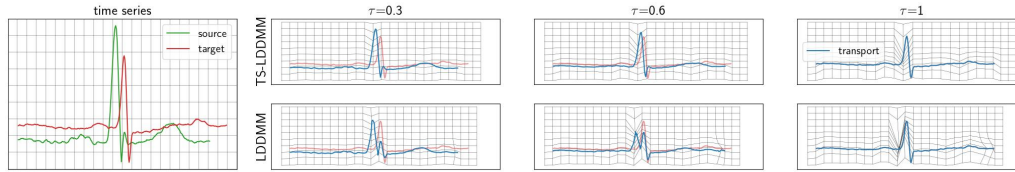


Figure 2: LDDMM and TS-LDDMM are applied to ECG data. We observe that LDDMM, using a general Gaussian kernel, does not learn the time translation of the first spike but changes the space values, i.e., one spike disappears before emerging at a translated position. At the same time, TS-LDDMM handles the time change in the shape. This difference of *deformations* implies differences in features *representations*.

3 Background on LDDMM

In this part, we expose how to learn the diffeomorphisms $(\phi_j)_{j \in [N]}$ using LDDMM, initially introduced in [5]. In a nutshell, for any $j \in [N]$, ϕ_j corresponds to a differential flow related to a learnable velocity field belonging to a well-chosen Reproducing Kernel Hilbert Space (RKHS).

In the next section, time series are going to be represented by diffeomorphism parameters $(\alpha_j)_{j \in [N]}$. That is why LDDMM is chosen since it offers a parametrization for diffeomorphisms that is sparse and interpretable, two features particularly relevant in the biomedical context.

The basic problem that we consider in this section is the following. Given a set of targets $\mathbf{y} = (y_i)_{i \in [T_2]}$ in $\mathbb{R}^{d'}$, a set of starting points $\mathbf{x} = (x_i)_{i \in [T_1]}$ in $\mathbb{R}^{d'}$, we aim to find a diffeomorphism ϕ such that the finite set of points $\mathbf{y}$ is similar in a certain sense to the set of finite sets of transformed points $\phi \cdot \mathbf{x} = (\phi(x_i))_{i \in [T_1]}$. The function ϕ is occasionally referred to as a *deformation*. In general, these sets $\mathbf{x}, \mathbf{y}$ are meshes of continuous objects, e.g., surfaces, curves, images, and so on.

Representing diffeomorphisms as deformations. Such *deformations* ϕ are constructed via differential flow equations, for any $x_0 \in \mathbb{R}^{d'}$ and $\tau \in [0, 1]$:

$$\frac{dX(\tau)}{d\tau} = v_\tau(X(\tau)), \quad X(0) = x_0, \phi_\tau^v(x_0) = X(\tau), \quad \phi^v = \phi_1^v, \quad (1)$$

where the velocity field is $v : \tau \in [0, 1] \mapsto v_\tau \in \mathbf{V}$ and $\mathbf{V}$ is a Hilbert space of continuously differentiable function on $\mathbb{R}^{d'}$. If $\|du\|_\infty + \|u\|_\infty \leq \|u\|_\mathbf{V}$ for any $u \in \mathbf{V}$ and $v \in L^2([0, 1], \mathbf{V}) = \{v \in C^0([0, 1], \mathbf{V}) : \int_0^1 \|v_\tau\|_\mathbf{V}^2 d\tau < \infty\}$, by [22, Theorem 5] ϕ^v exists and belongs to $\mathcal{D}(\mathbb{R}^{d'})$, where we denote by $\mathcal{D}(\mathbf{O})$ the set of diffeomorphism defined on an open set $\mathbf{O}$ to $\mathbf{O}$. Therefore, for any choice of v , ϕ^v defines a valid deformation. This offers a general recipe to construct diffeomorphism given a functional space $\mathbf{V}$.

With this in mind, the velocity field v fitting the data can be estimated by minimizing $v \in L^2([0, 1], \mathbf{V}) \mapsto \mathcal{L}(\phi^v \cdot \mathbf{x}, \mathbf{y})$, where $\mathcal{L}$ is an appropriate loss function. However, two computational challenges arise. First, this optimization problem is ill-posed, and a penalty term is needed to obtain a unique solution. In addition, a parametric family $\mathbf{V}_\Theta \subset L^2([0, 1], \mathbf{V})$, parameterized by Θ , is sought to efficiently solve this minimization problem.

From deformations to geodesics. It has been proposed in [40] to interpret $\mathbf{V}$ as a tangent space relative to the group of diffeomorphisms $\mathbf{H} = \{\phi^v : v \in L^2([0, 1], \mathbf{V})\}$. Following this geometric point of view, geodesics can be constructed on $\mathbf{H}$ by using the following squared norm

$$\mathcal{R}^2 : g \in \mathbf{H} \mapsto \inf_{v \in L^2([0, 1], \mathbf{V}) : g = \phi^v} \int_0^1 \|v_\tau\|_\mathbf{V}^2 d\tau \quad (2)$$

By deriving differential constraints related to the minimum of (2) and using Cauchy-Lipschitz conditions, geodesics can be defined only by giving the starting point and the initial velocity $v_0 \in \mathbf{V}$ [40], as straight lines in Euclidean space. Denoting by $\tau \mapsto \rho_{v_0}(\tau) \in \mathbf{H}$ the geodesic starting from

²Note that we denote by $d' \in \mathbb{N}$ the ambient space

the identity with initial velocity $v_0 \in \mathbf{V}$, the exponential map is defined as $\varphi^{\{v_0\}} \triangleq \rho_{v_0}(1)$. Using $\varphi^{\{v_0\}}$ instead of ϕ^v , the previous matching problem becomes a *geodesic shooting problem*:

$$\inf_{v_0 \in \mathbf{V}} \mathcal{L}(\varphi^{\{v_0\}}.\mathbf{x}, \mathbf{y}). \quad (3)$$

Using $\varphi^{\{v_0\}}$ instead of ϕ^v for any $v \in L^2([0, 1], \mathbf{V})$ regularizes the problem and induces a sparse representation for the learning diffeomorphisms. Moreover, by setting $\mathbf{V}$ as an RKHS, the geodesic shooting problem has a unique solution and becomes tractable, as described in the next section.

Discrete parametrization of diffeomorphism. In this part, $\mathbf{V}$ is chosen as an RKHS [6] generated by a smooth kernel K (e.g., Gaussian). We follow [17] and define a discrete parameterization of the velocity fields to perform geodesics shooting (3). The initial velocity field v_0 is chosen as a finite linear combination of the RKHS basis vector fields, $\mathbf{n}_0$ control points $\mathbf{X}_0 = (x_{k,0})_{k \in [\mathbf{n}_0]} \in (\mathbb{R}^{d'})^{\mathbf{n}_0}$ and momentum vectors $\alpha_0 = (\alpha_{k,0})_{k \in [\mathbf{n}_0]} \in (\mathbb{R}^{d'})^{\mathbf{n}_0}$ are defined such that for any $x \in \mathbb{R}^{d'}$,

$$v_0(\alpha_0, \mathbf{X}_0)(x) = \sum_{k=1}^{\mathbf{n}_0} K(x, x_{k,0}) \alpha_{k,0}. \quad (4)$$

In our applications, the control points $(x_{k,0})_{k \in [\mathbf{n}_0]}$ can be understood as the discretized graph $(t_k, \mathbf{s}_0(t_k))_{k \in [\mathbf{n}_0]}$ of a starting time series $\mathbf{s}_0$. With this parametrization of v_0 , [40] show that the velocity field v of the solution of (3) keeps the same structure along time, such that for any $x \in \mathbb{R}^{d'}$ and $\tau \in [0, 1]$,

$$\begin{aligned} v_\tau(x) &= \sum_{k=1}^{\mathbf{n}_0} K(x, x_k(\tau)) \alpha_k(\tau), \\ \begin{cases} \frac{dx_k(\tau)}{d\tau} = v_\tau(x_k(\tau)), & \frac{d\alpha_k(\tau)}{d\tau} = - \sum_{l=1}^{\mathbf{n}_0} \frac{d}{dx_k(\tau)} K(x_k(\tau), x_l(\tau)) \alpha_l(\tau)^\top \alpha_k(\tau) \\ \alpha_k(0) = \alpha_{k,0}, & x_k(0) = x_{k,0}, k \in [\mathbf{n}_0] \end{cases} \end{aligned} \quad (5)$$

These equations are derived from the hamiltonian $H : (\alpha_k, x_k)_{k \in [\mathbf{n}_0]} \mapsto \sum_{k,l=1}^{\mathbf{n}_0} \alpha_k^\top K(x_k, x_l) \alpha_l$, such that the velocity norm is preserved $\|v_\tau\|_{\mathbf{V}} = \|v_0\|_{\mathbf{V}}$ for any $\tau \in [0, 1]$. By (5), the velocity field related to a geodesic v^* is fully parametrized by its initial control points and momentum $(x_{k,0}, \alpha_{k,0})_{k \in [\mathbf{n}_0]}$. Thus, given a set of targets $\mathbf{y} = (y_i)_{i \in [T_2]}$ in $\mathbb{R}^{d'}$, a set of starting points $\mathbf{x} = (x_{i,0})_{i \in [T_1]}$ in $\mathbb{R}^{d'}$, a RKHS's kernel $K : \mathbb{R}^{d'} \times \mathbb{R}^{d'} \rightarrow \mathbb{R}^{d' \times d'}$, a distance on sets $\mathcal{L}$, a numerical integration scheme of ODE and a penalty factor $\lambda > 0$, the basic geodesic shooting step minimizes the following function using a gradient descent method:

$$\mathcal{F}_{\mathbf{x}, \mathbf{y}} : (\alpha_k)_{k \in [T_1]} \mapsto \mathcal{L}(\varphi^{\{v_0\}}.\mathbf{x}, \mathbf{y}) + \lambda \|v_0\|_{\mathbf{V}}^2, \quad (6)$$

where v_0 is defined by (4) and $\varphi^{\{v_0\}}.\mathbf{x}$ is the result of the numerical integration of (5) using control points $\mathbf{x}$ and initial momentums $(\alpha_k)_{k \in [T_1]}$.

Relation to Continuous Normalizing Flows. One particular popular choice to address the problem of learning a diffeomorphism or a velocity field is Normalizing Flows [47, 32] (NF) or their continuous counterpart [13, 24, 48] (CNF). However, we do not rely on this class of learning algorithms for several reasons. Indeed, existing and simple normalizing flows are not suitable for the type of data that we are interested in this paper [19, 16]. In addition, they are primarily designed to have tractable Jacobian functions, while we do not require such property in our applications. Finally, the use of a differential flow solution of an ODE (1) trick is also at the basis of CNF, which then consists of learning a velocity field to address in fitting the data through a loss aiming to address the problem at hand. Nevertheless, the main difference between CNF and LDDMM lies in the parametrization of the velocity field. LDDMM uses kernels to derive closed form formula and enhance interpretability while NF and CNF take advantage of deep neural networks to scale with large dataset in high dimensions.

4 Methodology

We consider in this paper observations which consist in a population of N multivariate time series, for any $j \in [N]$, $s^j \in C^1(I_j, \mathbb{R}^d)$. However, we can only access a n_j -samples $\tilde{s}^j = (\tilde{s}_i^j = s^j(t_i^j))_{i \in [n_j]}$

collected at timestamps $(t_i^j)_{i \in [n_j]}$ for any $j \in [N]$. Note that **the number of samples n_j is not necessarily the same across individuals** and the timestamps can be **irregularly sampled**. We assume the time series population is globally homogeneous regarding their "shapes" even if inter-individual variability exists. Intuitively speaking, the "shape" of a time series $s : \mathbb{I} \rightarrow \mathbb{R}^d$ is encoded in its graphs $G(s)$ defined as the set $\{(t, s(t)) : t \in \mathbb{I}\}$ and not only in its values $s(\mathbb{I}) = \{s(t) : t \in \mathbb{I}\}$ since the time axis is crucial. As a motivating use-case, s^j can be the time series of a heartbeat extracted from an individual's electrocardiogram (ECG), see Figure 2. The homogeneity in a resulting dataset comes from the fact that humans have similar shapes of heartbeat [61, 37].

The deformation problem. In this paper, we aim to study the inter-individual variability in the dataset by finding a relevant representation of each time series. Inspired from the framework of shape analysis [57], addressing similar problems in morphology, we suggest to represent each time series' graph $G(s^j)$ as the transformation of a reference graph $G(s_0)$, related to a time series $s_0 : \mathbb{I} \rightarrow \mathbb{R}^d$, by a diffeomorphism ϕ_j on $\mathbb{R}^{d+1}$, for any $j \in [N]$,

$$\phi_j.G(s_0) = \{\phi_j(t, s_0(t)), t \in \mathbb{I}\}. \quad (7)$$

s_0 will be understood as the typical representative shape common to the collection of time series $(s^j)_{j \in [N]}$. As s_0 is supposed to be fixed, then the representation of the time series $(s^j)_{j \in [N]}$ boils down to the one of the transformation $(\phi_j)_{j \in [N]}$. We aim to learn $G(s_0)$ and $(\phi_j)_{j \in [N]}$.

Optimization related to (7). Defining the *discretized graphs* of the time series $(s^j)_{j \in [N]}$ and a discretization of the reference graph $G(s_0)$ as, for any $j \in [N]$,

$$\mathbf{y}_j = G(\tilde{s}^j) = (t_i^j, \tilde{s}_i^j)_{i \in [n_j]} \in (\mathbb{R}^{d+1})^{n_j}, \quad \tilde{G}_0 = (t_i^0, \tilde{s}_i^0)_{i \in [\mathbf{n}_0]} \in (\mathbb{R}^{d+1})^{\mathbf{n}_0},$$

with $\mathbf{n}_0 = \text{median}((n_j)_{j \in [N]})$, the representation problem given in (7) boils down solving:

$$\arg\min_{\tilde{G}_0, (\alpha_k^j)_{k \in [\mathbf{n}_0]}} \sum_{j=1}^N \mathcal{F}_{\tilde{G}_0, \mathbf{y}_j} \left((\alpha_k^j)_{k \in [\mathbf{n}_0]} \right), \quad (8)$$

which is carried out by gradient descent on the control points $\tilde{G}_0$ and the momentums $\alpha_j = (\alpha_k^j)_{k \in [\mathbf{n}_0]}$ for any $j \in [N]$, initialized by a dataset's time series graph of size $\mathbf{n}_0$ and by $0_{(d+1)\mathbf{n}_0}$ respectively. The optimization hyperparameter details are given in Appendix E.1. The result of the minimization $\tilde{G}_0$ is then considered as the $\mathbf{n}_0$ -samples of a common time series s_0 and the momentums α_j encoding ϕ_j yields a feature vector in $\mathbb{R}^{d\mathbf{n}_0}$ of s^j for any $j \in [N]$. Finally, the vectors $(\alpha_j)_{j \in [N]}$ can be analyzed with any statistical or machine learning tools such as Principal Components Analysis (PCA), Latent Discriminant Analysis (LDA), longitudinal data analysis and so on.

Nevertheless, (8) asks to define a kernel and a loss in order to perform geodesic shooting (6), which is the purpose of the following subsection.

4.1 Application of LDDMM to time series analysis: TS-LDDMM

This section presents our theoretical contribution: we tailor the LDDMM framework to handle time series data. The reason is that applying a general diffeomorphism ϕ from $\mathbb{R}^{d+1}$ to a time series' graph $G(s)$ can result in a set $\phi.G(s)$ that does not correspond to the graph of any time series, as illustrated in the Figure 1. Thus, time series graphs have more structure than a simple 1D curve [23] and deserve their unique analysis, which will prove fruitful as demonstrated in Section 5.

To address this challenge, we need to identify an RKHS kernel $K : \mathbb{R}^{d+1} \times \mathbb{R}^{d+1} \rightarrow \mathbb{R}^{(d+1)^2}$ that generates deformations preserving the structure of the time series graph. This goal motivates us to clarify, in Theorem 1, the specific representation of diffeomorphisms we require before presenting a class of kernels that produce deformations with this representation.

Similarly, selecting a loss function on sets $\mathcal{L}$ that considers the temporal evolution in a time series' graph is crucial for meaningful comparisons with time series data. Consequently, we introduce the oriented Varifold distance.

A representation separating space and time. We prove that two time series graphs can always be linked by a time transformation composed with a space transformation. Moreover, a time series graph transformed by this kind of transformation is always a time series graph. We define $\Psi_\gamma \in$

$\mathcal{D}(\mathbb{R}^{d+1}) : (t, x) \in \mathbb{R}^{d+1} \rightarrow (\gamma(t), x)$ for any $\gamma \in \mathcal{D}(\mathbb{R})$ and $\Phi_f : (t, x) \in \mathbb{R}^{d+1} \rightarrow (t, f(t, x))$ for any $f \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$. We have the following representation theorem. All proofs are given in Appendix B.

Denote by $G(s) \triangleq \{(t, s(t)) : t \in I\}$ the graph of a time series $s : I \rightarrow \mathbb{R}^d$ and $\phi.G(s) \triangleq \{\phi(t, s(t)) : t \in I\}$ the action of $\phi \in \mathcal{D}(\mathbb{R}^{d+1})$ on $G(s)$.

Theorem 1. *Let $s : J \rightarrow \mathbb{R}^d$ and $s_0 : I \rightarrow \mathbb{R}^d$ be two continuously differentiable time series with I, J two intervals of $\mathbb{R}$. There exist $f \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ and $\gamma \in \mathcal{D}(\mathbb{R})$ such that $\gamma(I) = J$ and $\Phi_f \in \mathcal{D}(\mathbb{R}^{d+1})$,*

$$G(s) = \Pi_{\gamma, f}.G(s_0), \quad \Pi_{\gamma, f} = \Psi_\gamma \circ \Phi_f.$$

Moreover, for any $\bar{f} \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ and $\bar{\gamma} \in \mathcal{D}(\mathbb{R})$, there exists a continuously differentiable time series $\bar{s}$ such that $G(\bar{s}) = \Pi_{\bar{\gamma}, \bar{f}}.G(s_0)$

Remark 2. Note that for any $\gamma \in \mathcal{D}(\mathbb{R})$ and $s \in C^0(I, \mathbb{R}^d)$,

$$\{(\gamma(t), s(t)), t \in I\} = \{(t, s \circ \gamma^{-1}(t)) : t \in \gamma(I)\}.$$

As a result, Ψ_γ can be understood as a temporal reparametrization and Φ_f encodes the transformation about the space.

Choice for the kernel associated with the RKHS $\mathbb{V}$ As depicted on Figure 1-2, we can not use any kernel K to apply the previous methodology to learn deformations on time series' graphs. We describe and motivate our choice in this paragraph. Denote the one-dimensional Gaussian kernel by $K_\sigma^{(a)}(x, y) = \exp(-|x - y|^2/\sigma)$ for any $(x, y) \in (\mathbb{R}^a)^2$, $a \in \mathbb{N}$ and $\sigma > 0$. To solve the geodesic shooting problem (6) on $\mathbb{R}^{d+1}$, we consider for $\mathbb{V}$ the RKHS associated with the kernel defined for any $(t, x), (t', x') \in (\mathbb{R}^{d+1})^2$:

$$K_G((t, x), (t', x')) = \begin{pmatrix} c_0 K_{\text{time}} & 0 \\ 0 & c_1 K_{\text{space}} \end{pmatrix}, \quad (9)$$

$$K_{\text{space}} = K_{\sigma_{T,1}}^{(1)}(t, t') K_{\sigma_x}^{(d)}(x, x') I_d, \quad K_{\text{time}} = K_{\sigma_{T,0}}^{(1)}(t, t'),$$

parametrized by the widths $\sigma_{T,0}, \sigma_{T,1}, \sigma_x > 0$ and the constants $c_0, c_1 > 0$. This choice for K_G is motivated by the representation Theorem 1 and the following result.

Lemma 1. *If we denote by $\mathbb{V}$ the RKHS associated with the kernel K_G , then for any vector field v generated by (5) with v_0 satisfying (4), there exist $\gamma \in \mathcal{D}(\mathbb{R})$ and $f \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ such that $\phi^v = \Psi_\gamma \circ \Phi_f$.*

Instead of Gaussian kernels, other types of smooth kernels can be selected as long as the structure (9) is respected.

Remark 3. *With this choice of kernel, the features associated with the time transformation can be extracted from the momentums $(\alpha_{k,0})_{k \in [n_0]} \in (\mathbb{R}^{d+1})^{n_0}$ in (4) by taking the coordinates related to time. However, the features related to the space transformation are not only in the space coordinates since the related kernel K_{space} depends on time as well. The kernel's representation has been carefully designed to integrate both space and time, while ensuring that time remains independent of space. Initially, we considered separating the spatial and temporal components. However, post-hoc analysis of such a representation proved to be challenging. The separated spatial and temporal representations are correlated, and understanding this correlation is essential for interpreting the data. As a result, concatenating the two representations becomes necessary, though there is no straightforward method for doing so, as they are not commensurable. Consequently, we opted for a representation that inherently integrates both space and time.*

In Appendix D, we give guidelines for selecting the hyperparameters $(\sigma_{T,0}, \sigma_{T,1}, \sigma_x, c_0, c_1)$.

Loss This section specifies the distance function $\mathcal{L}$ introduced in the loss function defined in (6).

In practice, we can only access discretized graphs of time series, $(t_i^j, \tilde{s}_i^j)_{i \in [n_j]}$ for any $j \in [N]$, that are potentially of different sizes n_j and sampled at different timestamps $(t_i^j)_{i \in [n_j]}$ for any $j \in [N]$. Usual metrics, such as the Euclidean distance, are not appealing as they make the underlying assumptions of equal size sets and the existence of a pairing between points. Distances between measures on sets

(taking the empirical distribution), such as Maximum Mean Discrepancy (MMD) [18, 9], alleviate those issues; however, MMD only accounts for positional information and lacks information about the time evolution between sampled points. A classical data fidelity metric from shape analysis corresponding to the distance between *oriented varifolds* associated with curves alleviates this last issue [30]. Intuitively, an oriented varifold is a measure that accounts for positional and tangential information about the underlying curves at sample points. More details and information about *oriented varifolds* can be found in Appendix C.

More precisely, given two sets $G_0 = (g_i^0)_{i \in [T_0]}, G_1 = (g_i^1)_{i \in [T_1]} \in (\mathbb{R}^{d+1})^{T_1}$ and a kernel³ $k : (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2 \rightarrow \mathbb{R}$ verifying [30, Proposition 2 & 4], for any $\xi \in \{0, 1\}$ and $i \in [T_\xi - 1]$, denoting the center and length of the i^{th} segment $[g_i^\xi, g_{i+1}^\xi]$ by $c_i^\xi = (g_i^\xi + g_{i+1}^\xi)/2$, $l_i^\xi = \|g_{i+1}^\xi - g_i^\xi\|$, and $\vec{v}_i^\xi = (g_{i+1}^\xi - g_i^\xi)/l_i^\xi$, the varifold distance between G_0 and G_1 is defined as,

$$\begin{aligned} d_{W^*}^2(G_0, G_1) &= \sum_{i,j=1}^{T_0-1} l_i^0 k((c_i^0, \vec{v}_i^0), (c_j^0, \vec{v}_j^0)) l_j^0 - 2 \sum_{i=1}^{T_0-1} \sum_{j=1}^{T_1-1} l_i^0 k((c_i^0, \vec{v}_i^0), (c_j^1, \vec{v}_j^1)) l_j^1 \\ &+ \sum_{i,j=1}^{T_1-1} l_i^1 k((c_i^1, \vec{v}_i^1), (c_j^1, \vec{v}_j^1)) l_j^1 \end{aligned}$$

In practice, we set the kernel k as the product of two anisotropic Gaussian kernels, k_{pos} and k_{dir} , such that for any $(x, \vec{u}), (y, \vec{v}) \in (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2$

$$k((x, \vec{u}), (y, \vec{v})) = k_{\text{pos}}(x, y) k_{\text{dir}}(\vec{u}, \vec{v}).$$

Note that the loss kernel k has nothing to do with the velocity field kernel denoted by K_G or K specified in Section 4.1. Finally, we define the data fidelity loss function, $\mathcal{L}$, as a sum of $d_{W^*}^2$ using different kernel's width parameters σ to incorporate multiscale information. $\mathcal{L}$ is indeed differentiable with respect to its first variable. The specific kernels $k_{\text{pos}}, k_{\text{dir}}$ that we use in our experiments are given Appendix C.1. For further readings on curves and surface representation as varifolds, readers can refer to [30, 12].

A pedagogical online application is available to inspect the effect of hyperparameters on geodesic shooting (5) and registration (6).

5 Experiments

The source code is available on Github⁴. For conciseness, several experiments are relegated in appendix:

1. **TS-LDDMM representation identifiability, Appendix G:** On synthetic data, we evaluate the ability of our method to retrieve the parameter v_0^* that encodes the deformation $\varphi^{\{v_0^*\}}$ acting on a time series graph G by solving the geodesic shooting problem (6) between G and $\varphi^{\{v_0^*\}}.G$. **Results** show that TS-LDDMM representations are identifiable or weakly identifiable depending on the velocity field kernel K_G specification.
2. **Robustness to irregular sampling, Appendix H:** We compare the robustness of TS-LDDMM representation with 9 URL methods handling irregularly sampled multivariate time series on 15 shape-based datasets (7 univariates & 8 multivariates). We assess methods' classification performances under regular sampling (0% missing rate) and three irregular sampling regimes (30%, 50%, and 70% missing rates), according to the protocol depicted in [31]. **Results** show that our method, TS-LDDMM, outperforms all methods for sampling regimes with missing rates: 0%, 30%, and 50%.
3. **Classification benchmark on regularly sampled datasets, Appendix I:** We compare performances of a kernel support vector machine (SVC) algorithm based on TS-LDDMM representation with 3 state-of-the-art classification methods from shape analysis on 15 shape-based datasets (7 univariates & 8 multivariates). **Results** show that the TS-LDDMM-based method outperforms other methods (best performances over 13 datasets), making TS-LDDMM representation relevant for time series shape analysis.

³ $\mathbb{S}^d = \{x \in \mathbb{R}^{d+1} : |x| = 1\}$

⁴<https://github.com/thibaut-germain/TS-LDDMM>

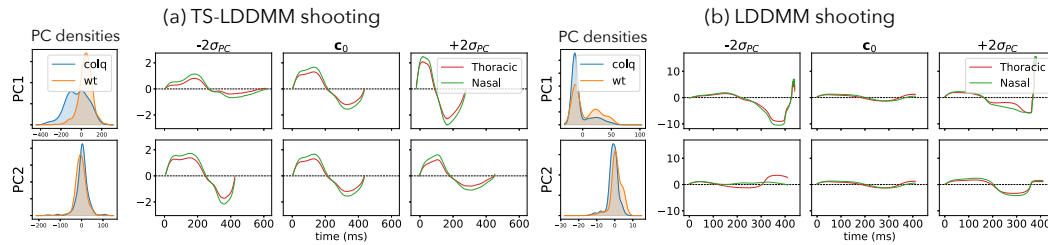


Figure 3: Analysis of the two principal components (PC) related to mice ventilation before exposure with TS-LDDMM representations (a), and LDDMM (b). In both cases and for all PC, the left plot displays PC densities according to mice genotype and right plot displays deformations of the reference graph c_0 along each PC.

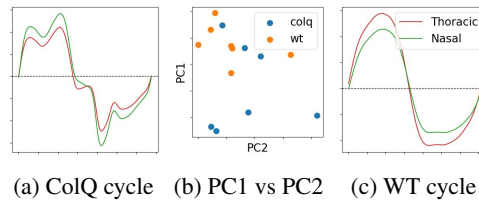


Figure 4: (a) a ColQ respiratory cycle sample. (b) Referent respiratory cycle of individual mouse c_0^j in the TS-LDDMM PC1-PC2 coordinates system of c_0 . (c) a WT respiratory cycle sample.

4. **Noise sensitivity for learning the reference graph, Appendix J:** We evaluate the noise sensitivity of TS-LDDMM and Shape-FPCA [60] for learning the reference graph on a synthetic dataset and for several levels of additive Gaussian noise. **Results** show that both methods are sensitive to noise. However, TS-LDDMM preserves the overall shape while shape-FPCA alters the shape depending on the noise level.

5.1 Interpretability: mice ventilation analysis

This experiment highlights the *interpretability* of TS-LDDMM representation for studying the inter-individual variability in biomedical applications. We consider a time series dataset monitoring the evolution of mice's nasal and thoracic airflow when exposed to a drug altering respiration [41]. The dataset includes recordings of 7 control mice (WT) and 7 mutant mice (ColQ) with an enzyme deficiency. The enzyme is involved in the respiration regulation, and the drug inhibits its activity. For each mouse, airflows were monitored for 15 to 20 minutes before the drug exposure and then for 35 to 40 minutes. A complete description of the dataset is given in the Appendix F.1.

Experimental protocol. We considered two experimental scenarios; the first focuses on mice ventilation before exposure to explore the inter-individual and genotype-specific variabilities. The second focuses on whole recordings to analyze the evolution of mice's ventilation after drug exposure. In both cases, the baseline protocol consists of first extracting N respiratory cycles from the datasets with the procedure described in [21]. Then, learning the referent respiratory cycle c_0 and the representations of respiratory cycles $(\alpha_0^j)_{j \in [N]}$ by solving (8) using TS-LDDMM. α_0^j being the momentum of the initial velocity field of the geodesic encodings the diffeomorphisms mapping c_0 to the j^{th} respiratory cycle. Finally, performing a Kernel-PCA on the initial velocity fields (4) belonging to V and encoded by the pairs $(\alpha_0^j, c_0)_{j \in [N]}$. The first experiment includes $N_1 = 700$ cycles collected before exposure. The second experiment includes $N_2 = 1400$ cycles with 25% (resp. 75%) before (resp. after) exposure. We also performed the first experimental scenario with LDDMM representation, and Appendix K describes the settings of both methods. Essentially, varifold losses are identical for both methods, and the velocity field kernels are set to encompass time and space scales. In addition, Appendix K presents a comparison between TS-LDDMM and Shape-FPCA on the second scenario.

Geodesic shooting along principal component directions. Any principal component (PC), noted v_0^{pc} , from a kernel-PCA in V , is itself an initial velocity field encoded by a pair (c_0, α_0^{pc}) . PCs encode the principal axis of deformations, and it is possible to shoot along the geodesic they encode with the differential equations (5), enabling interpretation of the main sources of deformations.

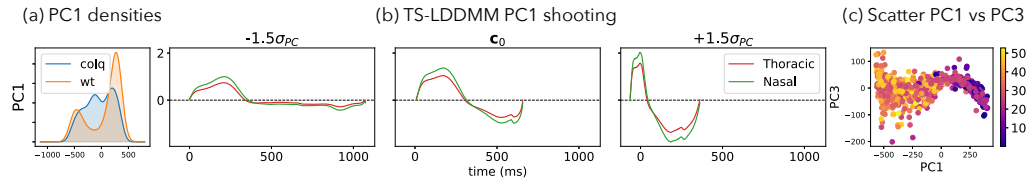


Figure 5: Analysis of the first Principal Component (PC1) related to mice ventilation before and after exposure with TS-LDDMM representations. **(a)** displays PC densities per mice genotype, **(b)** illustrates deformations of the reference respiratory cycle c_0 along PC1, and **(c)** displays all respiratory cycles with respect to time in PC1 and PC3 coordinates

Mice ventilation before exposure. We focus on the analysis of the two first Principal Components (PC) for TS-LDDMM (Figure 3a) and LDDMM (Figure 3b). Looking at the geodesic shooting along PCs, Figure 3 shows that principal components learned with TS-LDDMM lead to deformations that remain respiratory cycles. In contrast, deformations learned with LDDMM are challenging to interpret as respiratory cycles. The LDDMM velocity field kernel is a Gaussian anisotropic kernel that accounts for time and space scales; however, the entanglement of time and space dimensions in the kernel does not guarantee the graph structure, and it makes the convergence of the method complex (relative varifold loss error: TS-LDDMM: 0.06, LDDMM: 0.11).

Regarding TS-LDDMM Figure 3a, its PCs refer to deformations directions carrying different physiological meanings. Indeed, the geodesic shooting along these directions indicates that PC1 accounts for variations of the total duration of a respiratory cycle, while PC2 expresses the trade-off between inspiration and expiration duration. In addition, the distribution of ColQ respiratory cycles along PC1 is wider than in WT mice, indicating that the adaptation of mutant mice to their enzyme deficiency is variable. This observation can also be seen in Figure 4b where a referent respiratory cycle c_0^i is learned by solving (8) for each mouse and is encoded in the (PC1,PC2) coordinate system of c_0 by registration (3). Indeed, the average respiratory cycles of ColQ mice are more spread out than WT mice's. Going back to the densities of PC1, ColQ mice distribution has a heavier tail toward negative values compared to WT mice. When shooting in the opposite direction of PC1, we can observe that the inspiration is divided into two steps. Congruently with [21], such inspirations indicate motor control difficulties due to enzyme deficiency. Figure 4a is an example of ColQ respiratory cycle with negative PC1 coordinate.

Mice ventilation evolution after drug exposure. This experiment focuses on the first principal components learned from TS-LDDMM representations of respiratory cycles randomly sampled before and after drug exposure. Figure 5a illustrates the geodesic shootings along PC1. Again, PC1 accounts for variations in respiratory cycle duration, but more importantly, it can be observed on the deformation at $-1.5 \sigma_{PC}$ the apparition of a long pause after inspiration. Congruently, Figure 5c indicates that pauses appear after drug exposure as cycles with negative PC1 values mainly occur after 20 minutes and present more variability along PC3. In addition, Figure 5b shows a bimodal distribution for WT mice with one of the peaks in the negative values. This peak was not observed in the previous experiment Figure 3a. It indicates that pauses after inspiration are prevalent in WT mice after drug exposure. On the other hand, the distributions of ColQ mice's respiratory cycles along PC1 in both experiments are similar and account for the same deformation, suggesting that ColQ mice weakly react to the drug exposure as they already adapt their enzyme deficiency.

Experiment Conclusion. Analyzing mice ventilation with TS-LDDMM representation highlights the method's ability to create meaningful interaction between experts and the data. Indeed, combining statistical and visual results shows that main deformations carry physiological meaning, enabling the characterization of some mice genotypes and the effects of drug exposure.

6 Related Works

Shape analysis focuses on statistical analysis of mathematical objects invariant under some deformations like rotations, dilations, or time parameterization. The main idea is to represent these objects in a complete Riemannian manifold $(\mathcal{M}, g)$ with a metric g adapted to the geometry of the problem [40]. Then, any set of points in $\mathcal{M}$ can be represented as points in the tangent space of their Frechet mean m_0 [44, 33] by considering their logarithms. The goal is to find a well-suited Riemannian structure according to the nature of the studied object.

LDDMM framework is a relevant shape analysis tool to represent curves as depicted in [23]. However, graphs of time series are a well-structured type of curve due to the inclusion of the temporal dimension that requires specific care (Figure 1). In a similar vein, Qiu *et al* [46] proposes a method for tracking anatomical shape changes in serial images using LDDMM. They include temporal evolution, but not for the same purpose: the aim is to perform longitudinal modeling of brain images.

Leaving the LDDMM representation, the results of [53, 26] address the representation of curves with the Square-Root Velocity (SRV) representation. However, the SRV representation is applied after reparametrization of the temporal dimension of the unit length segment. Consequently, the graph structure of the time series is not respected, and the original time evolution of the time series is not encoded in the final representation. Very recently, in a functional data analysis (FDA) framework, a paper [60] (Shape-FPCA) improved by representing the original time evolution. However, the space and time representations remain correlated, complicating post-hoc analysis, as discussed in Remark 3. Additionally, this method is tailored for *continuous objects* and applies only to time series of the *same length*, making the estimation more sensitive to noise. This issue can be addressed through interpolation, but this approach is not always reliable in sparse and irregular sampling scenarios. Most FDA approaches, as seen in [50, 63, 59], address this challenge using interpolation or basis function expansion. In summary, FDA methods typically separate space and time representations for continuous objects, whereas TS-LDDMM algorithm maintain a discrete-to-discrete analysis, inherently integrating both space and time representations.

Balancing between discrete and continuous elements is a challenging task. In the deep learning literature [13, 31, 56, 29, 36, 1], Neural Ordinary Differential Equations (Neural ODEs) [13] learn continuous latent representations using a vector field parameterized by a neural network, serving as a continuous analog to Residual Networks [64]. This approach was further enhanced by Neural Controlled Differential Equations (Neural CDEs) [31] for handling irregular time series, functioning as continuous-time analogs of RNNs [49]. Extending Neural ODEs, Neural Stochastic Differential Equations (Neural SDEs) introduce regularization effects [36], although optimization remains challenging. Leveraging techniques from continuous-discrete filtering theory, Ansari *et al.* [1] applied successfully Neural SDEs to irregular time series. Oh *et al.* [43] improved these results by incorporating the concept of controlled paths into the drift term, similar to how Neural CDEs outperform Neural ODEs. With TS-LDDMM, the representation is also derived from an ODE, but the velocity field is parameterized with kernels and optimized to have a minimal norm, which enhances interpretability.

All these state-of-the-art methods previously mentioned [23, 43, 60, 26] are compared to TS-LDDMM in Appendix H and Appendix I.

Compared to the Metamorphosis framework [7], LDDMM framework has weaker assumptions. The 3DMM framework requires that each mesh be re-parametrized into a consistent form where the number of vertices, triangulation, and the anatomical meaning of each vertex are consistent across all meshes, as stated in the introduction of [8]. In our context, we do not need such pre-processing; the time series graph can have different sizes.

7 Limitations and conclusion

This paper proposes a feature representation method, TS-LDDMM, designed for shape comparison on homogeneous time series datasets. We show on a real dataset its ability to study, with high interpretability, the inter-individual shape variability. As an unsupervised approach, it is user-friendly and enables knowledge transfer for different supervised tasks such as classification.

Although TS-LDDMM is already competitive for classification, its performances can be leveraged on more heterogeneous datasets using a hierarchical clustering extension, which is relegated for future work.

TS-LDDMM employs kernel computations, which require specific libraries (e.g., KeOps [11]) to be efficient and scalable. However, in our experiments, the time complexity of TS-LDDMM is comparable to that of competitors. It is clear that TS-LDDMM needs to be extended to handle very large datasets with high-dimensional time series (such as videos).

Additionally, TS-LDDMM requires tuning several hyperparameters, though this is a common requirement among competitors [23, 43, 60, 26]. In future work, adaptive methods are expected to be developed to provide a more user-friendly interface.

Acknowledgments and Disclosure of Funding

This work was supported by grants from Région Ile-de-France (DIM MathInnov). Charles Truong is funded by the PhLAMMES chair of ENS Paris-Saclay.

References

- [1] Abdul Fatir Ansari, Alvin Heng, Andre Lim, and Harold Soh. Neural continuous-discrete state space models for irregularly-sampled time series. In *International Conference on Machine Learning*, pages 926–951. PMLR, 2023.
- [2] Asal Asgari. Clustering of clinical multivariate time-series utilizing recent advances in machine-learning. 2023.
- [3] Anthony Bagnall, Hoang Anh Dau, Jason Lines, Michael Flynn, James Large, Aaron Bostrom, Paul Southam, and Eamonn Keogh. The uea multivariate time series classification archive, 2018. *arXiv preprint arXiv:1811.00075*, 2018.
- [4] Ziv Bar-Joseph, Anthony Gitter, and Itamar Simon. Studying and modelling dynamic biological processes using time-series gene expression data. *Nature Reviews Genetics*, 13(8):552–564, 2012.
- [5] M Faisal Beg, Michael I Miller, Alain Trouvé, and Laurent Younes. Computing large deformation metric mappings via geodesic flows of diffeomorphisms. *International journal of computer vision*, 61:139–157, 2005.
- [6] Alain Berlinet and Christine Thomas-Agnan. *Reproducing kernel Hilbert spaces in probability and statistics*. Springer Science & Business Media, 2011.
- [7] Volker Blanz and Thomas Vetter. Face recognition based on fitting a 3d morphable model. *IEEE Transactions on pattern analysis and machine intelligence*, 25(9):1063–1074, 2003.
- [8] James Booth, Anastasios Roussos, Stefanos Zafeiriou, Allan Ponniah, and David Dunaway. A 3d morphable model learnt from 10,000 faces. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 5543–5552, 2016.
- [9] Karsten M Borgwardt, Arthur Gretton, Malte J Rasch, Hans-Peter Kriegel, Bernhard Schölkopf, and Alex J Smola. Integrating structured biological data by kernel maximum mean discrepancy. *Bioinformatics*, 22(14):e49–e57, 2006.
- [10] Claudio Carmeli, Ernesto De Vito, Alessandro Toigo, and Veronica Umanitá. Vector valued reproducing kernel hilbert spaces and universality. *Analysis and Applications*, 8(01):19–61, 2010.
- [11] Benjamin Charlier, Jean Feydy, Joan Alexis Glaunes, François-David Collin, and Ghislain Durif. Kernel operations on the gpu, with autodiff, without memory overflows. *Journal of Machine Learning Research*, 22(74):1–6, 2021.
- [12] Nicolas Charon and Alain Trouvé. The varifold representation of nonoriented shapes for diffeomorphic registration. *SIAM journal on Imaging Sciences*, 6(4):2547–2580, 2013.
- [13] Ricky TQ Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. *Advances in neural information processing systems*, 31, 2018.
- [14] Junyoung Chung, Caglar Gulcehre, KyungHyun Cho, and Yoshua Bengio. Empirical evaluation of gated recurrent neural networks on sequence modeling. *arXiv preprint arXiv:1412.3555*, 2014.
- [15] Hoang Anh Dau, Anthony Bagnall, Kaveh Kamgar, Chin-Chia Michael Yeh, Yan Zhu, Shaghayegh Gharghabi, Chotirat Ann Ratanamahatana, and Eamonn Keogh. The ucr time series archive. *IEEE/CAA Journal of Automatica Sinica*, 6(6):1293–1305, 2019.

- [16] Ruizhi Deng, Bo Chang, Marcus A Brubaker, Greg Mori, and Andreas Lehrmann. Modeling continuous stochastic processes with dynamic normalizing flows. *Advances in Neural Information Processing Systems*, 33:7805–7815, 2020.
- [17] Stanley Durrleman, Stéphanie Allasonnière, and Sarang Joshi. Sparse adaptive parameterization of variability in image ensembles. *International Journal of Computer Vision*, 101:161–183, 2013.
- [18] Gintare Karolina Dziugaite, Daniel M Roy, and Zoubin Ghahramani. Training generative neural networks via maximum mean discrepancy optimization. *arXiv preprint arXiv:1505.03906*, 2015.
- [19] Shibo Feng, Chunyan Miao, Ke Xu, Jiaxiang Wu, Pengcheng Wu, Yang Zhang, and Peilin Zhao. Multi-scale attention flow for probabilistic time series forecasting. *IEEE Transactions on Knowledge and Data Engineering*, 2023.
- [20] Jean-Yves Franceschi, Aymeric Dieuleveut, and Martin Jaggi. Unsupervised scalable representation learning for multivariate time series. *Advances in neural information processing systems*, 32, 2019.
- [21] Thibaut Germain, Charles Truong, Laurent Oudre, and Eric Krejci. Unsupervised classification of plethysmography signals with advanced visual representations. *Frontiers in Physiology*, 14:781, 2023.
- [22] Joan Glaunes. Transport par difféomorphismes de points, de mesures et de courants pour la comparaison de formes et l’anatomie numérique. *These de sciences, Université Paris*, 13, 2005.
- [23] Joan Glaunes, Anqi Qiu, Michael I Miller, and Laurent Younes. Large deformation diffeomorphic metric curve mapping. *International journal of computer vision*, 80:317–336, 2008.
- [24] Will Grathwohl, Ricky TQ Chen, Jesse Bettencourt, and David Duvenaud. Scalable reversible generative models with free-form continuous dynamics. In *International Conference on Learning Representations*, page 7, 2019.
- [25] Ella Guscelli, John I Spicer, and Piero Calosi. The importance of inter-individual variation in predicting species’ responses to global change drivers. *Ecology and Evolution*, 9(8):4327–4339, 2019.
- [26] Tae-Young Heo, Joon Myoung Lee, Myung Hun Woo, Hyeonseok Lee, and Min Ho Cho. Logistic regression models for elastic shape of curves based on tangent representations. *Journal of the Korean Statistical Society*, pages 1–19, 2024.
- [27] Sepp Hochreiter and Jürgen Schmidhuber. Long short-term memory. *Neural computation*, 9(8):1735–1780, 1997.
- [28] Heinz Gerd Hoymann. Lung function measurements in rodents in safety pharmacology studies. *Frontiers in pharmacology*, 3:156, 2012.
- [29] Junteng Jia and Austin R Benson. Neural jump stochastic differential equations. *Advances in Neural Information Processing Systems*, 32, 2019.
- [30] Irene Kaltenmark, Benjamin Charlier, and Nicolas Charon. A general framework for curve and surface comparison and registration with oriented varifolds. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3346–3355, 2017.
- [31] Patrick Kidger, James Morrill, James Foster, and Terry Lyons. Neural controlled differential equations for irregular time series. *Advances in Neural Information Processing Systems*, 33:6696–6707, 2020.
- [32] Ivan Kobyzev, Simon JD Prince, and Marcus A Brubaker. Normalizing flows: An introduction and review of current methods. *IEEE transactions on pattern analysis and machine intelligence*, 43(11):3964–3979, 2020.
- [33] Huiling Le. Locating fréchet means with application to shape spaces. *Advances in Applied Probability*, 33(2):324–338, 2001.

- [34] Mathias Lechner and Ramin Hasani. Learning long-term dependencies in irregularly-sampled time series. *arXiv preprint arXiv:2006.04418*, 2020.
- [35] Yurim Lee, Eunji Jun, Jaehun Choi, and Heung-Il Suk. Multi-view integrative attention-based deep representation learning for irregular clinical time-series data. *IEEE Journal of Biomedical and Health Informatics*, 26(8):4270–4280, 2022.
- [36] Xuanqing Liu, Tesi Xiao, Si Si, Qin Cao, Sanjiv Kumar, and Cho-Jui Hsieh. Neural ode: Stabilizing neural ode networks with stochastic noise. *arXiv preprint arXiv:1906.02355*, 2019.
- [37] Putri Madona, Rahmat Ilias Basti, and Muhammad Mahrus Zain. Pqrst wave detection on ecg signals. *Gaceta Sanitaria*, 35:S364–S369, 2021.
- [38] Larry Medsker and Lakhmi C Jain. *Recurrent neural networks: design and applications*. CRC press, 1999.
- [39] Qianwen Meng, Hangwei Qian, Yong Liu, Yonghui Xu, Zhiqi Shen, and Lizhen Cui. Unsupervised representation learning for time series: A review. *arXiv preprint arXiv:2308.01578*, 2023.
- [40] Michael I Miller, Alain Trouvé, and Laurent Younes. Geodesic shooting for computational anatomy. *Journal of mathematical imaging and vision*, 24:209–228, 2006.
- [41] Aurélie Nervo, André-Guilhem Calas, Florian Nachon, and Eric Krejci. Respiratory failure triggered by cholinesterase inhibitors may involve activation of a reflex sensory pathway by acetylcholine spillover. *Toxicology*, 424:152232, 2019.
- [42] Vit Niennattrakul and Chotirat Ann Ratanamahatana. Inaccuracies of shape averaging method using dynamic time warping for time series data. In *Computational Science–ICCS 2007: 7th International Conference, Beijing, China, May 27-30, 2007, Proceedings, Part I* 7, pages 513–520. Springer, 2007.
- [43] YongKyung Oh, Dongyoung Lim, and Sungil Kim. Stable neural stochastic differential equations in analyzing irregular time series data. In *The Twelfth International Conference on Learning Representations*, 2024.
- [44] Susovan Pal, Roger P Woods, Suchit Panjiyar, Elizabeth Sowell, Katherine L Narr, and Shantanu H Joshi. A riemannian framework for linear and quadratic discriminant analysis on the tangent space of shapes. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, pages 47–55, 2017.
- [45] John Paparrizos and Luis Gravano. k-shape: Efficient and accurate clustering of time series. In *Proceedings of the 2015 ACM SIGMOD international conference on management of data*, pages 1855–1870, 2015.
- [46] Anqi Qiu, Marilyn Albert, Laurent Younes, and Michael I Miller. Time sequence diffeomorphic metric mapping and parallel transport track time-dependent shape changes. *NeuroImage*, 45(1):S51–S60, 2009.
- [47] Danilo Rezende and Shakir Mohamed. Variational inference with normalizing flows. In *International conference on machine learning*, pages 1530–1538. PMLR, 2015.
- [48] Hadi Salman, Payman Yadollahpour, Tom Fletcher, and Kayhan Batmanghelich. Deep diffeomorphic normalizing flows. *arXiv preprint arXiv:1810.03256*, 2018.
- [49] Mike Schuster and Kuldeep K Paliwal. Bidirectional recurrent neural networks. *IEEE transactions on Signal Processing*, 45(11):2673–2681, 1997.
- [50] Han Lin Shang. A survey of functional principal component analysis. *ASTA Advances in Statistical Analysis*, 98:121–142, 2014.
- [51] Gota Shirato, Natalia Andrienko, and Gennady Andrienko. Identifying, exploring, and interpreting time series shapes in multivariate time intervals. *Visual Informatics*, 7(1):77–91, 2023.

- [52] Satya Narayan Shukla and Benjamin M Marlin. Multi-time attention networks for irregularly sampled time series. *arXiv preprint arXiv:2101.10318*, 2021.
- [53] Anuj Srivastava, Eric Klassen, Shantanu H Joshi, and Ian H Jermyn. Shape analysis of elastic curves in euclidean spaces. *IEEE transactions on pattern analysis and machine intelligence*, 33(7):1415–1428, 2010.
- [54] Sana Tonekaboni, Danny Eytan, and Anna Goldenberg. Unsupervised representation learning for time series with temporal neighborhood coding. *arXiv preprint arXiv:2106.00750*, 2021.
- [55] Patara Trirat, Yooju Shin, Junhyeok Kang, Youngeun Nam, Jihye Na, Minyoung Bae, Joeun Kim, Byunghyun Kim, and Jae-Gil Lee. Universal time-series representation learning: A survey. *arXiv preprint arXiv:2401.03717*, 2024.
- [56] Belinda Tzen and Maxim Raginsky. Neural stochastic differential equations: Deep latent gaussian models in the diffusion limit. *arXiv preprint arXiv:1905.09883*, 2019.
- [57] Marc Vaillant, Michael I Miller, Laurent Younes, and Alain Trouvé. Statistics on diffeomorphisms via tangent space representations. *NeuroImage*, 23:S161–S169, 2004.
- [58] Kai Wang, Youjin Zhao, Qingyu Xiong, Min Fan, Guotan Sun, Longkun Ma, Tong Liu, et al. Research on healthy anomaly detection model based on deep learning from multiple time-series physiological signals. *Scientific Programming*, 2016, 2016.
- [59] John Warmenhoven, Norma Bargary, Dominik Liebl, Andrew Harrison, Mark A Robinson, Edward Gunning, and Giles Hooker. Pca of waveforms and functional pca: A primer for biomechanics. *Journal of Biomechanics*, 116:110106, 2021.
- [60] Yuexuan Wu, Chao Huang, and Anuj Srivastava. Shape-based functional data analysis. *TEST*, 33(1):1–47, 2024.
- [61] Can Ye, BVK Vijaya Kumar, and Miguel Tavares Coimbra. Heartbeat classification using morphological and dynamic features of ecg signals. *IEEE Transactions on Biomedical Engineering*, 59(10):2930–2941, 2012.
- [62] Lexiang Ye and Eamonn Keogh. Time series shapelets: a new primitive for data mining. In *Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 947–956, 2009.
- [63] Qunqun Yu, Xiaosun Lu, and JS Marron. Principal nested spheres for time-warped functional data analysis. *Journal of Computational and Graphical Statistics*, 26(1):144–151, 2017.
- [64] Sergey Zagoruyko and Nikos Komodakis. Wide residual networks. *arXiv preprint arXiv:1605.07146*, 2016.
- [65] Juntang Zhuang, Tommy Tang, Yifan Ding, Sekhar C Tatikonda, Nicha Dvornek, Xenophon Papademetris, and James Duncan. Adabelief optimizer: Adapting stepsizes by the belief in observed gradients. *Advances in neural information processing systems*, 33:18795–18806, 2020.

A Societal impact

We believe that the paper has a positive societal impact for the following reasons:

- TS-LDDMM is an interpretable method for understanding inter-individual variability in biomedical datasets, potentially offering new insights in medicine.
- TS-LDDMM bridges the gap between the shape analysis community and the unsupervised representation learning (URL) community, fostering potential future collaborations between these fields.

However, the computational cost of the method may raise environmental concerns similar to those associated with deep learning [43]. Additionally, while TS-LDDMM has promising biomedical applications, it could also be misused for creating poison.

B Proofs

Denote by $G(s) \triangleq \{(t, s(t)) : t \in I\}$ the graph of a time series $s : I \rightarrow \mathbb{R}^d$ and $\phi.G(s) \triangleq \{\phi(t, s(t)) : t \in I\}$ the action of $\phi \in \mathcal{D}(\mathbb{R}^{d+1})$ on $G(s)$.

Theorem 4. *Let $s : J \rightarrow \mathbb{R}^d$ and $s_0 : I \rightarrow \mathbb{R}^d$ be two continuously differentiable time series with I, J two intervals of $\mathbb{R}$. There exist $f \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ and $\gamma \in \mathcal{D}(\mathbb{R})$ such that $\gamma(I) = J$ and $\Phi_f \in \mathcal{D}(\mathbb{R}^{d+1})$,*

$$G(s) = \Pi_{\gamma, f}.G(s_0), \quad \Pi_{\gamma, f} = \Psi_\gamma \circ \Phi_f.$$

Moreover, for any $\bar{f} \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ and $\bar{\gamma} \in \mathcal{D}(\mathbb{R})$, there exists a continuously differentiable time series $\bar{s}$ such that $G(\bar{s}) = \Pi_{\bar{\gamma}, \bar{f}}.G(s_0)$

Proof. Let $s : J \rightarrow \mathbb{R}^d$ and $s_0 : I \rightarrow \mathbb{R}^d$ be two continuously differentiable time series with $I = (a, b)$, $J = (\alpha, \beta)$ two intervals of $\mathbb{R}$. By setting $\gamma : t \in \mathbb{R} \mapsto (\beta - \alpha)(t - a)/(b - a) + \alpha \in \mathbb{R}$, we have $\gamma(I) = J$ and $\gamma \in \mathcal{D}(\mathbb{R})$. By defining $f : (t, x) \in \mathbb{R}^{d+1} \mapsto x - s_0(t) + s \circ \gamma(t)$, the map $\Phi_f \in \mathcal{D}(\mathbb{R}^{d+1})$, indeed, its inverse is $\Phi_f^{-1} : (t, x) \in \mathbb{R}^{d+1} \mapsto (t, x + s_0(t) - s(t))$ and is continuously differentiable. Moreover, we have $\Pi_{\gamma, f}.G(s_0) = \{(\gamma(t), s \circ \gamma(t)) : t \in I\} = G(s)$.

Let $\bar{f} \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$, $\bar{\gamma} \in \mathcal{D}(\mathbb{R})$ and $s_0 \in C^1(I, \mathbb{R}^d)$ with I an interval of $\mathbb{R}$. We have :

$$\begin{aligned} \Pi_{\gamma, f}.G(s_0) &= \{(\gamma(t), f(t, s_0(t))), t \in I\} \\ &= \{(t, f(\gamma^{-1}(t), s_0(\gamma^{-1}(t))), t \in \gamma(I)\}. \end{aligned} \quad (10)$$

By defining $\bar{s} : t \in \gamma(I) \rightarrow f(\gamma^{-1}(t), s_0(\gamma^{-1}(t)))$, we have $\bar{s} \in C^1(\gamma(I), \mathbb{R}^d)$ by composition of C^1 functions and $G(\bar{s}) = \Pi_{\gamma, f}.G(s_0)$ by (10), which concludes the proof. $\square$

Lemma 2. *If we denote by $\mathcal{V}$ the RKHS associated with the kernel K_G , then for any vector field v generated by (5) with v_0 satisfying (4), there exist $\gamma \in \mathcal{D}(\mathbb{R})$ and $f \in C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$ such that $\phi^v = \Psi_\gamma \circ \Phi_f$.*

Proof. Let v be a vector field generated by (5) with v_0 satisfying (4). We remark that the first coordinate of the velocity field v_τ denoted by v_τ^{time} only depends on the time variable t for any $\tau \in [0, 1]$. Thus, when computing the first coordinate of the deformation ϕ^v , denoted by γ , we integrate (1) with v_τ replaced by v_τ^{time} , thus γ is independent of the variable x . Moreover, $\gamma \in \mathcal{D}(\mathbb{R})$ since a Gaussian kernel induced an Hilbert space $\mathcal{V}$ satisfying $\|f\|_{\mathcal{V}} \leq \|f\|_\infty + \|df\|_\infty$ for any $f \in \mathcal{V}$ by [22, Theorem 9]. For the same reason, we have $\phi^v \in \mathcal{D}(\mathbb{R}^{d+1})$, and thus its last coordinates denoted by f belongs to $C^1(\mathbb{R}^{d+1}, \mathbb{R}^d)$, and by construction $\phi^v = \Psi_\gamma \circ \Phi_f$. $\square$

C Oriented varifold

In this section, we introduce the *oriented varifold* associated with curves. For further readings on curves and surfaces representation as varifolds, readers can refer to [30, 12]. We associate to

$\gamma \in C^1((a, b), \mathbb{R}^{d+1})$ an *oriented varifold* μ_γ , i.e. a distribution on the space $\mathbb{R}^{d+1} \times \mathbb{S}^d$ defined as follows, for any smooth test function $\omega : \mathbb{R}^{d+1} \times \mathbb{S}^d \rightarrow \mathbb{R}$,

$$\mathbb{E}_{Y \sim \mu_\gamma} [\omega(Y)] = \mu_\gamma(\omega) = \int_a^b \omega \left(\gamma(t), \frac{\dot{\gamma}(t)}{|\dot{\gamma}(t)|} \right) |\dot{\gamma}(t)| dt.$$

Denoting by W the space of smooth test function, we have that μ_γ belongs to its dual W^* . Thus, a distance on W^* is sufficient to set a distance on oriented varifolds associated to curve and thus on $C^1((a, b), \mathbb{R}^{d+1})$ by the identification $\gamma \rightarrow \mu_\gamma$. Remark that in (TS-LDDMM), γ should be the parametrization of a time series' graph $G(s)$, i.e. $\gamma : t \in I \rightarrow (t, s(t)) \in \mathbb{R}^{d+1}$ denoting by $s : I \rightarrow \mathbb{R}^d$ the time series. However, in practice, we work with discrete objects. That is why, we set W as an RKHS to use its representation theorem. More specifically [30, Proposition 2 & 4] encourages us to consider a kernel $k : (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2 \rightarrow \mathbb{R}$ such that there exist two positive and continuously differentiable kernels k_{pos} and k_{dir} , such that for any $(x, \vec{u}), (y, \vec{v}) \in (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2$

$$k((x, \vec{u}), (y, \vec{v})) = k_{\text{pos}}(x, y) k_{\text{dir}}(\vec{u}, \vec{v}),$$

with moreover $k_{\text{dir}} > 0$ and k_{pos} which admits an RKHS W_{pos} dense in the space of continuous function on $\mathbb{R}^{d+1}$ vanishing at infinite [10].

Given such a kernel $k : (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2 \rightarrow \mathbb{R}$ verifying [30, Proposition 2 & 4], we have that for any $(x, v) \in \mathbb{R}^{d+1} \times \mathbb{S}^d$, $\delta_{(x, \vec{v})}$ belongs to W^* as a distribution and that the dual metric $\langle \cdot, \cdot \rangle_{W^*}$ satisfies for any $(x_1, v_1), (x_2, v_2) \in (\mathbb{R}^{d+1} \times \mathbb{S}^d)^2$,

$$\langle \delta_{(x_1, \vec{v}_1)}, \delta_{(x_2, \vec{v}_2)} \rangle_{W^*} = k((x_1, \vec{v}_1), (x_2, \vec{v}_2)).$$

Thus, given two sets of triplets $X = (l_i, x_i, \vec{v}_i)_{i \in [T_0-1]} \in (\mathbb{R} \times \mathbb{R}^{d+1} \times \mathbb{S}^d)^{T_0-1}$, $Y = (l'_i, y_i, \vec{w}_i)_{i \in [T_1]} \in (\mathbb{R} \times \mathbb{R}^{d+1} \times \mathbb{S}^d)^{T_1-1}$ and denoting by

$$\mu_X = \sum_{i=1}^{T_0} l_i \delta_{(x_i, \vec{v}_i)}, \mu_Y = \sum_{i=1}^{T_1} l'_i \delta_{(y_i, \vec{w}_i)}, \quad (11)$$

we have,

$$|\mu_X - \mu_Y|_{W^*}^2 = \sum_{i,j=1}^{T_0-1} l_i k((x_i, \vec{v}_i), (x_j, \vec{v}_j)) l_j + \sum_{i,j=1}^{T_1-1} l'_i k((y_i, \vec{w}_i), (y_j, \vec{w}_j)) l'_j - 2 \sum_{i=1}^{T_0-1} \sum_{j=1}^{T_1-1} l_i k((x_i, \vec{v}_i), (y_j, \vec{w}_j)) l'_j.$$

Then, using the identification $X \rightarrow \mu_X, Y \rightarrow \mu_Y$, we can define a distance on sets of triplets as $d_{W^*,3}(X, Y) = |\mu_X - \mu_Y|_{W^*}^2$.

Now, we aim to discretize the oriented varifold μ_G related to a time series' graph $G(s)$ by using a set of triplets. This is carried out by using a discretized version of $G(s)$, i.e. $\tilde{G} = (g_i = (t_i, s(t_i)))_{i \in [T]} \in (\mathbb{R}^{d+1})^T$, in the following way: For any $i \in [T-1]$, denoting the center and length of the i^{th} segment $[g_i, g_{i+1}]$ by $c_i = (g_i + g_{i+1})/2$, $l_i = \|g_{i+1} - g_i\|$, and the unit norm vector of direction $\vec{g_i g_{i+1}}$ by $\vec{v}_i = (g_{i+1} - g_i)/l_i$, we define the set of triplets $X(\tilde{G}) = (l_i, c_i, \vec{v}_i)_{i \in [T-1]}$ and its related oriented varifold $\mu_{X(\tilde{G})} = \sum_{i=1}^{T-1} l_i \delta_{c_i, \vec{v}_i}$ as in (11). This is a valid discretization of the oriented varifold μ_G according to [30, Proposition 1]: $\mu_{X(\tilde{G})}$ converges towards μ_G as the size of the discretization mesh $\sup_{i \in [T-1]} |t_{i+1} - t_i|$ converges to 0.

Finally, we define a distance on discretized time series' graphs $\tilde{G}_1, \tilde{G}_2$ as $d_{W^*,3}(\tilde{G}_1, \tilde{G}_2) = d_{W^*,3}(X(\tilde{G}_1), X(\tilde{G}_2))$.

C.1 Varifold kernels

Denote the one-dimensional Gaussian kernel by $K_\sigma^{(a)}(x, y) = \exp(-|x - y|^2/\sigma)$ for any $(x, y) \in (\mathbb{R}^a)^2$, $a \in \mathbb{N}$ and $\sigma > 0$. In the implementation, we use the following kernels, for any $((t_1, x_1), (t_2, x_2)) \in (\mathbb{R}^{d+1})^2, ((w_1, v_1), (w_2, v_2)) \in (\mathbb{S}^d)^2$,

$$k_{\text{pos}}(x, y) = K_{\sigma_{\text{pos},t}}^{(1)}(t_1, t_2) K_{\sigma_{\text{pos},x}}^{(d)}(x_1, x_2), \quad k_{\text{pos}}(x, y) = K_{\sigma_{\text{dir},t}}^{(1)}(w_1, w_2) K_{\sigma_{\text{dir},x}}^{(d)}(v_1, v_2),$$

where $\sigma_{\text{pos},t}, \sigma_{\text{pos},x}, \sigma_{\text{dir},t}, \sigma_{\text{dir},x} > 0$ are hyperparameters. In practice, we select $\sigma_{\text{pos},x} \approx \sigma_{\text{dir},x} \approx 1$ when the time series are centered and normalized. Otherwise we select $\sigma_{\text{pos},x} \approx \sigma_{\text{dir},x} \approx \bar{\sigma}_s$ with $\bar{\sigma}_s$ the average standard deviation of the time series. We choose $\sigma_{\text{pos},t} \approx \sigma_{\text{dir},t} = m f_e$ with f_e the sampling frequency of the time series and $m \in [5]$ an integer depending on the time change between the starting and the target time series graph. The more significant the time change, the higher m should be. The intuition comes from the fact that the width $\sigma_{\text{pos},t}, \sigma_{\text{dir},t}$ rules the time windows used to perform the comparison, and $\sigma_{\text{pos},x}, \sigma_{\text{dir},x}$ affects the space window. The size of the windows should be selected depending on the variations in the data.

D Tuning the hyperparameters of the TS-LDDMM velocity field kernel

The parameter $\sigma_{T,0}$ should be chosen *large* compared the sampling frequency f_e and compared to average standard deviation $\bar{\sigma}_s$ of the time series, e.g $\sigma_{T,0} = 100$ as $\bar{\sigma}_s \approx f_e \approx 1$. It makes the time transformation smoother. If $\sigma_{T,0}$ is too small, for instance, $\sigma_{T,0} = f_e$, the effect of the time deformation is too localized, and there are not enough samples to make it visible.

The parameter $\sigma_{T,1}$ should be of the same order as f_e : two different points in time can have various space transformations. σ_x should be of the same order of $\bar{\sigma}_s$: two points with a big difference regarding space compared to $\bar{\sigma}_s$ can have very different space transformations.

We take $c_0 \approx 10c_1$, we want to encourage time transformation before space transformation. We take $(c_0, c_1) = (1, 0.1)$ in all experiments.

E Experimental settings

All experiments were performed on a Debian 6.1.69-1 server with NVIDIA RTX A2000 12GB GPU, Intel(R) Xeon(R) Gold 5220R CPU @ 2.20GHz, and 250 GB of RAM. The source code is available on Github⁵.

E.1 Optimization details of TS-LDDMM & LDDMM

We implemented TS-LDDMM in Python with the JAX library⁶.

Initialization. As initialization of (8), all momentum parameters are set to 0, and the initial graph of reference is picked from the dataset such that its length is equal to the median length observed in the dataset.

Gradient descent. The chosen gradient descent method is "adabelief" [65] implemented in the OPTAX library⁷. The gradient descent has two main parameters: the number of steps (nb_steps) and the maximum stepsize value (η_M). The stepsize has a scheduling scheme:

- Warmup period on $0.1 \times \text{nb_steps}$ steps: the stepsize increases linearly from 0 to η_M . The goal is to learn progressively the parameters. If the step size is too large at the start, smaller steps at the end cannot make up for the mistakes made at the beginning.
- Fine tuning periode on $0.9 \times \text{nb_steps}$: the stepsize decreases from η_M to 0 with a cosine decay implemented in the OPTAX scheduler, i.e. the decreasing factor as the form $0.5(1 + \cos(\pi t/T))$.

By default, we set nb_steps to 400 and η_M to 0.1.

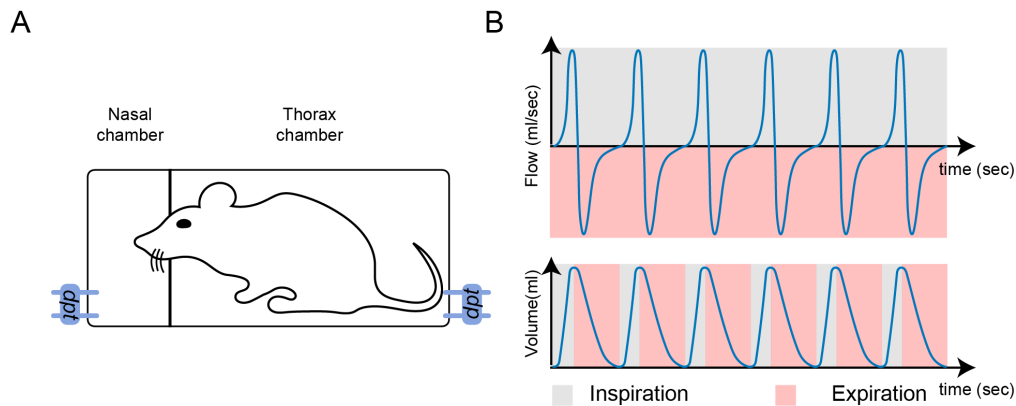


Figure 6: A: Illustration of a double-chamber plethysmograph. The term *dpt* stands for differential pressure transducer which measures the pressure in each compartment, the pressure then being converted to flow. B: Nasal airflow (top) and lung volume (bottom). During inspiration, airflow is positive (grey) and during expiration, airflow is negative (pink).

F Datasets

F.1 Mouse respiratory cycle dataset

Ventilation is a simple physiological function that ensures a vital supply of oxygen and the elimination of CO₂. Acetylcholine (ACh) is a neurotransmitter that plays an important role in muscular activity, notably for breathing. Indeed, muscle contraction information passes from the brain to the muscle through the nervous system. AChs are located in synapses of the nervous system (central and peripheral) and skeletal muscles. They ensure the information transmission from nerve to nerve. However, the transmission cannot end without the hydrolysis of ACh by the enzyme Acetylcholinesterase (AChE), allowing nerves to return to their resting state. Inhibition of (AChE) with, for instance, nerve gas, pesticide, or drug intoxication leads to respiratory arrests.

The dataset comes from the experiment [41], where they studied the consequences of partial deficits in AChE and AChE inhibition on mice respiration. AChE inhibition was induced with an irritant molecule called physostigmine (an AChE inhibitor). Mice nasal airflows were sampled at 2000Hz with a Double Chamber plethysmograph [28], as depicted in Figure 6-A). The flow is expressed in $ml.s^{-1}$; it has a positive value during inspiration and a negative value expiration Figure 6-B). Among the mice population, we selected 7 control mice (**wt**) and 7 ColQ mice (**colq**), which do not have AChE anchoring in muscles and some tissues. As described in [41], mice experiments were as follows:

1. The mouse is placed in a DCP for 15 or 20 min to serve as an internal control.
2. The mouse is removed from the DCP and injected with physostigmine.
3. The mouse is placed back into the DCP, and its nasal flow is recorded for 35 or 40 min.

Respiratory cycles were extracted following procedure [21]. We removed respiratory cycles whose duration exceeds 1 second; the average respiratory cycle duration is 300 ms. We randomly sampled 10 respiratory cycles per minute and mouse. It leads to a dataset of 12,732 (time, genotype)-annotated respiratory cycles.

⁵<https://github.com/thibaut-germain/TSLDDMM>

⁶<https://github.com/google/jax>

⁷<https://optax.readthedocs.io/en/latest/>

F.2 Shape-based UCR/UEA time series classification datasets

We selected 15 shape-based datasets (7 univariates and 8 multivariates) from the from the University of East Anglia (UEA) and the University of California Riverside (UCR) Time Series Classification Repository⁸ [15, 3]. All datasets were downloaded with the python package aeon⁹. Essential datasets information are summarized in Table 1 and further can be found in [15, 3].

Table 1: UCR/UEA shape-based time series datasets for classification.

	Dataset	Size	Length	Number of classes	Number of dimensions	Type
Univariate	ArrowHead	211	251	3	1	IMAGE
	BME	180	128	3	1	SIMULATED
	ECG200	200	96	2	1	ECG
	FacesUCR	2250	131	14	1	IMAGE
	GunPoint	200	150	2	1	MOTION
	PhalangesOutlinesCorrect	2658	80	2	1	IMAGE
	Trace	200	275	4	1	SENSOR
Multivariate	ArticularyWordRecognition	575	144	25	9	SENSOR
	Cricket	180	1197	12	6	MOTION
	ERing	60	65	6	4	SENSOR
	Handwriting	1000	152	26	3	MOTION
	Libras	360	45	15	2	VIDEO
	NATOPS	360	51	6	24	MOTION
	RacketSports	303	30	4	6	SENSOR
	UWaveGestureLibrary	240	315	8	3	SENSOR

G Appendix for experiment: TS-LDDMM representation identifiability

In this experiment, we evaluate the ability of TS-LDDMM to retrieve the parameter v_0^* that encodes the deformation $\varphi^{\{v_0^*\}}$ acting on a time series graph G by solving the geodesic shooting problem (6) between G and $\varphi^{\{v_0^*\}}.G$. Parameter identifiability is an important property for subsequent statistical analysis. Results show that TS-LDDMM representations are identifiable or weakly identifiable depending on the velocity field kernel K_G specification.

G.1 Settings

This experiment only involves the TS-LDDMM method in two different settings:

- **The velocity field kernel K_G is well-specified:** The velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 100, 1, 1)$, the varifold loss kernels (k_{pos}, k_{dir}) are set to $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x}) = (2, 1, 2, 0.6)$, and the optimizer has 400 steps with a maximum stepsize η_M of 0.05.
- **The velocity field kernel K_G is misspecified:** The velocity field kernel K_G is set with $(c_0, c_1, \sigma_{T,1}) = (1, 0.1, 1)$, $\sigma_{T,0}$ ranging in $(1, 5, 10, 50, 100, 200, 300)$, and σ_x ranging in $(0.1, 1, 10, 100)$. The varifold loss kernels (k_{pos}, k_{dir}) are set to $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x}) = (2, 1, 2, 0.6)$, and the optimizer has 400 steps with a maximum stepsize η_M of 0.05.

provided that the hyperparameters and the reference graph are wisely selected, i.e., the parameter v_0^* generating a deformation $\varphi^{\{v_0^*\}}$ of a time series graph G can be estimated from the data $G, \varphi^{\{v_0^*\}}.G$ by solving the geodesic shooting problem (6).

The velocity field kernel K_G is well specified. First, we show the model identifiability when the kernel K_G is well specified: the estimated parameter is a good approximation of the generating parameter when the generation and the estimation procedure use the same hyperparameters for the RKHS kernel K_G . All the hyperparameter values for generation and estimation are given in Appendix G.1.

⁸<https://timeseriesclassification.com>

⁹<https://www.aeon-toolkit.org/en/stable/>

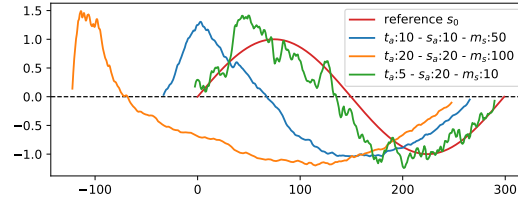


Figure 7: Plots of $\varphi^{\{v_0(\alpha^*, X)\}}.X$ for different values of α^* according to its sampling parameter t_a, s_a, m_s , taking $X = G(s_0)$ with $s_0 : k \in [300] \rightarrow \sin(2\pi k/300)$.

Table 2: Values of $\mathcal{L}(\varphi^{\{v_0(\alpha^*, X)\}}.X, \varphi^{\{\hat{v}_0\}}.X)$ as α^* is sampled according to $\text{Gen}(10, 10, 50)$ and $\hat{v}_0$ is estimated using K_G with varying parameters $\sigma_{T,1}, \sigma_x$.

$\sigma_{T,0} \backslash \sigma_x$	1	10	50	100	200	300
0.1	2e+0	3e-4	1e-5	4e-6	7e-4	4e-3
1	4e-2	1e-4	1e-5	4e-6	7e-4	4e-3
100	4e-2	2e-4	1e-5	4e-6	7e-4	4e-3

We fix the initial control points as $X = (x_k = (k, \sin(2\pi k/300)))_{k \in [300]}$. Given $m_s \in \mathbb{N}_{>0}$ and $t_a, s_a > 0$, we randomly generate initial momentums $\alpha^* = (\alpha_k^*)_{k \in [n_0]}$ with the following sampling, called $\text{Gen}(m_s, t_a, s_a)$: For any $k \in [n_0]$, α_k' is sampled according to a Gaussian normal distribution $\mathcal{N}(0_{d+1}, I_{d+1})$. Then, $(\alpha_k')_{k \in [n_0]}$ is regularized by a rolling average of size m_s , we get $\bar{\alpha}' = (\bar{\alpha}'_k)_{k \in [n_0]}$. Finally, we normalize $\bar{\alpha}'$ to derive α^* such that $|([\alpha_k^*]_t)_{k \in [n_0]}| = t_{\text{amp}}$ and $|([\alpha_k^*]_s)_{k \in [n_0]}| = s_{\text{amp}}$ for any $k \in [n_0]$, denoting by $[\alpha_k^*]_t, [\alpha_k^*]_s$ the time and space coordinates of α_k^* respectively. Note that the regularizing step $(\alpha_k')_{k \in [n_0]} \rightarrow \bar{\alpha}'$ is necessary to obtain realistic deformations which take into account the regularity induced by the RKHS V .

Then, using $v_0(\alpha^*, X)$ as defined in (4) with initial momentums α^* and control points X , we apply the induced deformation $\varphi^{\{v_0\}}$ by (5) to X and obtain $\varphi^{\{v_0\}}.X$. Finally, we solve (6) to recover an estimation $\hat{\alpha}$ of α^* and report the average relative error (ARE) $|v_0(\hat{\alpha}, X) - v_0(\alpha^*, X)|_V / |v_0(\alpha^*, X)|_V$ on 50 repetitions. This procedure is performed for any $m_s, t_a, s_a \in \{10, 50, 100\} \times \{5, 10, 15, 20\}^2$. Mean, standard deviation, and maximum of the ARE on all these hyperparameters choices are respectively **0.10, 0.03, 0.17**. Therefore, the estimation procedure (6) offers a good approximation of the true parameter when the kernel K_G is well specified. We observe that the estimation is difficult when $t_a \ll s_a$ because the time series can be very noisy as illustrated in Figure 7: this impacts the Varifold loss which is sensitive to tangents.

The velocity field kernel K_G is misspecified. We demonstrate a weak identifiability when the kernel K_G is misspecified: we can reconstruct the graph time series' after deformations even if the hyperparameters of K_G are different during the generation and the estimation. The hyperparameters of K_G during generation are $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 100, 1, 1)$ and we fix $\sigma_{T,1}, c_0, c_1 = (1, 1, 0.1)$ for K_G during estimation. We aim to understand the impact of $\sigma_{T,1}, \sigma_x$ on the reconstruction since they are encoding the smoothness of the transformation according to time and space.

For any choice of the hyperparameters $\sigma_{T,1}, \sigma_x \in \{1, 10, 50, 100, 200, 300\} \times \{0.1, 1, 100\}$ related to K_G in the estimation, we average $\mathcal{L}(\varphi^{\{v_0(\alpha^*, X)\}}.X, \varphi^{\{\hat{v}_0\}}.X)$ on 50 repetitions when α^* is sampled according to $\text{Gen}(10, 10, 50)$ and $\hat{v}_0 = v_0(\hat{\alpha}, X)$ denoting by $\hat{\alpha}$ the result of the minimization (6). We observe in Table 2 that the reconstruction is almost perfect except in the case when $\sigma_{T,0} = 1$ during estimation, while $\sigma_{T,0} = 100$ during generation. Compared to $\sigma_{T,0}$, σ_x has nearly no impact on the reconstruction. In Appendix C.1-D, we propose guidelines to drive future hyperparameters tuning and further discussions related to $\sigma_{T,1}, c_0, c_1$.

H Appendix for experiment: Robustness to irregular sampling

This experiment is inspired by [43] where the authors perform an extensive comparison of Neural Ordinary Differential Equations (Neural ODEs) methods [31]. We assess the classification performances of several methods under regular sampling (0% missing rate) and three irregular sampling regimes on 15 shape-based datasets (7 univariate & 8 multivariate). Methods and training strategy are taken from its associated Github¹⁰ and described in what follows. We conclude with the results, which show that our method, TS-LDDMM, outperforms all methods for sampling regimes with missing rates: 0%, 30%, and 50%.

H.1 Benchmark methods

In related work, we give an overview of Neural ODEs methods and their relation with TS-LDDMM.

- RNN-based methods: Baseline recurrent neural networks including RNN [38], LSTM [27], and GRU [14].
- Attention-based methods: Multi-Time Attention Networks (MTAN) [52] and Multi-Integration Attention Module (MIAM) [35]. Both handle multivariate time series irregularly sampled with attention mechanisms.
- Neural ODEs: ODE-LSTM [34] a form of Neural-ODEs used to learn continuous latent representations.
- Neural SDEs: Neural SDE [36] and Neural LNSDE [43] have been proposed to model randomness in time-series using drift and diffusion terms as an extension of Neural-ODEs.
- Shape-Analysis methods: TS-LDDMM (ours) and LDDMM [23]. From shape analysis, both methods learn representations by solving ODEs parametrized with Kernels. While both methods handle multivariate signals irregularly sampled, TS-LDDMM is specifically designed for time series.

H.2 Model architecture

Neural ODEs methods As depicted in [43], any Neural ODEs layer in Appendix H.1 is followed by an MLP with two fully connected layers with ReLU activations. The risk of overfitting and the model regularization are handled with a dropout rate of 10% and an early-stopping mechanism, ceasing the training when the validation loss does not improve for 10 successive epochs.

For each method and dataset, the learning rate, the hidden vector dimensions, and the number of layers are optimized to minimize the CrossEntropy loss on a validation set using the Ray¹¹ Python library. The learning rate varies from 10^{-4} to 10^{-1} using log uniform search, the hidden vector dimension ranges from 16, 32, 64, 128 using grid search, and the number of layers ranges from 1, 2, 3, 4 using grid search. The batch size was selected from 16, 32, 64, 128 according to the size of the dataset. All methods were trained for 100 epochs, and the best method was selected based on the lowest validation loss.

TS-LDDMM and LDDMM Representations learned with TS-LDDMM or LDDMM are fed to a Support Vector Classifier (SVC) from `scikit-learn`¹². All SVC's hyperparameters are set to default except the regularization term C, which is set through grid search on a validation set with the macro f1-score¹³.

To learn TS-LDDMM (resp. LDDMM) representations, the velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 0.33\bar{l}, 1, n_d)$, (resp. $(\sigma_T, \sigma_x) = (0.33\bar{l}, n_d)$) where $\bar{l}$ is the average time series length and n_d the number of dimensions. For both methods and all datasets, the varifold loss kernels (k_{pos}, k_{dir}) are identical and set to $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x}) = (2, n_d, 2, n_d)$. For TS-LDDMM (resp. LDDMM), the optimizer is set with 400 epochs (resp. 400) and a maximum learning rate $\eta_M = 0.1$ (resp. $\eta_M = 0.01$). In all cases, the initial reference graph is selected in the dataset as a time series with the median length.

¹⁰<https://github.com/yongkyung-oh/Stable-Neural-SDEs>

¹¹<https://github.com/ray-project/ray>

¹²<https://scikit-learn.org/stable/>

¹³https://scikit-learn.org/stable/modules/generated/sklearn.metrics.f1_score.html

H.3 Protocol

In this experiment, we investigate the robustness to missing samples and the classification performance of TS-LDDMM compared to Neural ODEs on 15 datasets described in Appendix F.2. For fairness between methods of different architectures, the evaluation protocol on each dataset and method is as follows:

1. Spilt the dataset in train 75%, validation 15%, and test 15%.
2. Tune hyperparameters with train and validation sets and a missing rate of 0%.
3. For each missing rate in [0%,30%,50%,70%]
 - Remove samples in time series in the train and test sets according to the missing rate and the drop procedure described in [31].
 - Train the model on the train set
 - Evaluate the macro f1-score on the test set

H.4 Results

In this experiment, we investigate the robustness to missing samples and the classification performance of TS-LDDMM representations. We compare TS-LDDMM with LDDMM and 8 neural ODEs networks. Performances are evaluated in terms of average macro f1-score and rank on four different regimes of missing rate 0%,30%,50%, and 70%. Results are aggregated in Table 3.

On three out of four regimes (0%,30%, and 50%) TS-LDDMM classifier is the best performer in terms of f1-score and rank. For missing rates of 0% and 30%, the score increases by 10% compared to the second-best performer, LDDMM. However, LDDMM is not the second-best performer in rank (Neural LNSDE), showing its sensitivity to parameterization, unlike TS-LDDMM, which remains consistent. Performances of Neural LNSDE remain constant with the increase of the missing rate as observed in [43], and it becomes the best performer for missing rate 70%. The decrease in TS-LDDMM performances with the increasing missing rate is due to the varifold loss, which poorly approximates the time series shape. Other losses might be more relevant for high missing rates.

Overall, TS-LDDMM is a relevant and consistent shape-based representation for irregularly sampled multivariate time series for missing rates up to 50% .

Table 3: Comparison of average macro f1-score and rank as the sample dropping rate increases. **First & second** best performers. TS-LDDMM is the best performer on three out of four regimes.

Methods	Regular		30 % dropped		50 % dropped		70 % dropped	
	F1-score	Rank	F1-score	Rank	F1-score	Rank	F1-score	Rank
RNN (1999)	0.64 ± 0.21	6.2	0.53 ± 0.23	6.6	0.48 ± 0.21	7.2	0.44 ± 0.21	6.07
LSTM (1997)	0.61 ± 0.29	6.0	0.57 ± 0.29	6.27	0.53 ± 0.25	6.07	0.51 ± 0.29	5.27
GRU (2014)	0.71 ± 0.26	4.2	0.68 ± 0.28	4.27	0.66 ± 0.28	3.73	0.59 ± 0.28	<u>3.67</u>
MTAN (2021)	0.59 ± 0.28	7.13	0.58 ± 0.28	5.8	0.54 ± 0.29	5.33	0.51 ± 0.28	5.0
MIAM (2022)	0.48 ± 0.35	6.93	0.42 ± 0.33	8.27	0.47 ± 0.31	6.93	0.35 ± 0.31	7.6
ODE-LSTM (2020)	0.63 ± 0.24	6.0	0.57 ± 0.25	6.53	0.51 ± 0.24	7.27	0.45 ± 0.23	6.73
Neural SDE (2019)	0.48 ± 0.28	7.67	0.47 ± 0.26	7.47	0.45 ± 0.27	7.13	0.45 ± 0.25	6.0
Neural LNSDE (2024)	0.7 ± 0.27	<u>3.87</u>	0.68 ± 0.29	<u>4.0</u>	0.67 ± 0.25	<u>3.53</u>	0.66 ± 0.23	2.47
LDDMM (2008)	<u>0.72 ± 0.2</u>	4.53	<u>0.7 ± 0.21</u>	4.2	0.57 ± 0.25	5.0	0.4 ± 0.25	7.13
TS-LDDMM (ours)	0.83 ± 0.18	2.93	0.8 ± 0.18	2.07	0.7 ± 0.26	3.33	0.51 ± 0.27	5.67

I Appendix for experiment: Classification benchmark on regularly sampled datasets

In this section, we compare the classification performances of TS-LDDMM with other methods from shape analysis on 15 shape-based datasets of time series regularly sampled. TS-LDDMM outperforms other methods on 12 out of 15, highlighting its relevance for shape analysis when dealing with time series.

I.1 Benchmark methods

- SRV-based method: we include TCLR [26] a logistic regression on the tangent space of the Frechet mean with Square Root Velocity (SRV representation). We also include Shape-FPCA [60] that encodes both the time series and its time parameterization.
- LDDMM-Based : TS-LDDMM (ours) and LDDMM [23]. Both methods learn representations by solving ODEs parametrized with Kernels. While both methods handle multivariate signals, TS-LDDMM is specifically designed for time series.

I.2 Model settings

TCLR & Shape-FPCA Shape-FPCA is available in the Python library FDASRSF¹⁴. Once the shape-FPCA representations are learned, they are fed to an SVC from `scikit-learn`. FDASRSF provides SRV representation methods that we combined with a logistic regression from `scikit-learn` to implement TCLR. For both methods, the number of steps to learn the Frechet mean is set to 50, and the regularization hyperparameter C is set through grid search on a validation set with the macro f1-score. Other parameters are set to default.

TS-LDDMM & LDDMM Representations learned with TS-LDDMM or LDDMM are fed to an SVC from `scikit-learn`. All SVC's hyperparameters are set to default except the regularization term C, which is set through grid search on a validation set with the macro f1-score.

To learn TS-LDDMM (resp. LDDMM) representations, the velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 0.33\bar{l}, 1, n_d)$, (resp. $(\sigma_T, \sigma_x) = (0.33\bar{l}, n_d)$) where $\bar{l}$ is the average time series length and n_d the number of dimensions. For both methods and all datasets, the varifold loss kernels (k_{pos}, k_{dir}) are identical and set to $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x}) = (2, n_d, 2, n_d)$. For TS-LDDMM (resp. LDDMM), the optimizer is set with 400 epochs (resp. 400) and a maximum learning rate $\eta_M = 0.1$ (resp. $\eta_M = 0.01$). In all cases, the initial reference graph is selected in the dataset as a time series with the median length.

I.3 Protocol

For each dataset and method, the evaluation protocol is a simple train, validation test with hyperparameter tuning:

1. Split The dataset in train 75%, validation 15%, and test 15%.
2. Training and hyperparameters tuning with train and validation sets
3. Evaluate the macro f1-score on the test set

I.4 Results

In this experiment, we investigate the classification performances of several methods from shape analysis on 15 shape-based time series datasets (7 univariate and 8 multivariate). The performances are evaluated in terms of macro f1-score. Results are aggregated in Table 4.

The TS-LDDMM-based classifier outperforms other methods on 12 out of 15 datasets. TCLR is the second-best performer on univariate datasets; however, its current implementation with FDASRSF does not extend to the multivariate case, which limits usage. LDDMM performances are lower than TCLR, and Shape-FPCA is the worst performer.

Overall, TS-LDDMM representations are well suited for shape-based time series classification, and its extension to multivariate irregularly sampled time series makes it a relevant option for time series shape analysis.

J Appendix for the experiment: Noise sensitivity

This experiment evaluates the influence of noise on the learning of the reference sequence for TS-LDDMM and SRVF Kacher-mean, a subroutine of Shape-FPCA [60].

¹⁴<https://fdasrsf-python.readthedocs.io/en/latest/>

Table 4: F1-score comparison between methods from shape analysis on 15 datasets. **First** and second best performers.

	Dataset	Shape-FPCA (2024)	TCLR (2024)	LDDMM (2008)	TS-LDDMM (ours)
Univariate	ArrowHead	0.18	0.75	0.84	0.91
	BME	0.16	<u>1.00</u>	0.82	1.00
	ECG200	0.40	0.67	0.81	0.79
	FacesUCR	0.08	<u>0.73</u>	0.69	0.86
	GunPoint	0.93	<u>0.97</u>	0.83	1.00
	PhalangesOutlinesCorrect	0.39	0.63	<u>0.53</u>	0.52
	Trace	0.55	<u>1.00</u>	0.46	1.00
Multivariate	ArticularyWordRecognition	–	–	<u>0.98</u>	1.00
	Cricket	–	–	<u>0.77</u>	0.93
	ERing	–	–	<u>0.95</u>	0.98
	Handwriting	–	–	<u>0.22</u>	0.44
	Libras	–	–	<u>0.56</u>	0.60
	NATOPS	–	–	<u>0.82</u>	0.82
	RacketSports	–	–	0.83	0.79
	UWaveGestureLibrary	–	–	<u>0.72</u>	0.81

J.1 Protocol

The dataset includes 100 sine waves with randomly generated time parametrization by following the procedure described in Appendix G.1 with Gen(50, 1, 0) and uniformly resampled. The dataset has been altered under four scenarios with an additive Gaussian noise centered and with standard deviation $\sigma_\epsilon \in \{0, 0.05, 0.1, 0.2\}$. The referent sequence is learned for each scenario, and the L_2 -norm error between the exact and the learned barycenter is computed.

J.2 Method settings

For SRVF, the number of steps to learn the Kacher-mean is set to 20. Regarding TS-LDDMM, the velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 65, 1, 1)$, and the varifold loss kernels (k_{pos}, k_{dir}) are set to $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x}) = (5, 1, 1, 1)$. For TS-LDDMM, the optimizer is set with 400 epochs and a maximum learning rate $\eta_M = 0.1$. In all cases, the initial reference graph is selected in the dataset as a time series with the median length.

J.3 Results

Figure 8 illustrates the results. Noise level affects the learning reference graph in both cases, as depicted by the increasing error and the illustrations. However, the overall sine wave shape is better preserved by TS-LDDMM compared to SRVF Kacher-mean, for which the sine wave amplitude decreases as the noise increases. In addition, for TS-LDDMM, the regularity of the reference graph can be controlled by penalizing the norm of the velocity fields in the loss function. Further work on penalization will be conducted to handle noisy data better.

K Appendix for experiment: Analysis of respiratory behavior in mice

K.1 Settings

This experiment involves TS-LDDMM, LDDMM [23] and Shape-FPCA [60] methods. Two scenarios are investigated: before drug exposure and before/after drug exposure. All methods are investigated on both scenarios.

TS-LDDMM parameters.

- **Before exposure:** The velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 150, 1, 2)$. The varifold loss is the sum of three varifolds to capture shapes variations at different scales with parameters: (Varifold 1, Varifold 2, Varifold 3): $((5, 2, 5, 1), (2, 1, 2, 0.6), (1, 0.6, 1, 0.6))$ and the mapper $(\sigma_{pos,t}, \sigma_{pos,t}, \sigma_{dir,t}, \sigma_{dir,x})$. The optimizer has 800 steps with a maximum stepsize η_M of 0.3.
- **Before/after exposure:** The velocity field kernel K_G is set to $(c_0, c_1, \sigma_{T,0}, \sigma_{T,1}, \sigma_x) = (1, 0.1, 220, 1, 2)$. The varifold loss is the sum of four varifolds to capture shapes

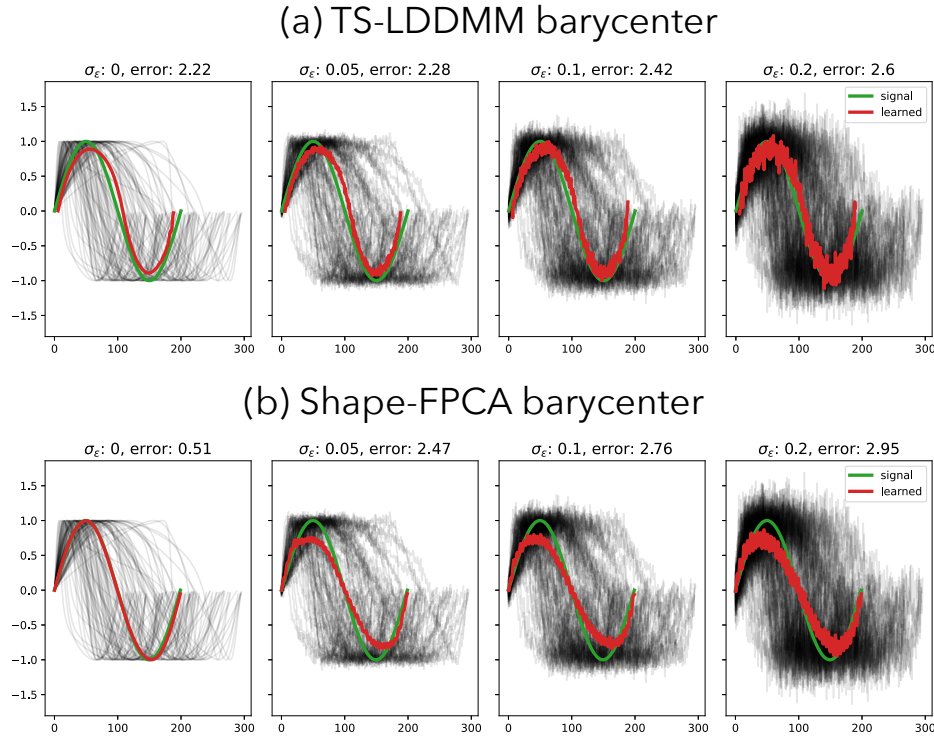


Figure 8: Illustration of the learned barycenter (red) compared to the exact barycenter (green) for both TS-LDDMM (a) and Shape-FPCA (b). The computation has been done for different level of noise $\epsilon \sim \mathcal{N}(0, \sigma_\epsilon)$ with $\sigma_\epsilon \in \{0, 0.05, 0.1, 0.2\}$.

variations at different scales with parameters: (Varifold 1, Varifold 2, Varifold 3, Varifold 4): $((30, 2, 30, 1), (5, 2, 5, 1), (2, 1, 2, 0.6), (1, 0.1, 1, 0.1))$ and the mapper $(\sigma_{\text{pos},t}, \sigma_{\text{pos},t}, \sigma_{\text{dir},t}, \sigma_{\text{dir},x})$. The optimizer has 800 steps with a maximum stepsize η_M of 0.3.

LDDMM parameters. Note that varifold losses are unchanged between TS-LDDMM and LDDMM. Compared to TS-LDDMM, the convergence of LDDMM is more sensitive to the maximum stepsize η_m , which must remain small for LDDMM to guarantee the convergence.

- **Before exposure:** The velocity field kernel K_G is an anisotropic Gaussian kernel with parameters $\sigma_T = 150$ for the time dimension and $\sigma_x = 2$ for space dimensions. The varifold loss is the sum of three varifolds to capture shapes variations at different scales with parameters: (Varifold 1, Varifold 2, Varifold 3): $((5, 2, 5, 1), (2, 1, 2, 0.6), (1, 0.6, 1, 0.6))$ and the mapper $(\sigma_{\text{pos},t}, \sigma_{\text{pos},t}, \sigma_{\text{dir},t}, \sigma_{\text{dir},x})$. The optimizer has 800 steps with a maximum stepsize η_M of 0.01.
- **Before/after exposure:** The velocity field kernel K_G is an anisotropic Gaussian kernel with parameters $\sigma_T = 220$ for the time dimension and $\sigma_x = 2$ for space dimensions. The varifold loss is the sum of four varifolds to capture shapes variations at different scales with parameters: (Varifold 1, Varifold 2, Varifold 3, Varifold 4): $((30, 2, 30, 1), (5, 2, 5, 1), (2, 1, 2, 0.6), (1, 0.1, 1, 0.1))$ and the mapper $(\sigma_{\text{pos},t}, \sigma_{\text{pos},t}, \sigma_{\text{dir},t}, \sigma_{\text{dir},x})$. The optimizer has 800 steps with a maximum stepsize η_M of 0.01.

Shape-FPCA parameters. For both scenarios, respiratory cycles are linearly interpolated and resampled to 200 points, and the length of the original time interval is kept. The computation of the Kacher-mean is done in a maximum of 50 iterations, and srv representations of the realigned time series and time parametrization are concatenated with cycle durations. When concatenating these vectors, the choice of amplitude factors is made to minimize the reconstruction error from the

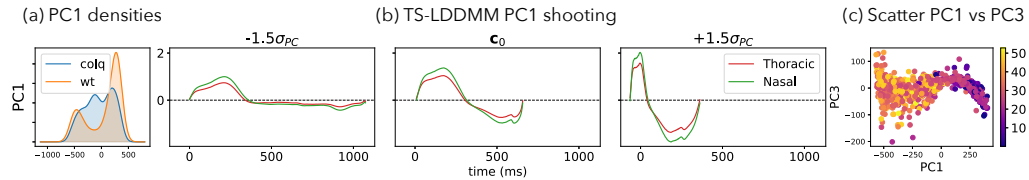


Figure 9: Analysis of the first Principal Component (PC1) related to mice ventilation before and after exposure with TS-LDDMM representations. **(a)** displays PC densities per mice genotype, **(b)** illustrates deformations of the reference respiratory cycle c_0 along PC1, and **(c)** displays all respiratory cycles with respect to time in PC1 and PC3 coordinates

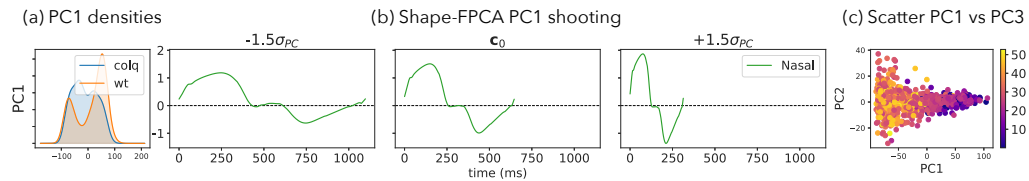


Figure 10: Analysis of the first Principal Component (PC1) related to mice ventilation before and after exposure with shape-FPCA representations. **(a)** displays PC densities per mice genotype, **(b)** illustrates deformations of the reference respiratory cycle c_0 along PC1, and **(c)** displays all respiratory cycles with respect to time in PC1 and PC2 coordinates

principal components analysis by following the procedure described in [60]. Shape-FPCA does not handle multivariate data, and we only kept the nasal airflow for this method.

K.2 Addiotinal results

Figure 9 presents results for TS-LDDMM and Figure 10 presents results for Shape-FPCA. The main components look similar. However, a subtle difference, yet important, can be noticed. With Shape-FPCA, the deformation tends to be a uniform time scaling, whereas, with TS-LDDMM, the time dilatation mainly occurs during the pause between inspiration and expiration. Qualitatively, this last deformation fits the physiological phenomenon: Mice's muscles cannot relax after exposure to the irritant molecule, leading to pauses between inspiration and expiration [41]. Qualitatively, contrary to Shape-FPCA, which manages to represent the main phenomena in the data, the deformations of TS-LDDMM capture subtle physiological behaviors essential for understanding the phenomenon at hand.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: Each claim in the introduction is referring to the part where it is tackled.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We have provided a special section for this purpose Section 7.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: All the proofs are given in Appendix B and each proof state all its arguments in a logical order.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The optimization methodology is described in Section 4 and all numerical details and protocol are given in Section 5 and in Appendix E.1-G-Appendix H-Appendix I-Appendix K.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: Synthetic data can be generated, benchmark data are publicly available (Appendix F.2), but the mouse dataset is not. The code is provided as supplementary material and will be made publicly available if the paper is accepted.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.

- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The experiments setting is clearly presented in Section 5 for the experiments on the mouse dataset. The full details of all experiments are presented in Appendix G-Appendix H-Appendix I-E.1.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: Error bars are given for the synthetic experiments, but not for the real experiments due to limited computational budget. There is no randomness in the train/test split but there is randomness in the optimization method (Adabelief), but its impact is not significant on the results.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: This is given at the beginning of Appendix E.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [\[Yes\]](#)

Justification: Every authors are funded for their work and preserve research integrity.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [\[Yes\]](#)

Justification: A special section is given on this subject Appendix A.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.

- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: There is no risk.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: All dataset and python library were cited when necessary.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: The mouse dataset is described in Appendix F.1.

Guidelines:

- The answer NA means that the paper does not release new assets.

- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and Research with Human Subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.