



Multi-view diffusion maps

Ofir Lindenbaum^{a,*}, Arie Yeredor^a, Moshe Salhov^b, Amir Averbuch^b

^a School of Electrical Engineering, Tel Aviv University, Israel

^b School of Computer Science, Tel Aviv University, Israel

ARTICLE INFO

Keywords:

Dimensionality reduction
Manifold learning
Diffusion maps
Multi-view

ABSTRACT

In this paper, we address the challenging task of achieving multi-view dimensionality reduction. The goal is to effectively use the availability of multiple views for extracting a coherent low-dimensional representation of the data. The proposed method exploits the intrinsic relation within each view, as well as the mutual relations between views. The multi-view dimensionality reduction is achieved by defining a cross-view model in which an implied random walk process is restrained to hop between objects in the different views. The method is robust to scaling and insensitive to small structural changes in the data. We define new diffusion distances and analyze the spectra of the proposed kernel. We show that the proposed framework is useful for various machine learning applications such as clustering, classification, and manifold learning. Finally, by fusing multi-sensor seismic data we present a method for automatic identification of seismic events.

1. Introduction

high-dimensional big data are becoming ubiquitous in a growing variety of fields, and pose new challenges in their analysis. Extracted features are useful in analyzing these datasets. However, some prior knowledge or modeling is required in order to identify the essential features. Unsupervised dimensionality reduction methods, on the other hand, aim to find low-dimensional representations based on the *intrinsic geometry* of the analyzed data. A “good” dimensionality reduction methodology reduces the complexity of the data, while preserving its coherency, thereby facilitating data analysis tasks (such as clustering, classification, manifold learning and more) in the reduced space. Many methods such as Principal Component Analysis (PCA) [1], Multidimensional Scaling (MDS) [2], Local Linear Embedding [3], Laplacian Eigenmaps (LE) [4], Diffusion Maps (DM) [5], p-Laplacian Regularization [6], T-distributed Stochastic Neighbor Embedding (t-SNE) [7], Uniform manifold approximation [8] and more have been proposed to achieve robust dimensionality reduction. low-dimensional representations have been shown to be useful in various applications such as face recognition based on LE [9], Non-linear independent component analysis using DM [10], Musical Key extraction employing DM [11], and many more.

Frameworks such as [1–8] do not consider the possibility to exploit more than one view representing the same process. Such multiple views can be obtained, for example, when the same underlying process is observed using several different modalities, or measured with different instrumentations. An additional view can provide meaningful insights

regarding the dynamical process behind the data, whenever the data is indeed generated or governed by such a latent process.

This study is devoted to the development of a framework for obtaining a low-dimensional parametrization from multiple measurements organized as “views”. Our approach essentially relies on quantifying of the speed of a random walk restricted to “hop” between views. The analysis of such a random walk provides a natural representation of the data using joint l organization of the multiple measurements.

The problem of learning from multiple views has been studied in several domains. Some prominent statistical approaches, for addressing this problem are Bilinear Models [12], Partial Least Squares [13] and Canonical Correlation Analysis (CCA) [14]. These methods are powerful for learning the relations among the different views, but are restricted to linear transformation for each view. Some methods such as [15–20] extend the CCA to nonlinear transformations by introducing a kernel, and demonstrate the advantages of fusing multiple views for clustering and classification. A Hessian multiset canonical correlations is presented in [21], the authors propose to overcome limitations of the Laplacian by defining local Hessians on the samples to capture structures of multiple views. Markov based methods such as [22–24] use diffusion distances for classification, clustering or retrieval tasks. The “agreement” (also called “consensus”) between different views is used in [25] to extract the geometric information from all views. A sparsity-based learning algorithm for cross-view dimensionality reduction is proposed in [26]. A neural network which extracts maximally correlated representations is studied in [27]. A review comprehensive article [28] presents recent progress and challenges in this domain.

* Corresponding author.

E-mail address: ofirlin@gmail.com (O. Lindenbaum).

In this work we extend the concepts established in [19,29] to devise a diffusion-based learning framework for fusing multiple views, by seeking the implied underlying low-dimensional structure in the ambient space. Our contributions can be summarized as follows.

1. We present a natural generalization of the DM method [5] for handling multiple views. The generalization is attained by combining the intrinsic relations within each view with the mutual relations between views, so as to construct a multi-view kernel matrix. The proposed kernel defines a cross-view diffusion process, and related diffusion distances, which impose a structured random walk between the various views. In addition, we show that the spectral decomposition of the proposed kernel can be used to find an embedding of the data that offers improved exploitation of the information encapsulated in the multiple views.
2. For the coupled views setting, we analyze theoretical properties of the proposed method. The associated parametrization is justified by being recast as a minimizer of a kernel-based objective function. Then, we relate the spectrum of the proposed kernel to the spectrum of a single-view kernel which is based on naïve concatenation of the views. Finally, under some simplification, we find the infinitesimal generator of the proposed kernel.
3. For practical use, we present an automated method for setting the kernel bandwidth parameter and a proposed procedure for recursively augmenting the representation with each new data point.
4. We demonstrate the applicability of our multi-view method to classification, clustering and manifold learning. Furthermore, by fusing data from multiple seismic sensors we demonstrate an automatic extraction of latent seismic parameters using real-world data.

The paper is structured as follows. Some essential background is provided in Section 2. In Section 3 we formulate the multi-view dimensionality reduction problem and discuss prior work. Then, in Section 4 we present our proposed multi-view DM method and discuss some of its basic properties. Section 5 studies additional theoretical properties of the proposed kernel for the particular (more simple) case of a coupled setting (where only two views are available). Section 6 presents the experimental results. Potential applications are described in Section 7 and concluding remarks appear in Section 8 and Appendix (respectively).

2. Background

2.1. General dimensionality reduction

Consider a high-dimensional dataset $\mathbf{X} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_M\} \in \mathbb{R}^{D \times M}$, $\mathbf{x}_i \in \mathbb{R}^D, i = 1, 2, \dots, M$. The goal is to find a low-dimensional representation $\mathbf{Z} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_M\} \in \mathbb{R}^{r \times M}$, $\mathbf{z}_i \in \mathbb{R}^r, i = 1, 2, \dots, M$, such that $r \ll D$, while latent “inner relations” (if any) among the multidimensional data points are preserved as closely as possible, in some sense. This problem setup is based on the assumption that the data is represented (viewed) in a single vector space (single view).

2.2. Diffusion maps (DM)

DM [5] is a dimensionality reduction method which aims at extracting the intrinsic geometry of the data. The DM framework is highly effective when the data is densely sampled from some low-dimensional manifold, so that its “inner relations” are the local connectivities (or proximities) on the manifold. Given a high-dimensional dataset $\mathbf{X}$, the DM framework consists of the following steps:

1. A kernel function $\mathcal{K} : \mathbf{X} \times \mathbf{X} \rightarrow \mathbb{R}$ is chosen, so as to construct a matrix $\mathbf{K} \in \mathbb{R}^{M \times M}$ with elements $K_{ij} \triangleq \mathcal{K}(\mathbf{x}_i, \mathbf{x}_j)$, satisfying the following properties: (i) Symmetry: $\mathbf{K} = \mathbf{K}^T$; (ii) Positive semi-definiteness: $\mathbf{K} \succeq \mathbf{0}$, namely $\forall \mathbf{v} \in \mathbb{R}^M : \mathbf{v}^T \mathbf{K} \mathbf{v} \geq 0$; and (iii) Non-negativity: $\mathbf{K} \succeq \mathbf{0}$, namely $K_{ij} \geq 0 \forall i, j \in [1, M]$. These properties

guarantee that the matrix $\mathbf{K}$ has real-valued eigenvectors and non-negative eigenvalues. A *Gaussian kernel* is a common example, in which $\mathcal{K}(\mathbf{x}_i, \mathbf{x}_j) \triangleq \exp \left\{ -\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma_x^2} \right\}$, where $\|\cdot\|$ denotes the L_2 norm and σ_x^2 is a user-selected width (scale) parameter.

2. By normalizing the rows of $\mathbf{K}$, the matrix

$$\mathbf{P}^x \triangleq \mathbf{D}^{-1} \mathbf{K} \in \mathbb{R}^{M \times M} \quad (1)$$

is obtained, where $\mathbf{D} \in \mathbb{R}^{M \times M}$ is a diagonal matrix with $D_{i,i} = \sum_j K_{i,j}$. $\mathbf{P}^x$ can be interpreted as the transition probabilities of a (fictitious) Markov chain on $\mathbf{X}$, such that $[(\mathbf{P}^x)^t]_{i,j} \triangleq p_t(\mathbf{x}_i, \mathbf{x}_j)$ (where t is an integer power) describes the implied probability of transition from point $\mathbf{x}_i$ to point $\mathbf{x}_j$ in t steps.

3. Spectral decomposition is applied to $\mathbf{P}^x$, obtaining a set of M eigenvalues $\{\lambda_m\}$ (in descending order) and associated normalized eigenvectors $\{\boldsymbol{\psi}_m\}$ satisfying $\mathbf{P}^x \boldsymbol{\psi}_m = \lambda_m \boldsymbol{\psi}_m$, $m = 0, \dots, M-1$, or $\mathbf{P}^x \boldsymbol{\Psi} = \boldsymbol{\Psi} \boldsymbol{\Lambda}$, where $\boldsymbol{\Psi}$ and $\boldsymbol{\Lambda}$ are (resp.) the eigenvectors and diagonal eigenvalues matrices.
4. A new representation is defined for the dataset $\mathbf{X}$, representing each $\mathbf{x}_i$ by the i th row of $(\mathbf{P}^x)^t \boldsymbol{\Psi} = \boldsymbol{\Psi} \boldsymbol{\Lambda}^t$, namely:

$$\boldsymbol{\Psi}_t(\mathbf{x}_i) : \mathbf{x}_i \mapsto [\lambda_1^t \psi_1[i], \dots, \lambda_{M-1}^t \psi_{M-1}[i]]^T \in \mathbb{R}^{M-1}, i \in [1, M] \quad (2)$$

where t is a selected number of steps and $\psi_m[i]$ denotes the i th element of $\boldsymbol{\psi}_m$. Note that the trivial eigenvector $\boldsymbol{\psi}_0 = \mathbf{1}$ (with corresponding eigenvalue $\lambda_0 = 1$) was omitted from the representation as it does not carry information about the data. The main idea behind this representation is that the Euclidean distance between two data points in the new representation equals a weighted L_2 distance between the conditional probability vectors $p_t(\mathbf{x}_i, :)$ and $p_t(\mathbf{x}_j, :)$, $i, j \in [1, M]$ (the i th and j th rows of $(\mathbf{P}^x)^t$). This weighted Euclidean distance is referred to as the *diffusion distance*, denoted

$$\begin{aligned} D_t^2(\mathbf{x}_i, \mathbf{x}_j) &\triangleq \|\boldsymbol{\Psi}_t(\mathbf{x}_i) - \boldsymbol{\Psi}_t(\mathbf{x}_j)\|^2 \\ &= \sum_{m=1}^{M-1} \lambda_m^{2t} (\psi_m[i] - \psi_m[j])^2 = \|p_t(\mathbf{x}_i, :) - p_t(\mathbf{x}_j, :)\|_{W^{-1}}^2, \end{aligned} \quad (3)$$

where $\mathbf{W} = \mathbf{D}/\text{Trace}\{\mathbf{D}\}$ and where $\|\mathbf{z}\|_Q \triangleq \mathbf{z}^T \mathbf{Q} \mathbf{z}$ denotes a $\mathbf{Q}$ -weighted Euclidean norm (the last equality is shown in [5]).

5. A desired accuracy level $\delta \geq 0$ is chosen for the diffusion distance defined by Eq. (3), such that $r(\delta, t) = \max\{\ell : |\lambda_\ell|^t > \delta |\lambda_1|^t\}$. The new (truncated) $r(\delta, t)$ -dimensional mapping, which leads to the desired representation $\mathbf{Z}$, is then defined as

$$\boldsymbol{\Psi}_t^\delta(\mathbf{x}_i) : \mathbf{x}_i \mapsto [\lambda_1^t \psi_1[i], \lambda_2^t \psi_2[i], \dots, \lambda_{r(\delta, t)}^t \psi_{r(\delta, t)}[i]]^T \triangleq \mathbf{z}_i \in \mathbb{R}^{r(\delta, t)}. \quad (4)$$

This dimensionality reduction approach was found to be useful in various applications in diverse fields. However, as previously noted, it is limited to a single-view representation.

3. Multi-view dimensionality reduction - Problem formulation and prior work

Assume now that we are given multiple sets (views) of observations $\mathbf{X}^\ell, \ell = 1, \dots, L$. Each view is a high-dimensional dataset $\mathbf{X}^\ell = \{\mathbf{x}_1^\ell, \mathbf{x}_2^\ell, \dots, \mathbf{x}_M^\ell\} \in \mathbb{R}^{D_\ell \times M}$, where D_ℓ is the dimension of each feature space. Note that a bijective correspondence between views is assumed. These views can be a result of different measurement devices or different types of features extracted from raw data. For each view $\ell = 1, \dots, L$, we seek a lower-dimensional representation that preserves the “inner relations” between multidimensional data points within each view $\mathbf{X}^\ell$, as well as among all views $\{\mathbf{X}^1, \dots, \mathbf{X}^L\}$. Our goal is to find L representations $\Phi_\ell(\mathbf{X}^1, \dots, \mathbf{X}^L) : \mathbb{R}^{D_\ell} \rightarrow \mathbb{R}^r$, such that $r \ll D_\ell, \ell = 1, \dots, L$.

Before turning to present our proposed approach, we describe existing alternative approaches for incorporating multiple views. In particular, we detail here the methods which we would later use as benchmarks

for comparison to our proposed framework. All of these methods begin by computing the kernel matrix $\mathbf{K}^\ell$ for each (ℓ th) view ($\ell = 1, \dots, L$), and differ in the way of combining these kernels for generating the DM.

3.1. Kernel product DM (KP)

A naïve generalization of DM [5] may be computed using an element-wise kernel product, namely by constructing $\mathbf{K}^\circ \triangleq \mathbf{K}^1 \circ \mathbf{K}^2 \circ \dots \circ \mathbf{K}^L \in \mathbb{R}^{M \times M}$, where $\circ$ denotes Hadamard's (element-wise) matrix product, such that $\mathbf{K}_{i,j}^\circ \triangleq \mathbf{K}_{i,j}^1 \cdot \mathbf{K}_{i,j}^2 \cdot \dots \cdot \mathbf{K}_{i,j}^L$, followed by row-normalization: $\mathbf{P} = (\mathbf{D}^\circ)^{-1} \mathbf{K}^\circ \in \mathbb{R}^{M \times M}$, where $\mathbf{D}^\circ$ is diagonal, with $D_{i,i}^\circ \triangleq \sum_j \mathbf{K}_{i,j}^\circ$.

Note that in the special case of using Gaussian kernels with equal width parameters σ , the resulting matrix $\mathbf{K}^\circ$ is equal to the matrix $\mathbf{K}^w$ constructed from the concatenated observations vector

$$\mathbf{w}_i = \left[(\mathbf{x}_i^1)^T, \dots, (\mathbf{x}_i^L)^T \right]^T \quad (5)$$

such that

$$\mathbf{K}_{i,j}^w = \exp \left\{ -\frac{\|\mathbf{w}_i - \mathbf{w}_j\|^2}{2\sigma^2} \right\}. \quad (6)$$

3.2. Kernel sum DM (KS)

An average diffusion process was used in [23], where the *sum kernel* is defined as

$$\mathbf{K}^+ \triangleq \sum_{\ell=1}^L \mathbf{K}^\ell, \quad (7)$$

followed by row-normalization: $\mathbf{P}^+ = (\mathbf{D}^+)^{-1} \mathbf{K}^+$, where $D_{i,i}^+ = \sum_j \mathbf{K}_{i,j}^+$.

The implied random walk sums up (and normalizes) the step probabilities from each view.

3.3. Kernel canonical correlation analysis (KCCA)

The frameworks [15,16] extend the well-known Canonical Correlation Analysis (CCA) by applying a kernel function prior to the application of CCA. Kernels $\mathbf{K}^1$ and $\mathbf{K}^2$ are constructed for each view, and the canonical vectors $\mathbf{v}_1$ and $\mathbf{v}_2$ are obtained by solving the following generalized eigenvalue problem

$$\begin{bmatrix} \mathbf{0}_{M \times M} & \mathbf{K}^1 \cdot \mathbf{K}^2 \\ \mathbf{K}^2 \cdot \mathbf{K}^1 & \mathbf{0}_{M \times M} \end{bmatrix} \begin{pmatrix} \mathbf{v}_1 \\ \mathbf{v}_2 \end{pmatrix} = \rho \cdot \begin{bmatrix} (\mathbf{K}^1 + \gamma \mathbf{I})^2 & \mathbf{0}_{M \times M} \\ \mathbf{0}_{M \times M} & (\mathbf{K}^2 + \gamma \mathbf{I})^2 \end{bmatrix} \begin{pmatrix} \mathbf{v}_1 \\ \mathbf{v}_2 \end{pmatrix}, \quad (8)$$

where $\gamma \mathbf{I}$ are regularization terms added to prevent overfitting and thus improve generalization of the method. Usually the Incomplete Cholesky Decomposition (ICD) [15,16,30] is used to reduce the run time required for solving (8). For clustering tasks, K -means clustering is applied to the set of generalized eigenvectors.

3.4. Spectral clustering with two views

The approach in [19] generalizes the traditional normalized graph Laplacian for two views. Kernels $\mathbf{K}^1$ and $\mathbf{K}^2$ are computed in each view and multiplied, yielding $\mathbf{W} = \mathbf{K}^1 \cdot \mathbf{K}^2$, from which

$$\mathbf{A} \triangleq \begin{bmatrix} \mathbf{0}_{M \times M} & \mathbf{W} \\ \mathbf{W}^T & \mathbf{0}_{M \times M} \end{bmatrix} \in \mathbb{R}^{2M \times 2M} \quad (9)$$

is obtained. Symmetric normalization is applied by using the diagonal matrix $\bar{\mathbf{D}}$ with diagonal elements $\bar{D}_{i,i} = \sum_j \mathbf{W}_{i,j}$, such that the normalized fused kernel is defined as

$$\bar{\mathbf{A}} = \bar{\mathbf{D}}^{-0.5} \cdot \mathbf{A} \cdot \bar{\mathbf{D}}^{-0.5}. \quad (10)$$

Assuming that the data can be clustered into N_c clusters, and denoting the $2M$ eigenvectors of $\bar{\mathbf{A}}$ as $\phi_i, i = 1, \dots, 2M$ (in descending order of the corresponding eigenvalues), the mapping

$$\Phi[i] \triangleq \frac{1}{s[i]} \left[\phi_1[i], \dots, \phi_{N_c}[i] \right]^T \in \mathbb{R}^{N_c} \quad (11)$$

(where $s[i] = \sum_{j=1}^{N_c} (\phi_j^2[i])$) is useful for applying K -means clustering. The work in [19] is focused on spectral clustering, however, a similar version of the kernel from Eq. (10) is suitable for manifold learning, as we demonstrate in the sequel. In Section 6 this approach is referred to as de Sa's.

As We show in the next section, our proposed approach essentially uses a stochastic matrix version of this kernel, and extends the construction to multiple views. We shall show that our stochastic matrix version is useful, as it provides various theoretical justifications for the implied multi-view diffusion process.

4. Our proposed approach for multi-view diffusion maps

Here we propose our generalization of the DM framework for handling a multi-view scenario. This is done by imposing an implied (fictitious) random walk model using the local connectivities between data points within all views. Our way to generalize the DM framework is by restraining the random walker to “hop” between different views in each step. The construction requires to choose a set of L symmetrical positive semi-definite, non-negative kernels, one for each view $\mathbf{K}^l : \mathbf{X}^l \times \mathbf{X}^l \rightarrow \mathbb{R}$, $l = 1, \dots, L$. We use the Gaussian kernel function, so that the (i, j) th element of each $\mathbf{K}^l \in \mathbb{R}^{M \times M}$ is given by

$$\mathbf{K}_{i,j}^l = \exp \left\{ -\frac{\|\mathbf{x}_i^l - \mathbf{x}_j^l\|^2}{2\sigma_l^2} \right\}, \quad i, j = 1, \dots, M, \quad l = 1, \dots, L, \quad (12)$$

where $\{\sigma_l^2\}_{l=1}^L$ are a set of selected parameters. The decaying property of the Gaussian is useful, as it removes the influence of large Euclidean distances. The *multi-view kernel* is formed by constructing the following matrix

$$\hat{\mathbf{K}} = \begin{bmatrix} \mathbf{0}_{M \times M} & \mathbf{K}^1 \mathbf{K}^2 & \mathbf{K}^1 \mathbf{K}^3 & \dots & \mathbf{K}^1 \mathbf{K}^L \\ \mathbf{K}^2 \mathbf{K}^1 & \mathbf{0}_{M \times M} & \mathbf{K}^2 \mathbf{K}^3 & \dots & \mathbf{K}^2 \mathbf{K}^L \\ \mathbf{K}^3 \mathbf{K}^1 & \mathbf{K}^3 \mathbf{K}^2 & \mathbf{0}_{M \times M} & \dots & \mathbf{K}^3 \mathbf{K}^L \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \mathbf{K}^L \mathbf{K}^1 & \mathbf{K}^L \mathbf{K}^2 & \mathbf{K}^L \mathbf{K}^3 & \dots & \mathbf{0}_{M \times M} \end{bmatrix} \in \mathbb{R}^{LM \times LM}. \quad (13)$$

Finally, using the diagonal row-normalization matrix $\hat{\mathbf{D}} \in \mathbb{R}^{LM \times LM}$ with $\hat{D}_{i,i} = \sum_j \hat{\mathbf{K}}_{i,j}$, the normalized row-stochastic matrix is defined as

$$\hat{\mathbf{P}} = \hat{\mathbf{D}}^{-1} \hat{\mathbf{K}} \in \mathbb{R}^{LM \times LM}. \quad (14)$$

We refer to the (l, m) th block of $\hat{\mathbf{P}}$ as the square $M \times M$ matrix starting at

$[1 + (l-1)M, 1 + (m-1)M]$, $l, m = 1, \dots, L$. Thus, the (i, j) th element of the (l, m) th block describes a (fictitious) probability of transition from $\mathbf{x}_i^l$ to $\mathbf{x}_j^m$. This construction takes into account all possibilities to “hop” between views, under the constraint that staying in the same view (namely, a transition from $\mathbf{x}_i^l$ to $\mathbf{x}_j^l$) is forbidden.

4.1. probabilistic interpretation of $\hat{\mathbf{P}}^t$

Subsequent to our proposed construction (Eqs. (12)–(14)), each element of the power- t matrix $\hat{\mathbf{P}}^t$,

$$\left[\hat{\mathbf{P}}^t \right]_{i+(l-1)M, j+(m-1)M} \triangleq \hat{p}_t(\mathbf{x}_i^l, \mathbf{x}_j^m) \quad (15)$$

denotes the probability of transition from $\mathbf{x}_i^l$ to $\mathbf{x}_j^m$ in t time-steps. Note that, as mentioned above, due to the block-off-diagonal structure of $\hat{\mathbf{P}}$, which forbids a transition into the same view within a single time-step, we have $\hat{p}_1(\mathbf{x}_i^l, \mathbf{x}_j^l) = 0$, $l = 1, \dots, L$, although for $t > 1$ $\hat{p}_t(\mathbf{x}_i^l, \mathbf{x}_j^l)$ may be non-zero.

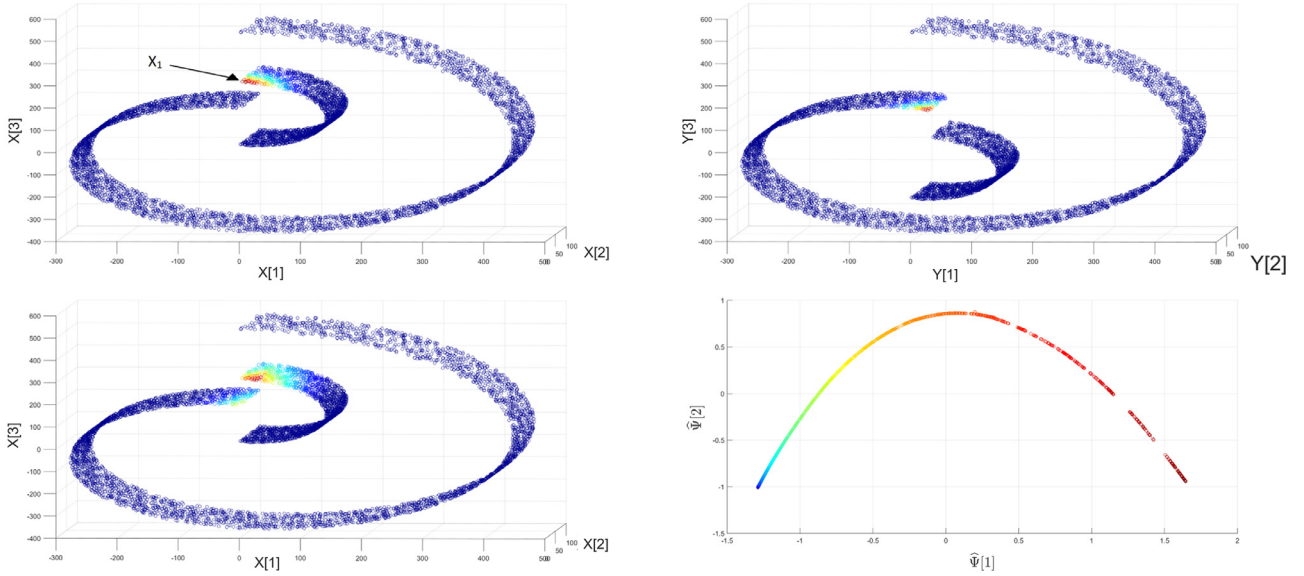


Fig. 1. Top-left: A non-smooth Swiss Roll sampled from View-I (X^1), colored by the single-view probability of transition ($t = 1$) from x_1 to $x_\cdot$. Top-right: A second Swiss Roll, sampled from View-II (X^2), colored by the single-view probability of transition ($t = 1$) from y_1 to $y_\cdot$. Bottom-left: the first Swiss Roll, colored by the multi-view probabilities of transition ($t = 1$) from x_i to $y_\cdot$. In the top-left figure, we highlight x_1 using an arrow. Bottom-right: a low-dimensional representation extracted based on the multi-view transition matrix $\hat{P}$ (Eq. (14)).

4.1.1. Smoothing effect $t = 1$

For simplicity let us examine the term $\hat{p}_t(x_i^l, x_j^m)$, $l \neq m$ for $t = 1$. The transition probability for $t = 1$ is

$$\hat{p}_1(x_i^l, x_j^m) = \frac{\sum_s K_{i,s}^l K_{s,j}^m}{\hat{D}_{i,i}}.$$

This probability takes into consideration all the various connectivities of node x_i^l to node x_s^l and the connectivities of the corresponding node x_s^m to the destination node x_j^m . The proposed multi-view approach has a “smoothing effect” in terms of the transition probabilities, meaning that the probability of transitioning from x_i^l to x_j^m could be positive even if $K_{i,j}^l = 0$ and $K_{i,j}^m = 0$: Assume that a non-empty subset $S = \{s_1, \dots, s_F\} \subseteq [1, M]$ exists, such that $K_{i,s_f}^l > 0$ and $K_{s_f,j}^m > 0$, $f = 1, \dots, F$. Then by definition of the multi-view probability we have $\hat{p}_1(x_i^l, x_j^m) > 0$. Note that although with the Gaussian kernel all elements of K^l are positive, the smoothing effect can still be significant when despite negligibly low values in $K_{i,j}^l$ and in $K_{i,j}^m$, $\hat{p}_1(x_i^l, x_j^m)$ can become relatively high.

Fig. 1 illustrates the multi-view transition probabilities compared to a single-view approach using two ($L = 2$) deformed “Swiss Roll” manifolds. In each view, there is no probability of transition from one side of its observed “gap” to the other side. And yet, the multi-view (one-step) transition probability is non-zero for points at both sides of the gaps. This smoothing effect occurs because the gaps are located at a different position (near different points) on each view, thus allowing the multi-view kernel to smooth out the nonlinear gaps.

4.1.2. Increasing the diffusion step t

Under the stochastic Markov model assumption, raising $\hat{P}$ to a higher power (by increasing the diffusion step t) spreads the probability mass function along its rows based on the connectivities in all views. As described in [5], this probability spread reduces the influence of eigenvectors associated with high-indexed (smaller) eigenvalues on the diffusion distance (Eq. (16)). This implies that the eigenvectors corresponding to low-indexed (large) eigenvalues have a low-frequency content, whereas the eigenvectors corresponding to the high-indexed (small) eigenvalues describe the oscillatory behavior of the data [5]. In Fig. 2, we present the eigenvalues of the matrix $\hat{P}^t$ with different values of t . For this experiment we generated $L = 3$ Swiss Rolls with $M = 1200$ data points

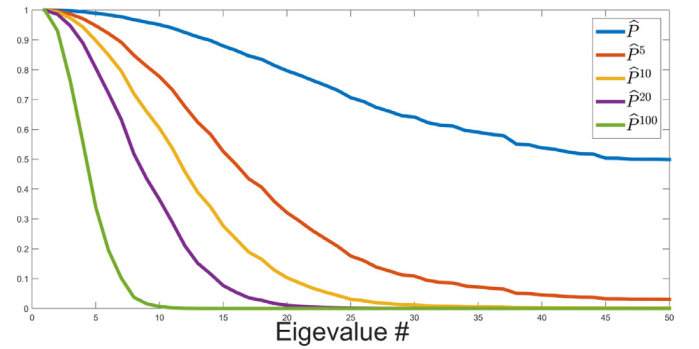


Fig. 2. The decay of the eigenvalues for increasing powers of the matrix $\hat{P}$.

each. It is evident that the numerical rank of $\hat{P}^t$ decreases for higher values of t .

4.2. Multi-view diffusion distance

In a variety of real-world data types, the Euclidean distance between observed data points (vectors) does not provide sufficient information about the intrinsic relations between these vectors, and is highly sensitive to non-unitary transformations. Common tasks such as classification, clustering or system identification often require a measure for the intrinsic connectivity between data points, which is only locally expressed by the Euclidean distance in the high-dimensional ambient space. The multi-view diffusion kernel (defined in Section 4) describes all the small local connections between data points. The row stochastic matrix $\hat{P}^t$ (Eq. (14)) accounts for all possible transitions between data points in t time steps while hopping between the views. For a fixed value $t > 0$, two data points are *intrinsically similar* if the conditional distributions $\hat{p}_t(x_i, \cdot) = [\hat{P}^t]_{i,:}$ and $\hat{p}_t(x_j, \cdot) = [\hat{P}^t]_{j,:}$ are similar. This type of similarity measure indicates that the points x_i and x_j are similarly connected (in the sense of similar probabilities of transitions) to several mutual points. Thus, they are connected by some geometrical path. In many cases, a small Euclidean distance can be misleading due to the fact that two data points can be “close” without having any geodesic path that connects them. Comparing the transition probabilities is more

robust, as it takes into consideration all of the local connectivities between the compared points. Therefore, even if two points seem far from one another in the Euclidean sense, they may still share many common neighbors and thus be “close” in the sense of having a small diffusion distance.

Based on this observation, by expanding the single-view construction given in [5], we define the weighted inner view diffusion distances for the first view as

$$D_t^2(\mathbf{x}_i^1, \mathbf{x}_j^1) \triangleq \sum_{k=1}^{L \cdot M} \frac{1}{\hat{\phi}_0(k)} \left(\left[\hat{\mathbf{P}}^t \right]_{i,k} - \left[\hat{\mathbf{P}}^t \right]_{j,k} \right)^2 = \left\| (\mathbf{e}_i - \mathbf{e}_j)^T \hat{\mathbf{P}}^t \right\|_{\hat{\mathbf{D}}^{-1}}^2, \quad (16)$$

where $1 \leq i, j \leq M$, $\mathbf{e}_i$ is the i th column of an $L \cdot M \times L \cdot M$ identity matrix, $\hat{\phi}_0$ is the (non-normalized) first left eigenvector of $\hat{\mathbf{P}}$, whose k th element is $\hat{\phi}_0(k) = \hat{D}_{k,k}$. A similarly weighted norm is defined for the l th view

$$D_t^2(\mathbf{x}_i^l, \mathbf{x}_j^l) \triangleq \sum_{k=1}^{L \cdot M} \frac{1}{\hat{\phi}_0(k)} \left(\left[\hat{\mathbf{P}}^l \right]_{i+\tilde{l},k} - \left[\hat{\mathbf{P}}^l \right]_{j+\tilde{l},k} \right)^2 = \left\| (\mathbf{e}_{l+i} - \mathbf{e}_{l+j})^T \hat{\mathbf{P}}^l \right\|_{\hat{\mathbf{D}}^{-1}}^2, \quad (17)$$

where $\tilde{l} = (l-1) \cdot M$. The main advantage of these distances (Eqs. (16) and (17)) is that they can be expressed in terms of the eigenvalues and eigenvectors of the matrix $\hat{\mathbf{P}}$. This insight allows us to use a representation (defined in Section 4.3 below) where the induced Euclidean distance is proportional to the diffusion distances defined in Eqs. (16) and (17). Indeed, let $\hat{\mathbf{P}} = \Psi \Lambda \Phi^T$ denote the eigenvalues decomposition of $\hat{\mathbf{P}}$, where $\Psi = [\psi_0, \dots, \psi_{L \cdot M-1}]$ and $\Phi = [\phi_0, \dots, \phi_{L \cdot M-1}]$ denote the matrices of (normalized) right- and left-eigenvectors (resp.), and $\Lambda \triangleq \text{Diag}(\lambda_0, \dots, \lambda_{L \cdot M-1})$ denotes the diagonal matrix of respective eigenvalues. Then

Theorem 1. *The inner view diffusion distance defined by Eqs. (16) and (17) can also be expressed as*

$$D_t^2(\mathbf{x}_i^l, \mathbf{x}_j^l) = \sum_{k=1}^{L \cdot M-1} \lambda_k^{2t} (\psi_k[i + \tilde{l}] - \psi_k[j + \tilde{l}])^2, \quad i, j = 1, \dots, M, \quad (18)$$

where $\tilde{l} = (l-1) \cdot M$.

Proof. Using the eigenvalues decomposition of $\hat{\mathbf{P}}$ we have

$$\hat{\mathbf{P}} \hat{\mathbf{D}}^{-1} (\hat{\mathbf{P}}^t)^T = \Psi \Lambda^t \Phi^T \hat{\mathbf{D}}^{-1} \Phi \Lambda^t \Psi^T. \quad (19)$$

Define the symmetric matrix $\hat{\mathbf{P}}_s \triangleq \hat{\mathbf{D}}^{-1/2} \hat{\mathbf{K}} \hat{\mathbf{D}}^{-1/2} = \hat{\mathbf{D}}^{1/2} \hat{\mathbf{P}} \hat{\mathbf{D}}^{-1/2}$, and note that this matrix is algebraically similar to $\hat{\mathbf{P}}$. Therefore, both matrices share the same set of eigenvalues Λ . Additionally, let Π denote the (left- and right-) orthonormal eigenvectors matrix of $\hat{\mathbf{P}}_s$, namely $\hat{\mathbf{P}}_s = \Pi \Lambda \Pi^T$. The left- and right-eigenvectors matrices of $\hat{\mathbf{P}} = \hat{\mathbf{D}}^{-1/2} \hat{\mathbf{P}}_s \hat{\mathbf{D}}^{1/2}$ are then easily identified as $\Phi = \hat{\mathbf{D}}^{1/2} \Pi$ and $\Psi = \hat{\mathbf{D}}^{-1/2} \Psi$ (resp.), so that $\Phi^T \hat{\mathbf{D}}^{-1} \Phi = \Pi^T \Pi = \mathbf{I}$. Therefore,

$$\begin{aligned} D_t^2(\mathbf{x}_i^l, \mathbf{x}_j^l) &= \left\| (\mathbf{e}_{l+i} - \mathbf{e}_{l+j})^T \hat{\mathbf{P}}^t \right\|_{\hat{\mathbf{D}}^{-1}}^2 \\ &= (\mathbf{e}_{l+i} - \mathbf{e}_{l+j})^T \hat{\mathbf{P}}^t \hat{\mathbf{D}}^{-1} (\hat{\mathbf{P}}^t)^T (\mathbf{e}_{l+i} - \mathbf{e}_{l+j}) \\ &= (\mathbf{e}_{l+i} - \mathbf{e}_{l+j})^T \Psi \Lambda^{2t} \Psi^T (\mathbf{e}_{l+i} - \mathbf{e}_{l+j}) \\ &= \sum_{k=1}^{L \cdot M-1} \lambda_k^{2t} (\psi_k[i + \tilde{l}] - \psi_k[j + \tilde{l}])^2. \end{aligned} \quad (20)$$

The term with $k = 0$ can be excluded from the sum since $\psi_0 = \mathbf{1}$ (an all-ones vector) is always the first right-eigenvector (with eigenvalue 1) for any stochastic matrix, so the term corresponding to $k = 0$ in the sum would vanish anyway. $\square$

4.3. Multi-view data parametrization

Tasks such as classification, clustering or regression in a high-dimensional feature space are considered to be computationally expensive. In addition, the performance in such tasks is highly dependent on

the distance measure used. As explained in Section 4.2, distance measures in the original ambient space are often meaningless in many real life situations. Interpreting Theorem 1 in terms of Euclidean distance enables us to define mappings for every view $\mathbf{X}^l, l = 1, \dots, L$, using the right eigenvectors of $\hat{\mathbf{P}}$ (Eq. (14)) weighted by λ_i^l . The representation for instances in $\mathbf{X}^l$ is given by

$$\hat{\Psi}_t(\mathbf{x}_i^l) : \mathbf{x}_i^l \mapsto [\lambda_1^l \psi_1[i + \tilde{l}], \dots, \lambda_{M-1}^l \psi_{M-1}[i + \tilde{l}]]^T \in \mathbb{R}^{M-1}, \quad (21)$$

where $\tilde{l} = (l-1) \cdot M$. These L mappings capture the intrinsic geometry of the views as well as the mutual relations between them. As shown in [31], the set of eigenvalues λ_m has a decaying property such that $1 = |\lambda_0| \geq |\lambda_1| \geq \dots \geq |\lambda_{M-1}|$. Exploiting the decaying property enables us to represent data up to a dimension r where $r \ll D_1, \dots, D_L$. The dimension $r = r(\delta)$ is determined by approximating the diffusion distance (Eq. (18)) up to a desired accuracy δ (we elaborate on this issue in Section 5.4). The reduced dimension version of $\hat{\Psi}_t(\mathbf{X})$ is denoted by $\hat{\Psi}_t^r(\mathbf{X})$.

Using the inner view diffusion distances defined in Eqs. (16) and (17), we define a multi-view diffusion distance as the sum of inner views distances,

$$D_t^{(MV)^2}(i, j) \triangleq \sum_{l=1}^L \left| \hat{\Psi}_t^r(\mathbf{x}_i^l) - \hat{\Psi}_t^r(\mathbf{x}_j^l) \right|^2. \quad (22)$$

This distance is the induced Euclidean distance in a space constructed from the concatenation of all low-dimensional multi-view mappings

$$\hat{\Psi}_t(\mathbf{X}) = [\hat{\Psi}_t^r(\mathbf{X}^1); \hat{\Psi}_t^r(\mathbf{X}^2); \dots; \hat{\Psi}_t^r(\mathbf{X}^L)] \in \mathbb{R}^{L \cdot r \times M}. \quad (23)$$

This mapping is used in Section 7.1 for the experimental evaluation of clustering.

4.4. Multi-view kernel bandwidth

When constructing the Gaussian kernels $\mathbf{K}^l, l = 1, \dots, L$, in Eq. (12), the values of the scale (width) parameter σ_l^2 have to be set. Setting these values to be too small may result in very small local neighborhoods that are unable to capture the local structures around the data points. Conversely, setting the values to be too large may result in a fully connected graph that may generate a too coarse description of the data. Several approaches have been proposed in the literature for determining the kernel width (scale):

In [32], a max-min measure is suggested such that the scale becomes

$$\sigma_l^2 = C \cdot \max_j \left[\min_{i, i \neq j} (||\mathbf{x}_i^l - \mathbf{x}_j^l||^2) \right]. \quad (24)$$

Based on empirical results, the authors in [32] suggest to set C within the range [1, 1.5]. The specific choice of C should be sufficient to capture all local connectivities. This single-view approach could be relaxed in the context of our multi-view scenario. The multi-view kernel $\hat{\mathbf{K}}$ (Eq. 13) consists of products of single-view kernel matrices $\mathbf{K}^l, l = 1, \dots, L$ (Eq. (12)). The diagonal values of each kernel matrix $\mathbf{K}^l$ are all 1's, therefore, in order to have connectivity in all rows of $\hat{\mathbf{K}}$ a connectivity in only one of the views is sufficient. By connectivity, we mean that there is at least one nonzero value in the affinity kernel. This insight suggests that a smaller value for the parameter C could be used in the multi view setting. In [33], the authors provide a probabilistic interpretation of the choice of C , and present a grid-search algorithm for optimizing the choice of C such that each point in the dataset is connected to at least one other point.

Another scheme, proposed in [34], aims to find a range of values for σ_l . The idea is to compute the kernel $\mathbf{K}^l$ (Eq. (12)) for various values of σ and search for the range of values where the Gaussian bell shape is more pronounced. To find such range of valid values for σ , first a logarithmic function is applied to the sum of all elements of the kernel matrix, next the range is identified as the maximal range where the logarithmic plot is linear. We expand this idea for a multi-view scenario based on the following algorithm:

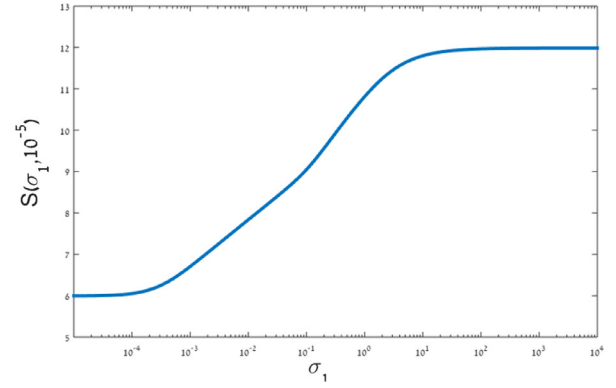
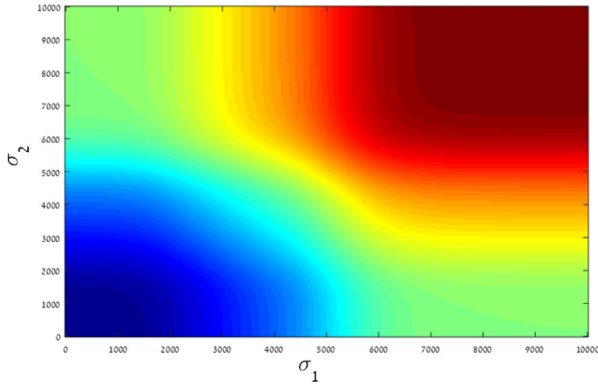


Fig. 3. Left: an example of the two dimensional function $S(\sigma_l, \sigma_m)$. Right: a slice at the first row ($\sigma_2 = 10^{-5}$). The asymptotes are clearly visible in both figures. **Algorithm 1** exploits the multi-view to set a small scale parameter for both views.

Note that the two dimensional function $S^{lm}(\sigma_l, \sigma_m)$ consists of two asymptotes, $S^{lm}(\sigma_l, \sigma_m) \xrightarrow{\sigma_l, \sigma_m \rightarrow 0} \log(N)$, and $S^{lm}(\sigma_l, \sigma_m) \xrightarrow{\sigma_l, \sigma_m \rightarrow \infty} \log(N^3) = 3\log(N)$, since for $\sigma_l, \sigma_m \rightarrow 0$, both K^l and K^m approach the Identity matrix, and for $\sigma_l, \sigma_m \rightarrow \infty$, both K^l and K^m approach all-ones matrices. An example of the plot $S^{lm}(\sigma_l, \sigma_m)$ for two views ($L = 2$) is presented in Fig. 3. The range of each σ_l should reflect the asymptotic behaviour, and may be determined empirically by choosing $\min(\sigma_l) \ll \text{median}\{\|\mathbf{x}_i^l - \mathbf{x}_j^l\|\}$ and $\max(\sigma_l) \gg \text{median}\{\|\mathbf{x}_i^l - \mathbf{x}_j^l\|\}$. In practice the range $[10^{-5}, 10^5]$ is sufficient.

4.5. Computational complexity

Let us now consider the computational complexity required for extracting $\Phi_l(\mathbf{X}^1, \dots, \mathbf{X}^L) \in \mathbb{R}^{r \times M}$ from $\mathbf{X}^l = \{\mathbf{x}_1^l, \mathbf{x}_2^l, \mathbf{x}_3^l, \dots, \mathbf{x}_M^l\} \in \mathbb{R}^{D_l \times M}$. The computational complexity of computing L kernels is $\mathcal{O}(\sum_l M^2 D_l)$. Computing the proposed normalized kernel $\hat{\mathbf{P}}$ (Eq. 14) adds a complexity of $\mathcal{O}(\frac{L(L-1)}{4} M^3)$. The final step requires the spectral decomposition of $\hat{\mathbf{P}}$, thus it is the most computationally expensive step and adds a complexity of $\mathcal{O}(L^3 M^3)$. However, due to the low rank, sparse nature of $\hat{\mathbf{P}}$, this spectral decomposition could be approximated using random projection methods such as in [35–37]. Using random projection the complexity of the spectral decomposition is improved to $\mathcal{O}(L^2 M^2 \log r)$. For the optional step proposed in Algorithm 1, if n_0 values are chosen for σ , the additional complexity cost is of $\mathcal{O}(M L n_0)$.

Algorithm 1 Multi-view kernel bandwidth selection.

Input: Multiple sets of observations (views) $\mathbf{X}^l, l = 1, \dots, L$.

Output: Scale parameters $\{\sigma_1, \dots, \sigma_L\}$.

- 1: Compute Gaussian kernels $K^l(\sigma_l), l = 1, \dots, L$ for several values of $\sigma_l \in [10^{-5}, 10^5]$.
 - 2: Compute for all pairs $l \neq m$: $S^{lm}(\sigma_l, \sigma_m) = \sum_i \sum_j K_{i,j}^{lm}(\sigma_l, \sigma_m)$, where

$$K^{lm}(\sigma_l, \sigma_m) = K^l(\sigma_l) \cdot K^m(\sigma_m).$$
 - 3: **for** $l = 1 : L$ **do**
 - 4: Find the minimal value for σ_l such that $S^{lm}(\sigma_l, \sigma_m)$ is linear for all $m \neq l$.
 - 5: **end for**
-

5. Coupled views $L = 2$

In this section we provide analytical results for the special case of a coupled data set (i.e $L = 2$). Some of the results could be expanded to a larger number of views but not in a straightforward manner. To simplify the notation in the rest of this section we denote $\mathbf{X} \triangleq \mathbf{X}^1$ and $\mathbf{Y} \triangleq \mathbf{X}^2$.

5.1. Coupled mapping

The mappings provided by our approach (Eq. 21) are justified by the relations given by Eq. 18). In this section we provide some intuition on the relation between the mappings of $\mathbf{X}$ and $\mathbf{Y}$. We focus on the analysis of a 1-dimensional mapping for each view. Let $\rho^x \triangleq \rho(\mathbf{X}) = [\rho(\mathbf{x}_1), \rho(\mathbf{x}_2), \dots, \rho(\mathbf{x}_M)]^T$ and $\rho^y \triangleq \rho(\mathbf{Y}) = [\rho(\mathbf{y}_1), \rho(\mathbf{y}_2), \dots, \rho(\mathbf{y}_M)]^T$ denote such 1-dimensional mappings (one for each view) and define $K^z \triangleq K^x \cdot K^y$, where K^x, K^y are computed from $\mathbf{X}$ and $\mathbf{Y}$ (resp.) based on Eq. 12). The following theorem characterizes the desirable 1-dimensional mappings ρ^x and ρ^y .

Theorem 2. A 1-dimensional zero-mean representation which minimizes the following objective function

$$\begin{aligned} \min_{\rho^x, \rho^y} \sum_{i,j} (\rho(\mathbf{x}_i) - \rho(\mathbf{y}_j))^2 K_{i,j}^z \\ \text{s.t. } \|\rho^x\|^2 + \|\rho^y\|^2 = 1, \quad \sum_i (\rho(\mathbf{x}_i) + \rho(\mathbf{y}_i)) = 0, \end{aligned} \quad (25)$$

is obtained by setting $\hat{\rho} \triangleq [\rho^{xT} \rho^{yT}]^T = \psi_1$, where ψ_1 is the first non-trivial normalized eigenvector of the graph Laplacian $\hat{\mathbf{D}} - \hat{\mathbf{K}}$. The scaling constraint avoids arbitrary scaling of the representation, while the zero-mean constraint (implying orthogonality of $\hat{\rho}$ to the constant all-ones vector $\mathbf{1}$) eliminates any presence of a trivial (constant) component in the representations.

Note that K^z is an affinity measure of points based on $K^x \cdot K^y$. This means that a small value of $K_{i,j}^z$ indicates weak connectivity between data points i and j (possibly through an intermediate point $\ell \in [1, M]$), implying that the distance between $\rho(\mathbf{x}_i)$ and $\rho(\mathbf{y}_j)$ can be large. Conversely, if $K_{i,j}^z$ is large, indicating strong connectivity (possibly through an intermediate point) between points i and j , the distance between $\rho(\mathbf{x}_i)$ and $\rho(\mathbf{y}_j)$ should be small when aiming to minimize the objective function.

Proof. By expanding the objective function in Eq. (25) we get

$$\begin{aligned} \sum_{i,j} (\rho(\mathbf{x}_i) - \rho(\mathbf{y}_j))^2 K_{i,j}^z &= \sum_{i,j} \rho^2(\mathbf{x}_i) K_{i,j}^z + \sum_{i,j} \rho^2(\mathbf{y}_j) K_{i,j}^z \\ &- 2 \sum_{i,j} \rho(\mathbf{x}_i) \rho(\mathbf{y}_j) K_{i,j}^z = \sum_i \rho^2(\mathbf{x}_i) \sum_j K_{i,j}^z + \sum_j \rho^2(\mathbf{y}_j) \sum_i K_{i,j}^z \\ &- \sum_{i,j} \rho(\mathbf{x}_i) \rho(\mathbf{y}_j) K_{i,j}^z - \sum_{i,j} \rho(\mathbf{x}_j) \rho(\mathbf{y}_i) K_{j,i}^z = \sum_i \rho^2(\mathbf{x}_i) D_{i,i}^{\text{rows}} \\ &+ \sum_j \rho^2(\mathbf{y}_j) D_{j,j}^{\text{cols}} - \sum_{i,j} \rho(\mathbf{x}_i) \rho(\mathbf{y}_j) K_{i,j}^z - \sum_{j,i} \rho(\mathbf{x}_j) \rho(\mathbf{y}_i) K_{j,i}^z \\ &= [\rho^{xT} \quad \rho^{yT}] \left[\begin{bmatrix} D^{\text{rows}} & \mathbf{0}_{M \times M} \\ \mathbf{0}_{M \times M} & D^{\text{cols}} \end{bmatrix} - \begin{bmatrix} \mathbf{0}_{M \times M} & K^z \\ (K^z)^T & \mathbf{0}_{M \times M} \end{bmatrix} \right] \begin{bmatrix} \rho^x \\ \rho^y \end{bmatrix}, \end{aligned} \quad (26)$$

where $D_{i,i}^{\text{rows}} = \sum_{j=1}^M K_{i,j}^z$ and $D_{j,j}^{\text{cols}} = \sum_{i=1}^M K_{i,j}^z$ are diagonal matrices.

The same minimization problem (Eq. (25)) can therefore be rewritten as

$$\min_{\hat{\mathbf{P}}} \hat{\mathbf{P}}^T (\hat{\mathbf{D}} - \hat{\mathbf{K}}) \hat{\mathbf{P}} \quad \text{s.t.} \quad \|\hat{\mathbf{P}}\| = 1, \quad \hat{\mathbf{P}}^T \mathbf{1} = 0. \quad (27)$$

Without the orthogonality constraint, this minimization problem could be solved by finding the minimal eigenvalue of $(\hat{\mathbf{D}} - \hat{\mathbf{K}})\hat{\mathbf{P}}^T = \bar{\lambda}\hat{\mathbf{P}}^T$. This eigenproblem has a trivial solution which is the all-ones eigenvector $\psi_0 = \mathbf{1}$ with $\bar{\lambda} = 0$. However, to satisfy the orthogonality constraint, we must use a different eigenvector (which would naturally be orthogonal to ψ_0 due to the symmetry of the matrix), leading to the second-smallest (smallest non-zero) eigenvalue λ_1 with its corresponding eigenvector ψ_1 . $\square$

5.2. Spectral decomposition

In this section, we show how to efficiently compute the spectral decomposition of $\hat{\mathbf{P}}$ (Eq. (14)) when only two view exist ($L = 2$). As already mentioned above, the matrix $\hat{\mathbf{P}}$ is algebraically similar (conjugation) to the symmetric matrix $\hat{\mathbf{P}}_s \triangleq \hat{\mathbf{D}}^{1/2} \hat{\mathbf{P}} \hat{\mathbf{D}}^{-1/2} = \hat{\mathbf{D}}^{-1/2} \hat{\mathbf{K}} \hat{\mathbf{D}}^{-1/2}$. Therefore, both $\hat{\mathbf{P}}$ and $\hat{\mathbf{P}}_s$ share the same set of eigenvalues. Due to symmetry of the matrix $\hat{\mathbf{P}}_s$, it has a set of $2M$ real eigenvalues $\{\lambda_i\}_{i=0}^{2M-1} \in \mathbb{R}$ and corresponding real orthogonal eigenvectors $\{\pi_m\}_{m=0}^{2M-1} \in \mathbb{R}^{2M}$, thus, $\hat{\mathbf{P}}_s = \Pi \Lambda \Pi^T$. By denoting $\Psi = \hat{\mathbf{D}}^{-1/2} \Pi$ and $\Phi = \hat{\mathbf{D}}^{1/2} \Pi$, we conclude that the set $\{\psi_m, \phi_m\}_{m=0}^{2M-1} \in \mathbb{R}^{2M}$ are the right and the left eigenvectors of $\hat{\mathbf{P}} = \Psi \Lambda \Phi^T$, respectively, satisfying $\psi_i^T \phi_j = \delta_{i,j}$ (Kronecker's delta function). In the sequel, we use the symmetric matrix $\hat{\mathbf{P}}_s$ to simplify the analysis.

To avoid the spectral decomposition of a $2M \times 2M$ matrix $\hat{\mathbf{P}}_s$, the spectral decomposition of $\hat{\mathbf{P}}_s$ can be computed using the Singular Value Decomposition (SVD) of the matrix $\hat{\mathbf{K}}^z = (\mathbf{D}^{\text{rows}})^{-1/2} \mathbf{K}^z (\mathbf{D}^{\text{cols}})^{-1/2}$ of size $M \times M$ where $D_{i,i}^{\text{rows}} = \sum_{j=1}^M K_{i,j}^z$ and $D_{j,j}^{\text{cols}} = \sum_{i=1}^M K_{i,j}^z$ are diagonal matrices. Theorem 3 enables us to form the eigenvectors of $\hat{\mathbf{P}}$ as a concatenation of the singular vectors of $\mathbf{K}^z = \mathbf{K}^x \cdot \mathbf{K}^y$.

Theorem 3. By using the left and right singular vectors of $\mathbf{K}^z = \mathbf{V} \Sigma \mathbf{U}^T$, the eigenvectors and the eigenvalues of $\hat{\mathbf{K}}$ are given explicitly by

$$\Pi = \frac{1}{\sqrt{2}} \begin{bmatrix} \mathbf{V} & \mathbf{V} \\ \mathbf{U} & -\mathbf{U} \end{bmatrix}, \Lambda = \begin{bmatrix} \Sigma & \mathbf{0}_{M \times M} \\ \mathbf{0}_{M \times M} & -\Sigma \end{bmatrix}. \quad (28)$$

Proof. Both $\mathbf{V}$ and $\mathbf{U}$ are orthonormal sets, therefore, $\mathbf{u}_i^T \mathbf{u}_j = \delta_{i,j}$, and $\mathbf{v}_i^T \mathbf{v}_j = \delta_{i,j}$, thus, the set $\{\pi_m\}$ is orthonormal. Therefore, $\Pi \Pi^T = \mathbf{I}$. By direct substitution of Eq. (28), $\Pi \Lambda \Pi^T$ can be computed explicitly as

$$\begin{aligned} \Pi \Lambda \Pi^T &= \frac{1}{2} \begin{bmatrix} \mathbf{V} & \mathbf{V} \\ \mathbf{U} & -\mathbf{U} \end{bmatrix} \begin{bmatrix} \Sigma & \mathbf{0}_{M \times M} \\ \mathbf{0}_{M \times M} & -\Sigma \end{bmatrix} \begin{bmatrix} \mathbf{V}^T & \mathbf{U}^T \\ \mathbf{V}^T & -\mathbf{U}^T \end{bmatrix} \\ &= \frac{1}{2} \begin{bmatrix} \mathbf{V} \Sigma & -\mathbf{V} \Sigma \\ \mathbf{U} \Sigma & \mathbf{U} \Sigma \end{bmatrix} \begin{bmatrix} \mathbf{V}^T & \mathbf{U}^T \\ \mathbf{V}^T & -\mathbf{U}^T \end{bmatrix} = \frac{1}{2} \begin{bmatrix} \mathbf{0}_{M \times M} & 2\mathbf{K}^z \\ (2\mathbf{K}^z)^T & \mathbf{0}_{M \times M} \end{bmatrix} = \hat{\mathbf{K}}. \end{aligned} \quad (29)$$

$\square$

Thus the proposed mapping in Eq. (21) could be computed for $L = 2$ using the SVD of $\mathbf{K}^z$, Eq. (28) and $\Psi = \hat{\mathbf{D}}^{-1/2} \Pi$.

5.3. Cross-view (CV) diffusion distance

In some physical systems the observed dataset $\mathbf{X}$ may depend on some underlying parameter, say α . Under this model, multiple snapshots (views) can typically be obtained for various values of α , each denoted $\mathbf{X}^\alpha$. An example of such a scenario occurs in hyper-spectral images that change over time. Quantifying the amount of change in the datasets due to a change in α in such models is often desired, but if the datasets are high-dimensional, this can be quite a challenging task. This scenario was recently studied in [38], where the DM framework was applied to each

fixed value of α . Then, by using the extracted low-dimensional mappings, the Euclidean distance enables to quantify the extent of changes due to α . However, this approach may be rather unstable, since every small change in the data can result in different mappings and the mappings are extracted independently of any mutual influence.

Our multiview approach, on the other hand, incorporates the mutual relations of data within each view, as well as the relations between views. This point of view facilitates a more robust measure of the extent of variation between two datasets that correspond to a small variation in α . To this end, we define a new diffusion distance, which measures the relation between two views, i.e. between all the data points obtained for different values of α . We measure the distance between all the coupled data points among the mappings of the snapshots $\mathbf{X}^{\alpha_l}$ and $\mathbf{X}^{\alpha_m}$ by using the expression

$$D_i^{(\text{CV})^2}(\mathbf{X}^{\alpha_l}, \mathbf{X}^{\alpha_m}) \triangleq \sum_{i=1}^M \left\| \hat{\Psi}_i(\mathbf{x}_i^{\alpha_l}) - \hat{\Psi}_i(\mathbf{x}_i^{\alpha_m}) \right\|^2. \quad (30)$$

Our kernel matrix is a product of the Gaussian kernel matrices in each view. If these values of the kernel matrices ($\mathbf{K}^{\mathbf{x}^{\alpha_l}}, \mathbf{K}^{\mathbf{x}^{\alpha_m}}$) are similar, this corresponds to similarity between the views' inner geometry. The right- and left-singular vectors of the matrix $\mathbf{K}^{\mathbf{x}^{\alpha_l}} \mathbf{K}^{\mathbf{x}^{\alpha_m}}$ will be similar, thus, $D_i^{(\text{CV})}$ will be small.

Theorem 4. Using Gaussian kernels with identical σ values, the cross manifold distance (defined in Eq. (30)) is invariant to orthonormal transformations between the ambient spaces $\mathbf{X}^{\alpha_l}$ and $\mathbf{X}^{\alpha_m}$.

Proof. Denote an orthonormal transformation matrix $\mathbf{R} : \mathbf{X}^{\alpha_l} \rightarrow \mathbf{X}^{\alpha_m}$, w.l.o.g. by $\mathbf{x}_i^{\alpha_m} = \mathbf{R} \mathbf{x}_i^{\alpha_l}$. Then

$$\begin{aligned} K_{i,j}^{\mathbf{x}^{\alpha_m}} &= \exp \left\{ -\frac{\|\mathbf{x}_i^{\alpha_m} - \mathbf{x}_j^{\alpha_m}\|^2}{2\sigma_m^2} \right\} = \exp \left\{ -\frac{\|\mathbf{R} \mathbf{x}_i^{\alpha_l} - \mathbf{R} \mathbf{x}_j^{\alpha_l}\|^2}{2\sigma_m^2} \right\} \\ &= \exp \left\{ -\frac{\|\mathbf{x}_i^{\alpha_l} - \mathbf{x}_j^{\alpha_l}\|^2}{2\sigma_l^2} \right\} = K_{i,j}^{\mathbf{x}^{\alpha_l}}. \end{aligned} \quad (31)$$

The penultimate transition is due to the orthonormality of $\mathbf{R}$ and to the identical $\sigma_l = \sigma_m$ values. Therefore, the matrix $\mathbf{K}^z = (\mathbf{K}^{\mathbf{x}^{\alpha_m}})^2$ from Eq. (12) is symmetric and its right- and left-singular vectors are equal, i.e. $\mathbf{U} = \mathbf{V}$ in Eq. (28). This induces a repetitive form in $\Psi = \hat{\mathbf{D}}^{-1/2} \Pi \rightarrow \psi_l[i] = \psi_l[M+i]$, $1 \leq i, l \leq M-1 \rightarrow \Psi_i(\mathbf{x}_i^{\alpha_l}) = \Psi_i(\mathbf{x}_i^{\alpha_m})$, thus, $D_i^{(\text{CV})^2}(\mathbf{X}^{\alpha_l}, \mathbf{X}^{\alpha_m}) = 0$. $\square$

5.4. Spectral decay of $\hat{\mathbf{K}}$

The power of kernel based methods for dimensionality reduction stems from the spectral decay of the kernel matrix' eigenvalues. In this section we study the relation between the spectral decay of the Kernel Product (Eq. (6)) and our multi-view kernel (Eq. (14)).

Theorem 5. All $2M$ eigenvalues of $\hat{\mathbf{P}}$ (Eq. (14)) are real-valued and bounded, $|\lambda_i| \leq 1$, $i = 0, \dots, 2M-1$.

Proof. ¹ As shown in Section 5.2, $\hat{\mathbf{P}}$ is algebraically similar to a symmetric matrix, thus its eigenvalues are guaranteed to be real-valued. Denote by λ and ψ an eigenvalue and an eigenvector, resp., such that $\lambda \psi = \hat{\mathbf{P}} \psi$.

Define $i_0 \triangleq \arg \max_i |\psi[i]|$ (the index of the largest element of ψ). The maximal value $\psi[i_0]$ can be computed using $\hat{\mathbf{P}}$ from Eq. (14),

$$\begin{aligned} \lambda \psi[i_0] &= \sum_{j=0}^{2M-1} \hat{P}_{i_0 j} \psi[j] \Rightarrow |\lambda| = \left| \sum_{j=0}^{2M-1} \hat{P}_{i_0 j} \frac{\psi[j]}{\psi[i_0]} \right| \leq \sum_{j=0}^{2M-1} \hat{P}_{i_0 j} \frac{|\psi[j]|}{|\psi[i_0]|} \\ &\leq \sum_{j=1}^{2M} \hat{P}_{i_0 j} = 1. \end{aligned} \quad (32)$$

¹ A similar proof is given in <https://sites.google.com/site/yoelshkolnisky/teaching>.

The first inequality is due to the triangle inequality and the second equality is due to the definition of i_0 .

□

Although, according to this Theorem, the eigenvalues are all real-valued and bounded, their mere boundedness is generally insufficient for dimensionality reduction. Dimensionality reduction is meaningful only in the presence of a significant spectral decay.

Definition 1. Let $\mathcal{M}$ be a manifold. The intrinsic dimension d of the manifold is a positive integer determined by how many independent “coordinates” are needed to describe $\mathcal{M}$. Using a parametrization to describe a manifold, the dimension of $\mathcal{M}$ is the smallest integer d such that a smooth map $f(\xi) = \mathcal{M}$ describes the manifold, where $\xi \in \mathbb{R}^d$.

Our framework employs a Gaussian kernel, the spectral decay of which was studied in [5]. We use Lemma 1 (which is based on Weyl’s asymptotic law and appears in [31]) to evaluate the spectral decay of our kernel.

Lemma 1. Assume that the data is sampled from a manifold with intrinsic dimension $d \ll M$. Let $\mathbf{K}^\circ \in \mathbb{R}^{M \times M}$ (a kernel constructed based on the concatenation of views) denote the kernel with an exponential decay as a function of the Euclidean distance. For $\delta > 0$, the number of eigenvalues of $\mathbf{K}^\circ$ above δ is proportional to $(\log(\frac{1}{\delta}))^d$.

Assume that the eigenvalues of $\mathbf{K}^\circ$ are arranged in descending order of their absolute values, $|\lambda_0| \geq |\lambda_1| \geq \dots \geq |\lambda_{M-1}|$, set $\delta \in (0, 1)$ and define $r_\delta \triangleq \max\{\ell \in [1, M] : |\lambda_{\ell-1}| > \delta\}$ (denoting the number of eigenvalues of $\mathbf{K}^\circ$ above δ). Recall now that $\mathbf{K}^\circ = \mathbf{K}^x \circ \mathbf{K}^y$ corresponds to a single DM view (formed of the concatenation of two views) as addressed in [5]. Theorem 6 relates the spectral decay of our proposed kernel $\hat{\mathbf{P}}$ (Eq. (14)) to the decay of the Kernel Product-based DM ($\mathbf{P}^\circ$).

Lemma 2. Let $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{M \times M}$ be any two positive semi-definite (PSD) matrices. Then for any $1 \leq k \leq M-1$

$$\prod_{\ell=k}^{M-1} \lambda_\ell(\mathbf{A} \cdot \mathbf{B}) \leq \prod_{\ell=k}^{M-1} \lambda_\ell(\mathbf{A} \circ \mathbf{B}) \quad (33)$$

where $\lambda_\ell(\cdot)$ denotes the ℓ th eigenvalue (in descending order) of the enclosed matrix.

This inequality is proved in [39] and [40].

Theorem 6. The product of the last $M-1-r_\delta$ eigenvalues of $\mathbf{K}^x$ is smaller or equal to δ^{M-1-r_δ} . Formally, $\prod_{\ell=r_\delta}^{M-1} \lambda_\ell(\mathbf{K}^x \cdot \mathbf{K}^y) \leq \delta^{M-1-r_\delta}$.

Proof. Substitute the PSD matrices $\mathbf{A} = \mathbf{K}^x$ and $\mathbf{B} = \mathbf{K}^y$ in Lemma 2 and choose $\ell = r_\delta$ in Eq. (33) to obtain

$$\prod_{\ell=r_\delta}^{M-1} \lambda_\ell(\mathbf{K}^x \cdot \mathbf{K}^y) \leq \prod_{\ell=r_\delta}^{M-1} \lambda_\ell(\mathbf{K}^\circ) \leq \delta^{M-1-r_\delta}. \quad (34)$$

□

Relying on the spectral decay of the kernel matrix, we can approximate Eq. (18) by neglecting all eigenvalues smaller than δ . Thus, we can compute a low dimensional mapping such that

$$\hat{\Psi}_t^r(\mathbf{x}_i) : \mathbf{x}_i \mapsto [\lambda_1^t \psi_1[i], \lambda_2^t \psi_2[i], \lambda_3^t \psi_3[i], \dots, \lambda_{r-1}^t \psi_{r-1}[i]]^T \in \mathbb{R}^{r-1}. \quad (35)$$

The following Lemma introduces an error bound for using this low-dimensional mapping:

Lemma 3. The truncated diffusion distance up to coordinate r defined as

$$[D_t^r(\mathbf{x}_i, \mathbf{x}_j)]^2 \triangleq \|\hat{\Psi}_t^r(\mathbf{x}_i) - \hat{\Psi}_t^r(\mathbf{x}_j)\|^2 = \sum_{s=1}^r \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2, \quad (36)$$

is bounded by the inner view diffusion distance (defined in Eq. (16))

$$\begin{aligned} & 2 \cdot \left[\sum_{s=1}^{M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2 - \delta^{2t} \cdot \left(\frac{1 - \Delta_{i,j}}{\hat{D}_{i,j}^{\min}} \right) \right] \leq [D_t^r(\mathbf{x}_i, \mathbf{x}_j)]^2 \\ & \leq 2 \cdot \sum_{s=1}^{M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2, \end{aligned}$$

where $\hat{D}_{i,j}^{\min}$ is the minimal value of $\hat{D}_{i,i}$ and $\hat{D}_{j,j}$, $\Delta_{i,j}$ is the Kronecker delta function and δ is the accuracy threshold.

Proof. For the last inequality, clearly

$$[D_t^r(\mathbf{x}_i, \mathbf{x}_j)]^2 \leq [D_t^{2M}(\mathbf{x}_i, \mathbf{x}_j)]^2 = 2 \cdot \sum_{s=1}^{M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2,$$

the equality is a result of the repetitive form of Ψ and Λ , which were defined in Theorem 3 for $L = 2$. Note that $s = 0$ was excluded from the sum as $\psi_0 = 1$ is constant. For the first inequality, using $\Psi = \hat{\mathbf{D}}^{-1/2} \Pi$ where Π is an orthonormal basis defined in Eq. (28), we get

$$\Psi \Psi^T = \hat{\mathbf{D}}^{-1/2} \Pi \Pi^T \hat{\mathbf{D}}^{-1/2} = \hat{\mathbf{D}}^{-1},$$

which means that

$$\sum_{s=0}^{2M-1} (\psi_s[i] - \psi_s[j])^2 = \frac{1}{\hat{D}_{i,i}} + \frac{1}{\hat{D}_{j,j}} - \frac{2\Delta_{i,j}}{\hat{D}_{i,i}}. \quad (37)$$

Using the definition of the truncated diffusion distance we have

$$\begin{aligned} [D_t^r(\mathbf{x}_i, \mathbf{x}_j)]^2 &= \sum_{s=0}^{2M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2 - \sum_{s=r+1}^{2M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2 \\ &\geq 2 \sum_{s=1}^{M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2 - \delta^{2t} \sum_{s=0}^{2M-1} (\psi_s[i] - \psi_s[j])^2 \\ &\geq 2 \left[\sum_{s=1}^{2M-1} \lambda_s^{2t} (\psi_s[i] - \psi_s[j])^2 - \delta^{2t} \cdot \left(\frac{1 - \Delta_{i,j}}{\hat{D}_{i,j}^{\min}} \right) \right]. \end{aligned} \quad (38)$$

In the same way, a similar bound for the truncated diffusion distance between $\mathbf{y}_i$ and $\mathbf{y}_j$ can be derived. □

The dimension r is typically determined by looking at the decay rate of the eigenvalues and on choosing a threshold level δ . In this section, we demonstrated how the dimension r and the threshold δ are related to the approximation of the diffusion distance. By removing all coordinates with eigenvalues smaller than δ , the error of the diffusion distance is bounded according to Lemma 3. The decaying property of the eigenvalues could be quantified by the numerical rank of the Gaussian kernel matrix. Next, we present some existing results that shed some light onto the practicality of the truncation of the diffusion coordinates up to dimension r . In practice, there usually is no single optimal value for r , and the number of necessary coordinates depends on the application.

5.4.1. Intrinsic dimension and numerical rank

In [41] Bermanis et al. prove that the numerical rank of a Gaussian kernel matrix is typically independent of its size. Their results state that the numerical rank is proportional to the volume of the data. By using boxes with side-length of $\epsilon^{D/2} = \sigma^D$ (D is the dimension of the data), it is shown that the numerical rank is bounded from above by the minimal number of such cubes required to cover the data. As the number of eigenvalues $|\lambda_\ell| > \delta$ is proportional to the numerical rank, this means that for a prescribed δ , the number of coordinates r_δ required is independent of the number of samples N .

Practically speaking, this result is useful if the numerical rank is small. Furthermore, as there are several methods for estimating the numerical rank, using the numerical rank to choose r_δ is not very expensive. Another practical method, which we found useful for setting r without directly choosing δ , is based on estimating the intrinsic dimension of the data d . Various methods have been proposed for estimating

the intrinsic dimension of high-dimensional data. Here, we provide a brief description of a few. Popular methods, such as [42,43] apply local PCA to the data, and estimate the intrinsic dimension as the median of the number of principal components corresponding to singular values larger than some threshold. Based on our experience, we found that a method called Dimensionality from Angle and Norm Concentration (DANCo) [44] provides a robust estimation of the intrinsic dimension d . DANCo generates artificial data of dimension d and attempts to minimize the KullbackLeibler divergence between the estimated probability distribution functions (pdf-s) of the observed data and the artificially-generated samples.

5.5. Out-of-sample extension

To extend the diffusion coordinates to new data points without re-applying a large-scale eigendecomposition [32], the Nyström extension [45] is widely used. Here we formulate the extension method for a multi-view scenario. Given the data sets $\mathbf{X}$ and $\mathbf{Y}$ and new points $\tilde{\mathbf{x}} \notin \mathbf{X}$ and $\tilde{\mathbf{y}} \notin \mathbf{Y}$, we want to extend the multi-view diffusion mapping to $\tilde{\mathbf{x}}$ and $\tilde{\mathbf{y}}$ without re-applying the batch procedure. First, we describe the explicit form for the eigenvalue problem for two views $\mathbf{X}$ and $\mathbf{Y}$. The eigenvector ψ_k and its associated eigenvalue λ_k satisfy $\lambda_k \psi_k = \hat{\mathbf{P}} \psi_k$. Substituting the definition of $\hat{\mathbf{P}}$ from Eqs. (12) and (14) we get that

$$\lambda_k \psi_k = \hat{\mathbf{D}}^{-1} \begin{bmatrix} \mathbf{0}_{M \times M} & \mathbf{K}^x \cdot \mathbf{K}^y \\ \mathbf{K}^y \cdot \mathbf{K}^x & \mathbf{0}_{M \times M} \end{bmatrix} \cdot \psi_k \in \mathbb{R}^{2M},$$

due to the block form of the matrix $\hat{\mathbf{P}}$

$$\lambda_k \psi_k^x[i] = \sum_j \hat{p}(\mathbf{x}_i, \mathbf{y}_j) \psi_k^y[j] \in \mathbb{R}^M,$$

$$\lambda_k \psi_k^y[i] = \sum_j \hat{p}(\mathbf{y}_i, \mathbf{x}_j) \psi_k^x[j] \in \mathbb{R}^M,$$

where $\psi_k^x[i] \triangleq \psi_k[i], i = 1, \dots, M$ and $\psi_k^y[i] \triangleq \psi_k[i + M], i = 1, \dots, M$. The transition matrices are

$$\hat{p}(\mathbf{x}_i, \mathbf{y}_j) = \frac{\sum_s K_{i,s}^x K_{s,j}^y}{\hat{\mathbf{D}}_{i,i}^{\text{rows}}}, \text{ and } \hat{p}(\mathbf{y}_i, \mathbf{x}_j) = \frac{\sum_s K_{i,s}^y K_{s,j}^x}{\hat{\mathbf{D}}_{i,i}^{\text{cols}}}.$$

The Nyström extension is similarly obtained by computing a weighted sum of the original eigenvectors. The weights are computed by applying the kernel $\mathcal{K}$ to the extended data points, followed by row-normalization. For the proposed mapping, the extension is defined by

$$\hat{\psi}_k(\tilde{\mathbf{x}}) = \frac{1}{\lambda_k} \sum_j \hat{p}(\tilde{\mathbf{x}}, \mathbf{y}_j) \psi_k^y[j] \quad (39)$$

$$\hat{\psi}_k(\tilde{\mathbf{y}}) = \frac{1}{\lambda_k} \sum_j \hat{p}(\tilde{\mathbf{y}}, \mathbf{x}_j) \psi_k^x[j] \quad (40)$$

where the “missing” probabilities $\hat{p}(\tilde{\mathbf{x}}, \mathbf{y}_j)$ and $\hat{p}(\tilde{\mathbf{y}}, \mathbf{x}_j)$ are approximated using the original kernel function as

$$\hat{p}(\tilde{\mathbf{x}}, \mathbf{y}_j) = \sum_s \exp \left\{ -\frac{\|\tilde{\mathbf{x}} - \mathbf{x}_s\|^2}{2\sigma_x^2} \right\} \exp \left\{ -\frac{\|\mathbf{y}_s - \mathbf{y}_j\|^2}{2\sigma_y^2} \right\} \frac{1}{\hat{\mathbf{D}}_{j+M,j+M}} \quad (41)$$

$$\hat{p}(\tilde{\mathbf{y}}, \mathbf{x}_j) = \sum_s \exp \left\{ -\frac{\|\tilde{\mathbf{y}} - \mathbf{y}_s\|^2}{2\sigma_y^2} \right\} \exp \left\{ -\frac{\|\mathbf{x}_s - \mathbf{x}_j\|^2}{2\sigma_x^2} \right\} \frac{1}{\hat{\mathbf{D}}_{j+M,j+M}} \quad (42)$$

and where $\psi_k^x[j] = \psi_k[j]$ and $\psi_k^y[j] = \psi_k[m + j]$.

The new mapping vector for the new data point is then given by

$$\hat{\Psi}(\tilde{\mathbf{x}}) = [\lambda_1 \hat{\psi}_1(\tilde{\mathbf{x}}), \lambda_2 \hat{\psi}_2(\tilde{\mathbf{x}}), \lambda_3 \hat{\psi}_3(\tilde{\mathbf{x}}), \dots, \lambda_{M-1} \hat{\psi}_{M-1}(\tilde{\mathbf{x}})] \in \mathbb{R}^{M-1}. \quad (43)$$

The new coordinates in the diffusion space are approximated and the new data points $\tilde{\mathbf{x}}, \tilde{\mathbf{y}}$ have no effect on the original map's structure.

5.6. Infinitesimal generator

In the subsequent analysis we only consider kernel functions operating on the scaled difference between their two arguments, namely kernel functions which can be written as

$$\mathcal{K}(\mathbf{x}_i, \mathbf{x}_j) = K \left(\frac{\mathbf{x}_i - \mathbf{x}_j}{\sqrt{\epsilon}} \right), \quad (44)$$

where ϵ is a positive constant, and $K(\mathbf{z}) : \mathbb{R}^D \mapsto \mathbb{R}$ is a non-negative symmetric function, namely $R(\mathbf{z}) \geq 0$ and $R(\mathbf{z}) = R(-\mathbf{z}) \forall \mathbf{z} \in \mathbb{R}^D$. For example, the Gaussian kernel (see Section 2.2) $\mathcal{K}(\mathbf{x}_i, \mathbf{x}_j) = \exp \left\{ -\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma_x^2} \right\}$ satisfies this property with $\epsilon = \sigma^2$.

A family of differently normalized diffusion operators was introduced in [5]. If appropriate limits are taken such that $M \rightarrow \infty$, $\epsilon \rightarrow 0$, then from Coifman and Lafon [5] it follows that the DM kernel operator will converge to one of the following differential operators: 1. Normalized graph Laplacian; 2. Laplace-Beltrami diffusion; or 3. Heat kernel equation. These are proved in [5]. The operators are all special cases of the diffusion equation. This convergence provides not only a physical justification for the DM framework, but allows in some cases to distinguish between the geometry and the density of the data points. In this section we study the asymptotic properties of the proposed kernel $\hat{\mathbf{K}}$ (Eq. (13)), limiting the discussion to only two views, i.e. $L = 2$.

We are interested in understanding the properties of the eigenfunctions of the proposed multi-view kernel $\hat{\mathbf{P}}$ (Eqs. (13) and (14)) for two views. We assume that there is some unknown mapping $\beta : \mathbb{R}^d \rightarrow \mathbb{R}^d$ from view $\mathbf{X}$ to view $\mathbf{Y}$ that satisfies $\mathbf{y}_i = \beta(\mathbf{x}_i), i = 1, \dots, M$. Each view-specific kernel applies the same function, namely $K^x(\mathbf{z}) = K^y(\mathbf{z}) = K(\mathbf{z})$, and $K(\mathbf{z})$ is normalized such that $\int_{\mathbb{R}^d} K(\mathbf{z}) d\mathbf{z} = 1$. Note that the Gaussian kernel can be properly normalized to satisfy this requirement. The analysis relates to data points $\{\mathbf{x}_1, \dots, \mathbf{x}_M\} \in \mathbb{R}^D$ sampled from a uniform distribution over a bounded domain in $\mathbb{R}^d$. The image of the function β is a bounded domain in $\mathbb{R}^d$ with distribution $\alpha(\mathbf{z})$.

Theorem 7. *The infinitesimal generator induced by the proposed kernel matrix $\hat{\mathbf{K}}$ (Eq. (13)) after row-normalization, denoted in here as $\hat{\mathbf{P}}$, converges when $M \rightarrow \infty$, $\epsilon \rightarrow 0$ (with $\epsilon = \sigma_x^2 = \sigma_y^2$) to a “cross domain Laplacian operator”. The functions $f(\mathbf{x})$ and $g(\mathbf{y})$ converge to eigenfunctions of $\hat{\mathbf{P}}$. These functions are the solutions of the following diffusion-like equations:*

$$(\hat{\mathbf{P}}f)(\mathbf{x}_i) = g(\beta(\mathbf{x}_i)) + \epsilon \triangle \gamma(\beta(\mathbf{x}_i))/\alpha(\beta(\mathbf{x}_i)) + \mathcal{O}(\epsilon^{3/2}), \quad (45)$$

$$(\hat{\mathbf{P}}g)(\mathbf{y}_i) = f(\beta^{-1}(\mathbf{y}_i)) + \epsilon \triangle \eta(\beta^{-1}(\mathbf{y}_i))/\alpha(\beta(\mathbf{y}_i)) + \mathcal{O}(\epsilon^{3/2}), \quad (46)$$

where the functions γ, η are defined as $\gamma(\mathbf{z}) \triangleq g(\mathbf{z})\alpha(\mathbf{z}), \eta(\mathbf{z}) \triangleq f(\mathbf{z})\alpha(\mathbf{z})$.

The proof of Theorem 7 is deferred to the Appendix. Although the interpretation of the result is hardly intuitive, it provides evidence that the mapping $\mathbf{f}$ and $\mathbf{g}$ are indeed coupled and are related to the second derivative on the manifold. However, the distribution of the points α has an impact on the mapping, which we hope to reduce in future research.

5.7. The convergence rate

In Theorem 7, we assume the number of data points $M \rightarrow \infty$ while the scale parameter $\epsilon \rightarrow 0$. In practice we cannot expect to have an infinite number of data points. It was shown in [5,46] (and elsewhere) that a single-view graph Laplacian converges to the laplacian operator on a manifold. It is demonstrated in [47,48] that the variance of the error for such an operator decreases as $M \rightarrow \infty$, but increases as $\epsilon \rightarrow 0$. The study in [47] proves that for a uniform distribution of data points, the variance of the error is bounded by $\mathcal{O}(\frac{1}{M^{1/2}\epsilon^{1+d/4}}, \epsilon^{1/2})$. This bound was improved in [48] by an asymptotic factor of $\sqrt{\epsilon}$ based on the correlation between $\mathbf{D}^{-1}$ and $\mathbf{K}$.

We now turn our attention to the variance of the multi-view kernel for a finite number of points. Given $\mathbf{x}_1, \dots, \mathbf{x}_M$ independent uniformly

distributed data points sampled from a bounded domain in $\mathbb{R}^d$, define the multi-view Parzen Window density estimator by

$$\hat{K}_{M,\epsilon}(\mathbf{x}) \triangleq \frac{1}{M^2} \sum_{\ell=1}^M \sum_{j=1}^M \frac{1}{\epsilon^d} K\left(\frac{\mathbf{x} - \mathbf{x}_\ell}{\sqrt{\epsilon}}\right) K\left(\frac{\mathbf{y}_\ell - \mathbf{y}_j}{\sqrt{\epsilon}}\right). \quad (47)$$

We are interested in finding a bound for the variance of $\hat{K}_{M,\epsilon}(\mathbf{x})$ for a finite number of data points:

$$\begin{aligned} \text{var}(\hat{K}_{M,\epsilon}(\mathbf{x})) &= \frac{1}{M^4 \epsilon^{2d}} \cdot M \cdot \text{var}\left(\sum_{\ell=1}^M K\left(\frac{\mathbf{x} - \mathbf{x}_\ell}{\sqrt{\epsilon}}\right) K\left(\frac{\mathbf{y}_\ell - \mathbf{y}_j}{\sqrt{\epsilon}}\right)\right) \\ &\leq \frac{1}{M^4 \epsilon^{2d}} \cdot M^3 \cdot \text{var}\left(K_\epsilon^x K_\epsilon^y\right) \leq \frac{1}{M \epsilon^{2d}} \left[\text{var}\left(K_\epsilon^x\right) \cdot \|K_\epsilon^y\|_\infty + \text{var}\left(K_\epsilon^y\right) \cdot \|K_\epsilon^x\|_\infty \right] \\ &\leq \frac{1}{M \epsilon^{2d}} \cdot [\epsilon^{d/2} \cdot m_1 \cdot 1 + \epsilon^{d/2} \cdot m_2 \cdot \alpha(\mathbf{x})] \leq \frac{m_1 + m_2 \cdot \alpha(\mathbf{x})}{M \cdot \epsilon^{1.5d}}, \end{aligned} \quad (48)$$

where the constants m_1 and m_2 are functions of the chosen kernels and of the pdf $\alpha(\cdot)$ of the points $\mathbf{y}_i, i = 1, \dots, M$. This bound helps to choose an optimal value for the scaling factor ϵ given the number of data points M and the intrinsic dimension d .

5.8. Generalized multi-view kernel

One can consider a more general multi-view kernel which does not preclude a transition within views $\mathbf{X}$ and $\mathbf{Y}$ in each time step. Such a kernel will take the form

$$\hat{K} = \begin{bmatrix} \eta \cdot (K^x)^2 & (1-\eta) \cdot K^x K^y \\ (1-\eta) \cdot K^y K^x & \eta \cdot (K^y)^2 \end{bmatrix}, \quad (49)$$

where the parameter $\eta \in [0, 1]$ prescribes the (implied) probability of within-view transitions. This kernel is normalized using the sum of rows diagonal matrix $\hat{D}$, such that $\hat{P} = \hat{D}^{-1} \hat{K}$. For large values of η , the kernel favors the within-view transitions, thereby sharing characteristics with the single-view diffusion process. For small values of η , the kernel tends to behave like our multi-view kernel $\hat{K}$ (Eq. (13)).

5.9. Additive noise

In this section we explore the effect of noise by following the analysis presented in [18]: Assuming that the noise is additive, we evaluate how the proposed method preserves the statistics of the latent model parameters $\Theta \in \mathbb{R}^{d \times M}$. We begin by considering the case of a single-view linear model, then consider a multi-view (linear model) and finally use the kernel trick to relax the linearity assumption.

The single-view linear model is defined by $\mathbf{x}_i = \mathbf{A}\theta_i + \xi_i, i = 1, \dots, M$, where the vectors $\theta_i \in \mathbb{R}^d, i = 1, \dots, M$ (columns of Θ) describe latent parameters, modeled as i.i.d., zero-mean random vectors with a finite covariance matrix Σ_θ , and $\xi_i \in \mathbb{R}^d$ are i.i.d., zero-mean random “noise” vectors with a finite covariance matrix Σ_ξ . The observations are $\mathbf{x}_i \in \mathbb{R}^D$ and the matrix $\mathbf{A} \in \mathbb{R}^{D \times d}$ has full column rank.

A standard approach for reducing the dimension is to apply PCA. However, in the presence of noise, even if the covariance of Θ is full rank PCA will provide a biased representation of Θ . In applying PCA to the set $\mathbf{X}$, one first computes the sample covariance $\frac{1}{M} \mathbf{X} \mathbf{X}^T$, which at the limit $M \rightarrow \infty$ converges to $\mathbf{A} \Sigma_\theta \mathbf{A}^T + \Sigma_\xi$. This means that using the top principal components of $\mathbf{X}$ embeds the data into a biased representation of Θ . Nonetheless, by introducing a second set of measurements $\mathbf{y}_i = \mathbf{B}\theta_i + \eta_i, i = 1, \dots, M$, where $\mathbf{B} \in \mathbb{R}^{D \times d}$ is another full-rank matrix and η_i are i.i.d. zero-mean noise vectors which are all uncorrelated with the respective ξ_i (namely, $E[\xi_i \cdot \eta_i^T] = \mathbf{0}, i = 1, \dots, M$), one can remove the bias term. Essentially, the corresponding additional measurements, along with the noise decorrelation assumption, provide the information required for retrieving an unbiased representation of Θ . Removing the bias term is possible by applying CCA [49] to $\mathbf{X}$ and $\mathbf{Y}$. CCA not only removes the bias, but enables to uncover a reduced representation that preserves the statistics of θ . This is because the sample cross covariance

of $\mathbf{X}$ and $\mathbf{Y}$ is an unbiased estimate of $\mathbf{A} \Sigma_\theta \mathbf{B}^T$. The top d right- and left-singular vectors of $\frac{1}{M} \mathbf{X} \mathbf{Y}^T$ are used to embed the data into a d dimensional space, such that the statistics of Θ are preserved. By denoting the top d singular vectors as $(\hat{\mathbf{U}}, \hat{\mathbf{S}}, \hat{\mathbf{V}}^T) = \text{SVD}_d(\frac{1}{M} \mathbf{X} \mathbf{Y}^T)$, the reduced representations are defined as $\hat{\mathbf{U}}^T \mathbf{X}$ and $\hat{\mathbf{V}}^T \mathbf{Y}$, respectively.

Although CCA is an effective, powerful tool in this context, it is based on a linear model, and thus limited to linear transformations of the data. A widely used solution to capture non-linear relations in the data are Kernel matrices [50–52]. Using a kernel matrix in the ambient space is a natural way to capture the sample-covariance in some unknown high-dimensional feature space [53]. Therefore, using a kernel generalizes the results provided by CCA to the nonlinear case. We now demonstrate how under mild assumptions the proposed kernel $\mathbf{K}^z = \mathbf{K}^x \mathbf{K}^y$ mitigates the bias effect of the additive uncorrelated noise.

Assume that $\mathbf{x}_i, \mathbf{y}_i, i = 1, \dots, M$ are noisy observations of the same low-dimensional latent variable θ_i , such that $\mathbf{x}_i = \mathbf{r}(\theta_i) + \xi_i$ and $\mathbf{y}_i = \mathbf{h}(\theta_i) + \eta_i$. The functions $\mathbf{r}, \mathbf{h} : \mathbb{R}^d \rightarrow \mathbb{R}^D$. The affinity values of $\mathbf{K}^x, \mathbf{K}^y$ are the sample covariance in two high-dimensional feature spaces Γ, Δ . The inaccessible feature maps for both views are defined by $\gamma(\mathbf{x}_i)$ and $\delta(\mathbf{y}_i)$. This means that by computing the SVD of $\mathbf{K}^z = \mathbf{K}^x \mathbf{K}^y$, we would essentially be applying CCA in the inaccessible features spaces Γ, Δ . This implies that if the following conditions hold

1. The functions $\mathbf{r}, \mathbf{h} \in \mathbb{C}^\infty$ are invertible.
2. The feature representations $\Gamma(\mathbf{X})$ and $\Delta(\mathbf{Y})$ are centered.
3. The noise terms are uncorrelated: $E[\xi_i \cdot \eta_i^T] = \mathbf{0}, i = 1, \dots, M$.
4. $E[\gamma(\mathbf{x}_i)|\theta_i] = a_0 \theta_i + a_1$, with some real valued constants a_0 and a_1 .
5. $E[\delta(\mathbf{y}_i)|\theta_i] = b_0 \theta_i + b_1$, with some real valued constants b_0 and b_1 .

then, applying SVD to $\mathbf{K}^z$ enables to cancel out the view-specific noise terms. Condition (1) implies that the latent parameters Θ lie on a noisy manifold in both observed spaces. Based on condition (2) and on the choice of kernels, the empirical covariance matrices in the feature space are $\mathbf{K}^x = \frac{1}{M} \Gamma \Gamma^T$ and $\mathbf{K}^y = \frac{1}{M} \Delta \Delta^T$. We remind that Γ and Δ are inaccessible and the kernels are computed based on $\mathbf{X}$ and $\mathbf{Y}$. As in the linear case, using $\mathbf{K}^x$ or $\mathbf{K}^y$ alone is insufficient for obtaining an unbiased estimate of the representation of Θ .

We now turn our attention to demonstrate that the matrix $\mathbf{K}^z = \mathbf{K}^x \mathbf{K}^y = \frac{1}{M^2} \Gamma \Gamma^T \Delta \Delta^T$ enables to remove the uncorrelated additive noise. A right eigenvector of $\mathbf{K}^z$, denoted by τ_i , satisfies

$$\mathbf{K}^z \tau_i = \frac{1}{M^2} \Gamma \Gamma^T \Delta \Delta^T \tau_i = \lambda_i \tau_i, \quad (50)$$

by multiplying Eq. (50) on the left by Δ^T we get

$$\frac{1}{M^2} \Delta^T \Gamma \Gamma^T \Delta \Delta^T \tau_i = \lambda_i \Delta^T \tau_i. \quad (51)$$

Substituting $\mathbf{u}_i = \Delta^T \tau_i$, $\bar{\Sigma}_{\gamma\delta} = \frac{1}{M} \Gamma^T \Delta$ and $\bar{\Sigma}_{\delta\gamma} = \frac{1}{M} \Delta^T \Gamma$ yields

$$\bar{\Sigma}_{\delta\gamma} \bar{\Sigma}_{\gamma\delta} \mathbf{u}_i = \lambda_i \mathbf{u}_i, \quad (52)$$

so up to scaling the left- and right-singular vectors of $\bar{\Sigma}_{\gamma\delta}$ provide a low-dimensional representation that captures the statistics of Θ . Thus, based on conditions (3–5), by taking the number of points to infinity, one can extract an unbiased estimate of the representation of Θ .

6. Experimental results

In this section we present experimental results to evaluate our proposed framework. First we empirically evaluate the theoretical properties derived in Section 5. Then, we demonstrate how the proposed framework can be used for learning coupled manifolds even in the presence of noise.

6.1. Empirical evaluations of theoretical aspects

In the first group of experiments we provide empirical evidence corroborating the theoretical analysis from Section 5.

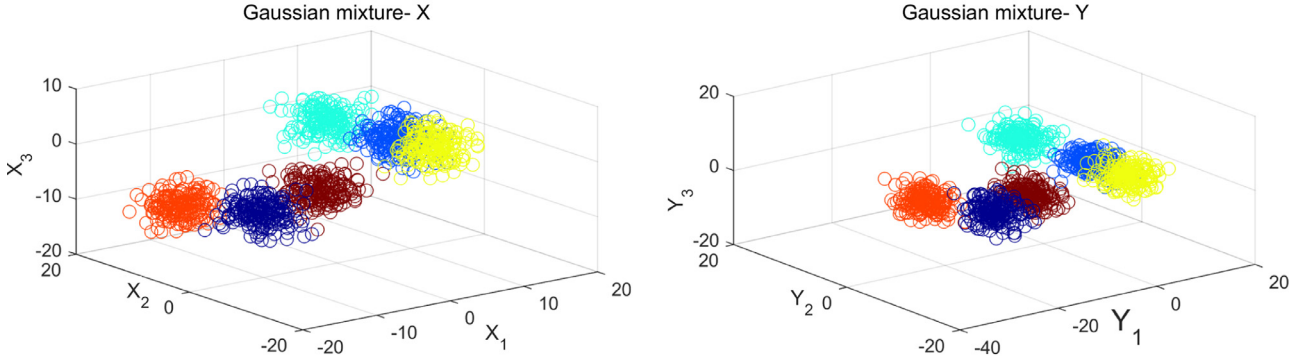


Fig. 4. The first 3 dimensions of the Gaussian mixture. Both views share the center of masses of the Gaussian spread. Left: first view denoted as $\mathbf{X}$. Right: second view denoted as $\mathbf{Y}$. The variance of the Gaussian in each dimension is 8.

6.1.1. Spectral decay

In Section 5.4, an upper bound on the eigenvalues' rate of decay for our multi-view-based approach (matrix $\hat{\mathbf{P}}$ Eq. (14)) was presented. In order to empirically evaluate the spectral decay for $\hat{\mathbf{P}}$, $\mathbf{P}^*$ (Eq. (6) and [5]) and $\mathbf{P}^+$, we generated synthetically-clustered data drawn from Gaussian-Mixtures distributions. The following steps describe the generation of both views, denoted $(\mathbf{X}, \mathbf{Y})$ and referred to as View-I ($\mathbf{X}$) and View-II ($\mathbf{Y}$), resp.:

1. Six vectors $\mu_j \in \mathbb{R}^9$, $j = 1, \dots, 6$ were drawn from a Gaussian distribution $N(\mathbf{0}, 8 \cdot \mathbf{I}_{9 \times 9})$. These vectors would serve as the centers of masses of the generated classes.
2. One hundred data points were drawn for each cluster $j = 1, \dots, 6$ from a Gaussian distribution $N(\mu_j, 2 \cdot \mathbf{I}_{9 \times 9})$. Denote these 600 data points $\mathbf{X} \in \mathbb{R}^{9 \times 600}$.
3. One hundred additional data points were similarly drawn from each of the six Gaussian distributions $N(\mu_j, 2 \cdot \mathbf{I}_{9 \times 9})$. Denote these 600 data points $\mathbf{Y} \in \mathbb{R}^{9 \times 600}$.

The first 3 dimensions of both views are depicted in Fig. 4. We compute the probability matrix for each view $\mathbf{P}^x$ and $\mathbf{P}^y$, the Kernel Sum approach probability matrix $\mathbf{P}^+$, the Kernel Product approach $\mathbf{P}^*$ (Eq. (6)) and the proposed approach $\hat{\mathbf{P}}$. The eigendecomposition is computed for all matrices. The resulting eigenvalues' decay rate are compared with the eigenvalues product from both views. To get a fair comparison between all the methods, we set the Gaussian scale parameters σ_x and σ_y in each view and then use these scales in all the methods. The vectors' variance in the concatenation approach is the sum of variances since we assume statistical independence. Therefore, the following scale parameters $\sigma_o^2 = \sigma_x^2 + \sigma_y^2$ are used.

The experiment is repeated but this time $\mathbf{X}$ contains 6 clusters whereas $\mathbf{Y}$ contains only 3. For $\mathbf{Y}$, we use only the first 3 centers of masses and generate 200 points in each cluster. Fig. 5 presents a logarithmic scale of the spectral decay for eigenvalues extracted from all methods. It is evident that our proposed kernel has the strongest spectral decay.

6.1.2. Cross view diffusion distance

In this section, we examine the proposed Cross View Diffusion Distance (Section 5.3). A Swiss Roll is generated by using the function

$$\text{View I: } \mathbf{x}_i = \begin{bmatrix} x_i[1] \\ x_i[2] \\ x_i[3] \end{bmatrix} = \begin{bmatrix} 6\theta_i \cos(\theta_i) \\ h_i \\ 6\theta_i \sin(\theta_i) \end{bmatrix} + \mathbf{n}_i^{(1)}, \quad (53)$$

with $\theta_i = (1.5\pi)s_i$, $i = 1, 2, \dots, 1,000$, where s_i are 1000 data points evenly spread along the segment $[1,3]$, and where $\mathbf{n}_i^{(1)}$ are i.i.d. zero-mean Gaussian noise vectors with covariance $\sigma_N^2 \cdot \mathbf{I}_{3 \times 3}$. The second view is generated by applying an orthonormal transformation to the (noiseless) Swiss

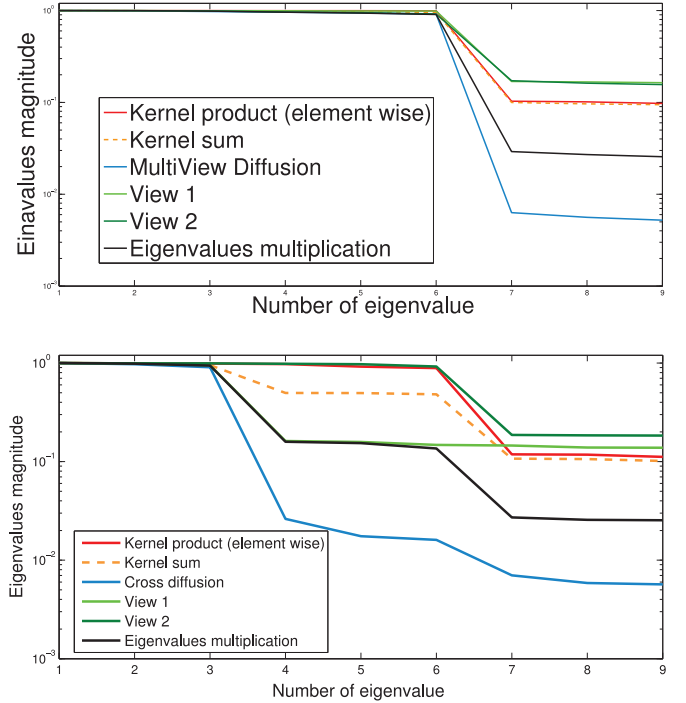


Fig. 5. Eigenvalues decay rate. Comparison between different mapping methods. Top: 6 clusters in each view. Bottom: 6 clusters in $\mathbf{X}$ and 3 clusters in $\mathbf{Y}$.

Roll and similarly adding Gaussian noise:

$$\text{View II: } \mathbf{y}_i = \begin{bmatrix} y_i[1] \\ y_i[2] \\ y_i[3] \end{bmatrix} = \mathbf{R} \begin{bmatrix} 6\theta_i \cos(\theta_i) \\ h_i \\ 6\theta_i \sin(\theta_i) \end{bmatrix} + \mathbf{n}_i^{(2)}, \quad (54)$$

where $\mathbf{R} \in \mathbb{R}^{3 \times 3}$ is a random orthonormal transformation matrix, and where $\mathbf{n}_i^{(2)}$ are i.i.d. $N(\mathbf{0}, \sigma_N^2 \cdot \mathbf{I}_{3 \times 3})$. The matrix $\mathbf{R}$ is generated by independently drawing its elements from a standard Gaussian distribution, followed by applying the Gram-Schmidt orthogonalization procedure. The variables h_i , $i = 1, \dots, 1000$ are drawn from a uniform distribution in the interval $[0,100]$. An example for both Swiss Rolls is shown in Fig. 6.

A standard DM is applied to each view and a 2-dimensional embedding of the Swiss Roll is extracted. The sum of distances between all the data points in the embedding spaces is denoted as a single-view diffusion distance (SVDD). The distance is computed using the measure

$$D_t^{(\text{SV})^2}(\mathbf{X}, \mathbf{Y}) = \sum_{i=1}^M \|\Psi_t(\mathbf{x}_i) - \Psi_t(\mathbf{y}_i)\|^2, \quad (55)$$

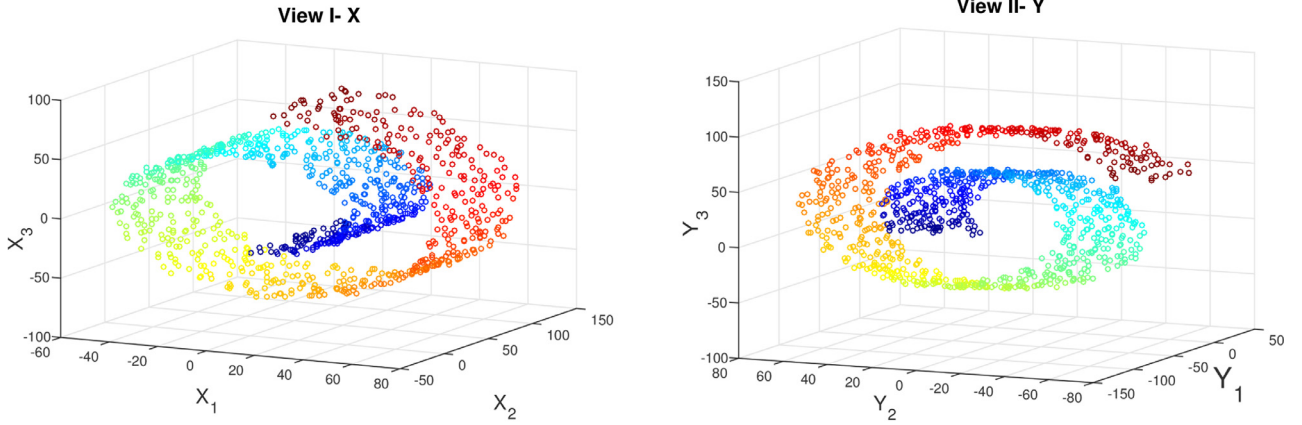


Fig. 6. The two Swiss Rolls, Left - a Swiss Roll generated by Eq. (53), Right - a Swiss Roll generated by Eq. (54).

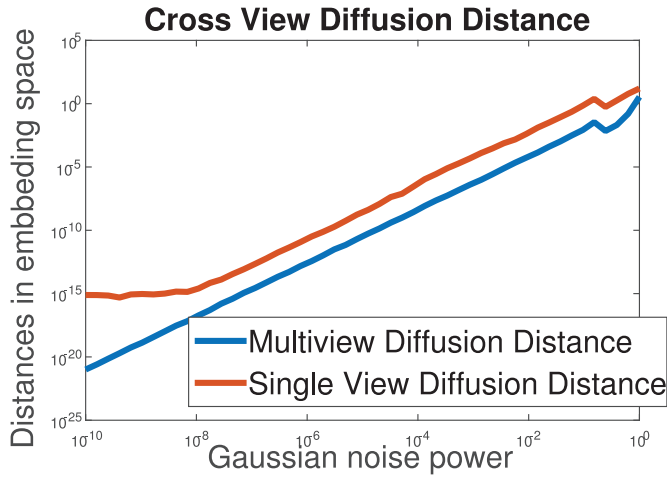


Fig. 7. Comparison between two cross view diffusion based distances. Simulated on two Swiss Rolls with additive Gaussian noise. The results are the median of 100 simulations.

where $\Psi_i(x_i), \Psi_i(y_i), i = 1, \dots, M$ are the single-view diffusion mappings. Then, the proposed framework is applied to extract the coupled embedding. A Cross View Diffusion Distance (CVDD) is computed using Eq. (30). This experiment was executed 100 times for various values of the Gaussian noise variance σ_N^2 .

In about 10% of the single-view trials the embeddings' axis are flipped. This generates a large SVDD although the embeddings share similar structures. In order to mitigate the effect of this type of errors we used the Median of the measures taken from the 100 trials, presented in Fig. 7.

6.2. Manifold learning

In this section we demonstrate how our proposed Multi-view DM framework allows to simultaneously extract L low-dimensional representations for L datasets.

6.2.1. Artificial manifold learning

The general DM approach is based on an underlying assumption, that the sampled space describes a single low-dimensional manifold. However, this assumption may be incorrect if the sampled space describes the existence of redundancy in the manifold, or more generally, if the sampled space can describe two or more manifolds generated by a common physical process. In this section we consider such cases. We examine the

extracted embedding computed using our method and compare it to the Kernel Product approach (Section 3.1).

Helix A

Two coupled manifolds with a common underlying open circular structure are generated. The helix shaped manifolds were generated by the application of a 3-dimensional function to $M = 1,000$ data points $\{a_i, b_i\}_{i=1}^M$, such that the $\{a_i\}$ are evenly spread in $[0, 2\pi]$ and $b_i = (a_i + 0.5\pi) \bmod 2\pi, i = 1, \dots, M$. The following functions are used to generate the datasets for View-I and View-II denoted as $\mathbf{X}$ and $\mathbf{Y}$, resp.:

$$\text{View I: } \mathbf{x}_i = \begin{bmatrix} x_i[1] \\ x_i[2] \\ x_i[3] \end{bmatrix} = \begin{bmatrix} 4 \cos(0.9a_i) + 0.3 \cos(20a_i) \\ 4 \sin(0.9a_i) + 0.3 \sin(20a_i) \\ 0.1(6.3a_i^2 - a_i^3) \end{bmatrix}, i = 1, 2, \dots, 1,000, \quad (56)$$

$$\text{View II: } \mathbf{y}_i = \begin{bmatrix} y_i[1] \\ y_i[2] \\ y_i[3] \end{bmatrix} = \begin{bmatrix} 4 \cos(0.9b_i) + 0.3 \cos(20b_i) \\ 4 \sin(0.9b_i) + 0.3 \sin(20b_i) \\ 0.1(6.3b_i - b_i^2) \end{bmatrix}, i = 1, 2, \dots, 1,000. \quad (57)$$

The resulting 3-dimensional Helix-shaped manifolds $\mathbf{X}$ and $\mathbf{Y}$ are shown in Fig. 8.

The Kernel Product mapping (Eq. (6)) separates the manifold to a bow and a point as shown in Fig. 10. This structure neither represents any of the original structures nor reveals the underlying parameters a_i, b_i . On the other hand, our embedding (Eq. (21)) captures the two structures. one for each view. As shown in Fig. 9, one structure represents the angle of a_i while the other represents the angle of b_i . The Euclidean distance in the new spaces preserves the mutual relations between data points based on the geometrical relation in both views. Moreover, both manifolds are in the same coordinate system and this is a strong advantage as it enables to compare the manifolds in the lower-dimensional space. The Euclidean distance in the new spaces preserves the mutual relations between data points that are based on the geometrical structure of both views.

Helix B

The previous experiment was repeated using the following alternative functions:

$$\text{View I: } \mathbf{x}_i = \begin{bmatrix} x_i[1] \\ x_i[2] \\ x_i[3] \end{bmatrix} = \begin{bmatrix} 4 \cos(5a_i) \\ 4 \sin(5a_i) \\ 4a_i \end{bmatrix}, \quad (58)$$

$$\text{View II: } \mathbf{y}_i = \begin{bmatrix} y_i[1] \\ y_i[2] \\ y_i[3] \end{bmatrix} = \begin{bmatrix} 4 \cos(5b_i) \\ 4 \sin(5b_i) \\ 4b_i \end{bmatrix}. \quad (59)$$

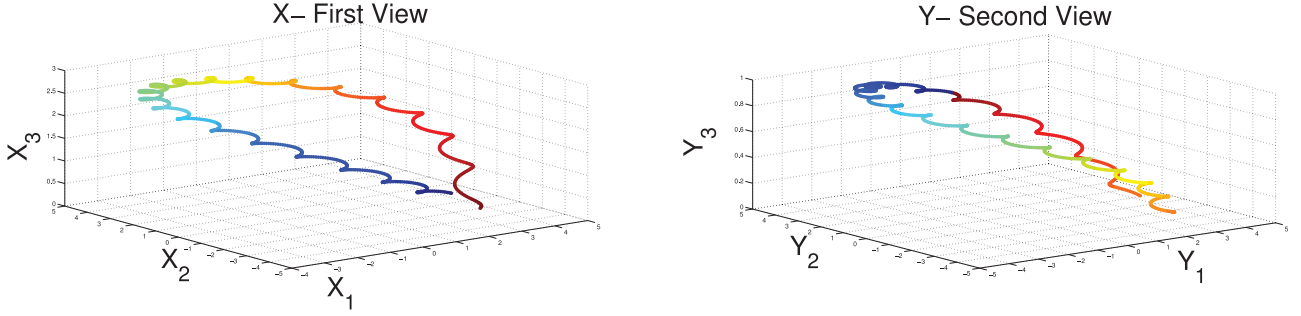


Fig. 8. Left: first Helix $\mathbf{X}$ (Eq. (56)). Right: second Helix $\mathbf{Y}$ (Eq. (57)). Both manifolds have some circular structure governed by the angle parameter $a[i]$ and $b[i]$, $i = 1, 2, \dots, 1000$ colored by the points index i .

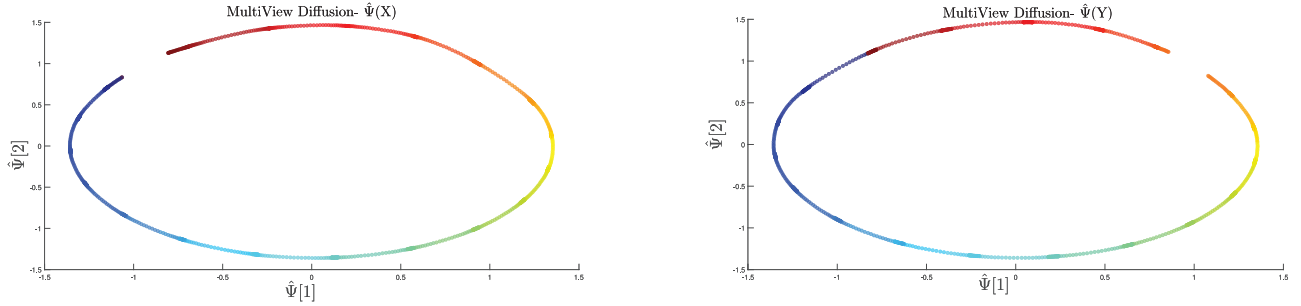


Fig. 9. Left: Multi-View based embedding of the first view $\hat{\Psi}(\mathbf{X})$. Right: Multi-View based embedding of the second view $\hat{\Psi}(\mathbf{Y})$. They were computed by using Eq. (21), respectively.

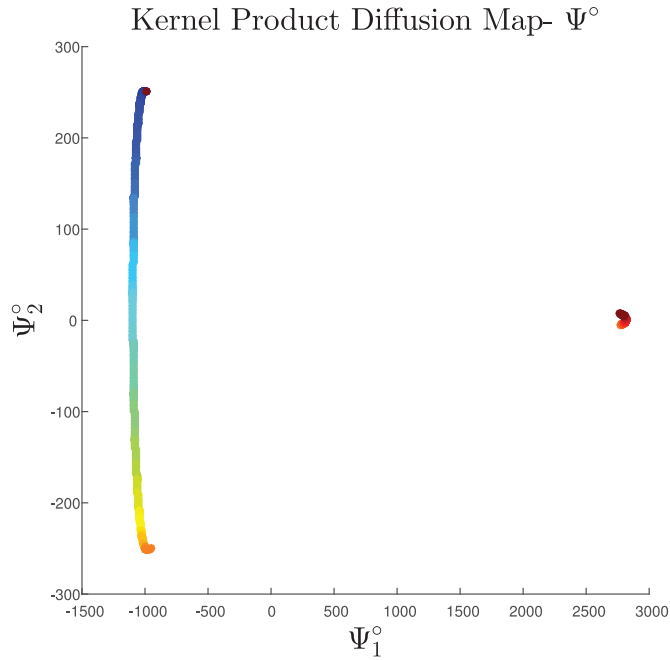


Fig. 10. 2-dimensional DM-based mapping of the Helix computed using the concatenated vector from both views that correspond to the kernel $\mathbf{P}^*$ (Eq. (6)).

Again, $M = 1000$ points were generated using $a_i \in [0, 2\pi]$, $b_i = (a_i + 0.5\pi) \bmod 2\pi$, $i = 1, \dots, M$. The generated manifolds are presented in Fig. 11.

As can be viewed in Fig. 12, our proposed embeddings (Eq. (21)) has successfully captured the governing parameters a_i and b_i . The Kernel Product based embedding (Eq. (6)), as evident in Fig. 13, again sepa-

rated the data points into two unconnected structures that do not represent well the parameters.

6.2.2. Multiview video sequence

Various examples involving datasets in diverse fields, such as images, audio, MRI ([11,54,55], resp.) have demonstrated the power of DM for the extraction of underlying changing physical parameters from real datasets. In this experiment, the multi-view approach is tested on a real video data, in what can literally be termed a “toy example”:

Two web cameras and a toy train with preset tracks are used. The train’s tracks have an “eight” shape structure. Extracting the underlying manifold from the set of images enables to organize the images according to the location along the train’s path and thus reveals the true underlying parameters of the processes.

The setting of the experiment is as follows: each camera records a set of images from a different angle. A sample frame from each view is shown in Fig. 14. The video is sampled at 30 frames per second with a resolution of 640×480 pixels per frame. $M = 220$ images were collected from each view (camera). Then, the R,G,B values were averaged and downsampled to 160×120 pixels resolution. The matrices were reshaped into column vectors. The resulted set of vectors are denoted by $\mathbf{X}$ and $\mathbf{Y}$ where $x_i, y_i \in \mathbb{R}^{19200}$, $1 \leq i \leq 220$. The sequential order of the images is not important for the algorithm. In a normal setting, one view is sufficient to extract the parameters that govern the movement of the train and thus extract the natural order of the images. However, we use two types of interferences to create a scenario in which each view by itself is insufficient for the extraction of the underlying parameters. The first interference is a gap in the recording of each camera. We remove 20 consecutive frames from each view at different time intervals. By removing frames, the bijective correspondence of some of the images in the sequence is broken. However, even an approximated correspondence is sufficient for our proposed manifold extraction. A standard 2-dimensional DM based mapping of each view was extracted. The results are bow-shaped manifolds as presented in Fig. 15. Applying DM

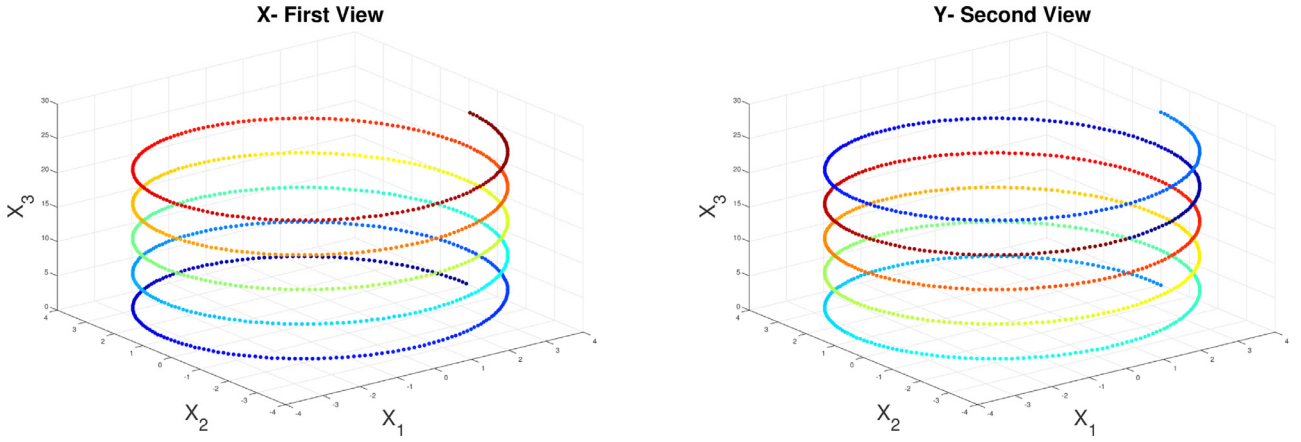


Fig. 11. Left: first Helix $\mathbf{X}$ (Eq. (58)). Right: second Helix $\mathbf{Y}$ (Eq. (59)). Both manifolds have some circular structure governed by the angle parameter $a[i]$ and b_i , $i = 1, 2, \dots, 1000$, as colored by the point's index i .

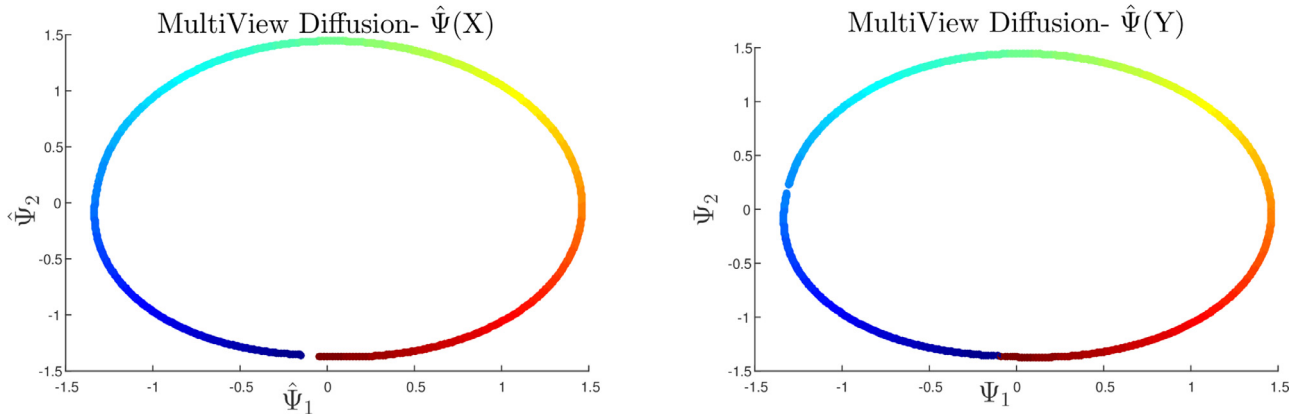


Fig. 12. The coupled mappings computed using our proposed parametrization in Eq. (21).

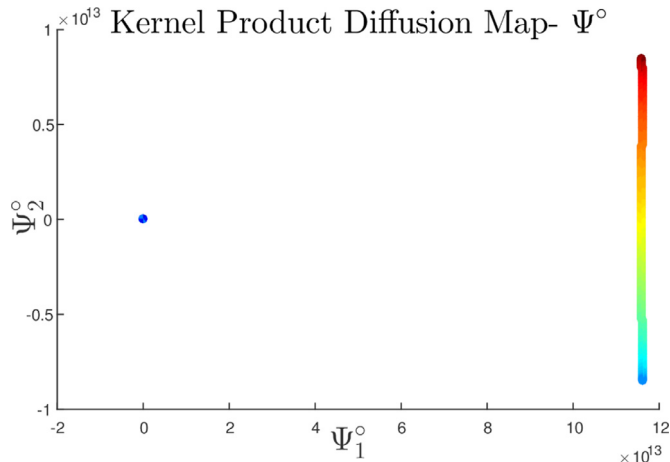


Fig. 13. A 2-dimensional mapping, extracted based on $\mathbf{P}^\circ$ (Eq. (6)).

separately to each view extracts the correct order of the data points (images) along the path. However, the “missing” data points broke the circular structure of the expected manifold and resulted in a bow-shaped embedding. We use the multi-view based methodology to overcome this interference by application of the multi-view framework to extract two coupled mappings (Eq. (21)). The results are shown in Fig. 16. The pro-

posed approach overcomes the interferences by smoothing out the gap inherited in each view through the use of connectivities from the “unobstructed” view. Finally, we concatenate the vectors from both views and compute the Kernel Product embedding. The results are presented in Fig. 17. Again, the structure of the manifold is distorted and incomplete due to the missing images.

This experiment was then repeated, replacing 10 frames from each view with “noise frames” consisting of pixel-wise i.i.d. zero-mean Gaussian noise with variance 10. A single-view DM-based mapping was computed. The Kernel Product-based DM and the multi-view based DM mappings were computed as well. As presented in Fig. 18, the Gaussian noise distorted the manifolds extracted in each view. The multi-view approach extracted two circular structures presented in Fig. 19. Again, the data points are ordered according to the position along the path. This time, the circular structure is unfolded and the gaps are visible in both embeddings. Applying the Kernel Product approach (Eq. (6)) has yielded a distorted manifold as presented in Fig. 20.

7. Applications

7.1. Multi-view clustering

The task of clustering has been in the core of machine learning for many years. The goal is to divide a given dataset into subsets based on the inherited structure of the data. We use the multi-view construction to extract low-dimensional mappings from multiple sets of



Fig. 14. Left: a sample image from the first camera (X). Right: a sample image from the second camera (Y).

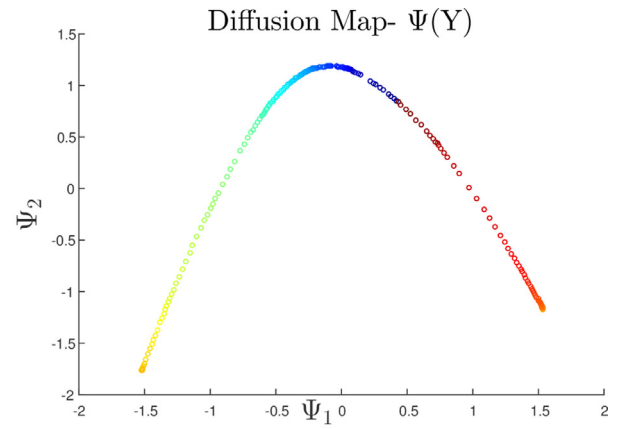
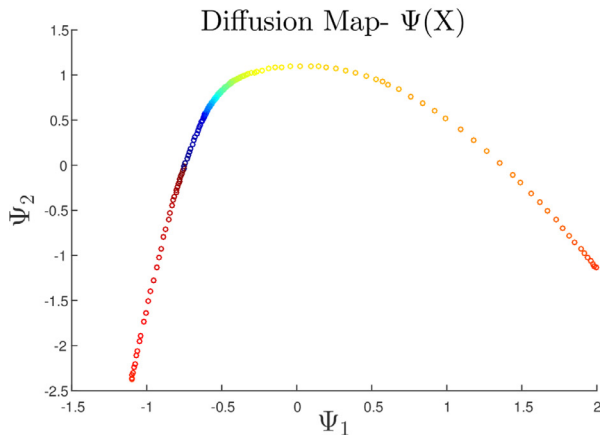


Fig. 15. Left: DM-based single-view mapping $\Psi(X)$. Right: DM-based single-view mapping $\Psi(Y)$. The removed images caused a bow shaped structure.

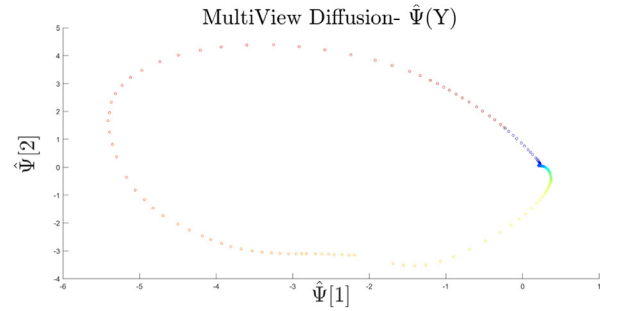
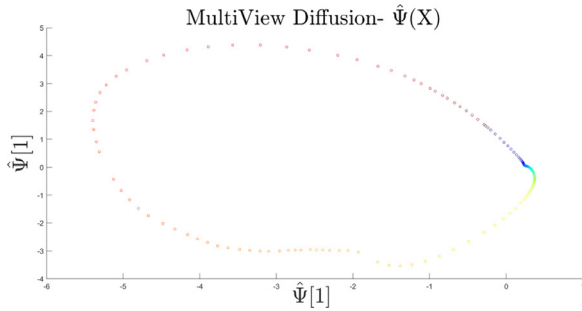


Fig. 16. Left: Mapping $\hat{\Psi}(X)$. Right: Mapping $\hat{\Psi}(Y)$ as extracted by the multi-view based framework. Two small gaps, which correspond to the removed images, are visible.

high-dimensional data points. In the following experiments we expand the examples presented in [56] to cluster artificial and real data sets. For the real data sets applying the multi-view approach requires an eigen decomposition of large matrices. To reduce the runtime of experiments we use an approximate matrix decomposition based on sparse random projections [35].

7.1.1. Two circles clustering

Spectral properties of data sets are useful for clustering since they reveal information about the unknown number of clusters. The characteristic of the eigenvalues of $\hat{P}$ (Eq. (14)) can provide insight into the number of clusters within the data set. The study in [57] relates the

number of clusters to the multiplicity of the eigenvalue 1. A different approach in [58] provides an analysis about the relation between the eigenvalue drop to the number of clusters. In this section, we evaluate how our proposed method captures the clusters' structure when two views are available.

We generate two circles that represent the original clusters using the function

$$z_i = \begin{bmatrix} z_i[1] \\ z_i[2] \end{bmatrix} = \begin{bmatrix} r \cdot \cos(\theta_i) \\ r \cdot \sin(\theta_i) \end{bmatrix}, \quad (60)$$

where $M = 1,600$ points θ_i , $1 \leq i \leq M$, are evenly spread in $[0, 4\pi]$. The clusters are created by changing the radius as follows:

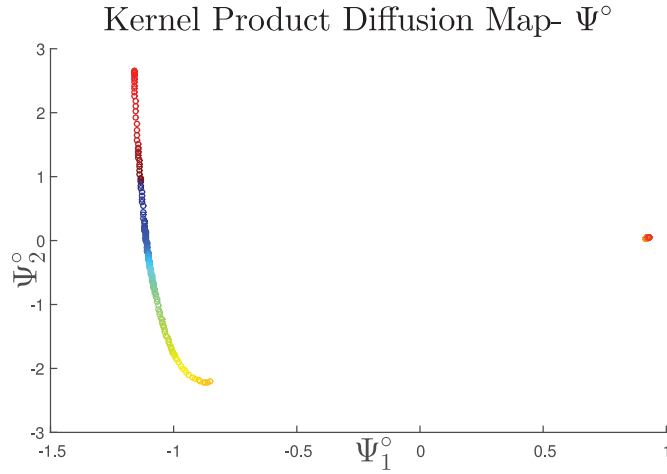


Fig. 17. A standard diffusion mapping (Kernel Product-based) that was computed by using the concatenated vector from both views that correspond to kernel $\mathbf{K}$.

$r = 2, 1 \leq i \leq 800$ (first cluster), $r = 4,801 \leq i \leq 1,600$ (second cluster). The views $\mathbf{X}$ (Eq. (61)) and $\mathbf{Y}$ (Eq. (62)) are generated by the application of the following non-linear functions that produce the distorted views

$$x_i[1] = \begin{cases} z_i[1] + 1 + n_i[2] | z_i[2] \geq 0 \\ z_i[1] + n_i[3] | z_i[2] < 0 \end{cases}, x_i[2] = z_i[2] + n_i[1] \quad (61)$$

and

$$y_i[1] = z_i[1] + n_i[4], y_i[2] = \begin{cases} z_i[2] + 1 + n_i[6] | z_i[1] \geq 0 \\ z_i[2] + n_i[6] | z_i[1] < 0 \end{cases}, \quad (62)$$

where $n_i[\ell]$, $1 \leq \ell \leq 6$, are i.i.d. random variables drawn from a Gaussian distribution with $\mu = 0$ and $\sigma_n^2 \in [0.03, 0.6]$. This data is referred to as the Coupled Circles dataset.

In Fig. 21, the views $\mathbf{X}$ and $\mathbf{Y}$, which were generated by Eqs. (61) and (62), are shown. Color and shape indicate the ground truth clusters. Initially, DM is applied to each view and clustering is applied using K-means ($K = 2$) within the first diffusion coordinate. The kernel bandwidths σ_x and σ_y for all methods are set using the min-max method described in Eq. (24). We use $t = 1$ since it is optimal for clustering tasks. For the kernel product method we use $\sigma_o = \sqrt{\sigma_x^2 + \sigma_y^2}$. We further extract a 1-dimensional representation using the proposed multi-view framework (Eq. (21)), the Kernel Sum DM (Eq. (7)), Kernel Product

Table 1

Accuracy of clustering and the normalized mutual information (NMI) on the infinity MNIST dataset. Each view consists a noisy version based on 40K images of handwritten 2's and 3's. The first view $\mathbf{X}^1$ is generated by adding Gaussian noise, while $\mathbf{X}^2$ is generated by zeroing out random pixels with probability 0.5.

Method	NMI	Accuracy
DM $\mathbf{X}^1$	0.38	84.4%
DM $\mathbf{X}^2$	0.38	84.1%
Kernel Prod	0.41	85.2%
Kernel Sum	0.39	84.6%
Kernel CCA	0.54	90.5%
de Sa	0.59	91.2%
Multiview	0.70	94.7%

DM (Eq. (6)), de Sa's approach (Eq. (10)) and Kernel CCA (Eq. (8)) described in Section 3. The regularization parameter is $\gamma = 0.01$ for KCCA and we use 100 components for the Incomplete Cholesky Decomposition [15,16]. Clustering is performed in the representation space by the application of K-means where $K = 2$. To evaluate the performance of our proposed map 100 simulations with various values of the Gaussian's noise variance (all with zero mean) were performed. The average clustering success rate is presented in Fig. 22. It is evident that the multi-view based approach outperforms the DM-based single-view and the Kernel Product approaches.

The performance of kernel methods is highly dependent on setting an appropriate kernel bandwidth σ_x, σ_y , in Algorithm 1 we have presented a method for setting such parameters. To evaluate the influence of these parameters on the clustering quality we set $\sigma_n = 0.16$ and extract the multi-view, Kernel Sum and Kernel Product diffusion mapping for various values of σ_x, σ_y . The average clustering performance using K-means ($K = 2$) are presented in Fig. 23.

7.1.2. Handwritten digits

For the following clustering experiment, we use the Multiple Features database² from the UCI repository. The data set consists of 2,000 handwritten digits from 0 to 9 that are equally spread. The extracted features from these images are the profile correlations (FAC), Karhunen-Loève coefficients (KAR), Zerkine moment (ZER), morphological (MOR), pixel averages in 2×3 windows and the Fourier coefficients (Fou) as our feature spaces $\mathbf{X}^1, \mathbf{X}^2, \mathbf{X}^3, \mathbf{X}^4, \mathbf{X}^5, \mathbf{X}^6$ respectively. We apply dimensionality reduction using a single-view DM, Kernel Product DM, Kernel Sum DM and the proposed Multi-view. We apply K-means to the

² <http://archive.ics.uci.edu/ml>.

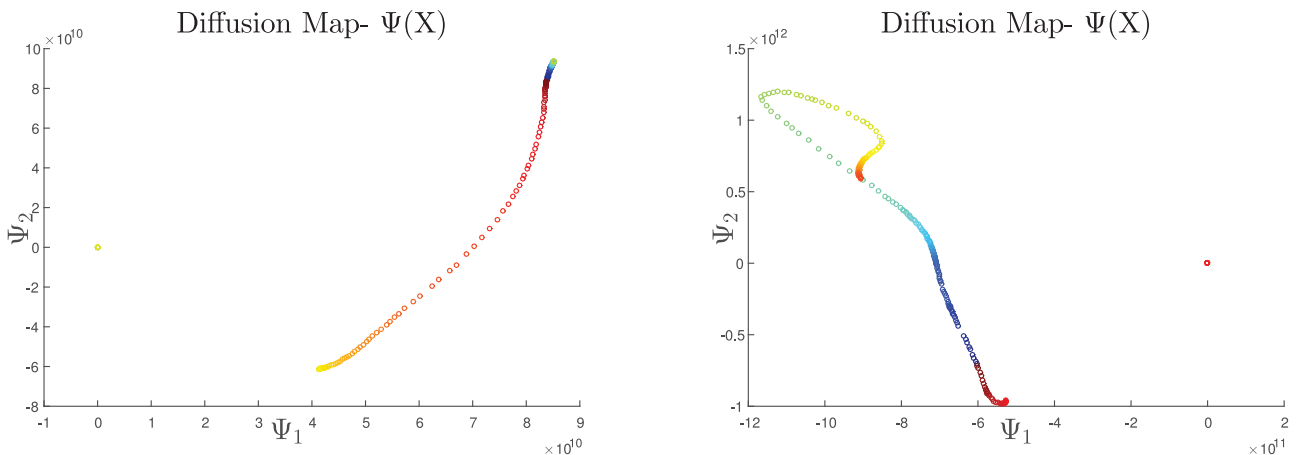


Fig. 18. Left: DM-based single-view mapping $\Psi(\mathbf{X})$. Right: DM-based single-view mapping $\Psi(\mathbf{Y})$. The Gaussian noise deformed the circular structure.

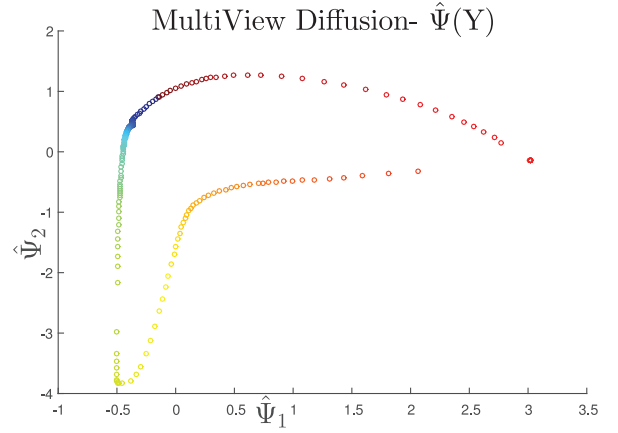
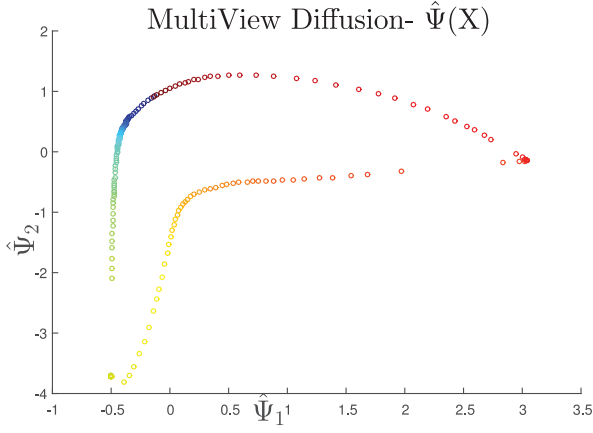


Fig. 19. Left: Mapping $\hat{\Psi}(X)$. Right: Mapping $\hat{\Psi}(Y)$ as extracted by the multi-view framework. Two gaps are visible that correspond to Gaussian noise.

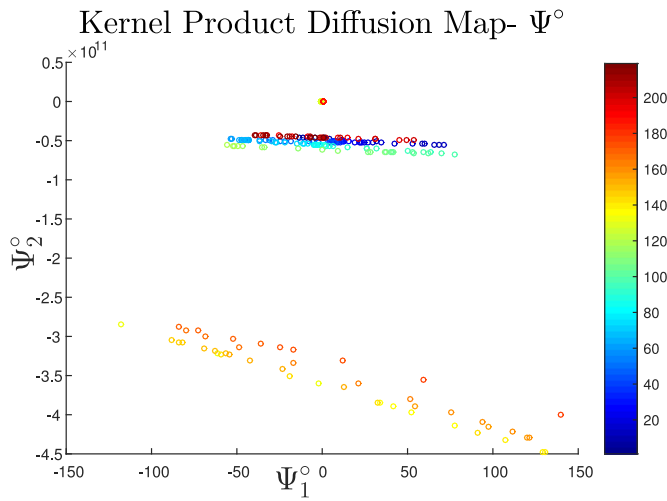


Fig. 20. Computation of a standard diffusion mapping (Kernel Product) by using the concatenation vector from both views (corresponding to kernel P^* Eq. (6)).

duced mapping using 6 to 20 coordinates. The clustering performance is measured using the Normalized Mutual Information [59] (NMI). Fig. 24 presents the average clustering results using K-Means.

Next, we attempt to cluster 40K grey scale images of handwritten 2's and 3's. The images with dimension 28×28 were collected from the infinity MNIST dataset [60]. We generate two independent noisy versions of the 40K samples. By adding pixel Gaussian noise $N(0, 0.5)$ to each image we create the first noisy view which is denoted as X^1 . The second view X^2 is created by randomly and independently zeroing each pixel with probability 0.5. In Fig. 25. we present 25 examples from

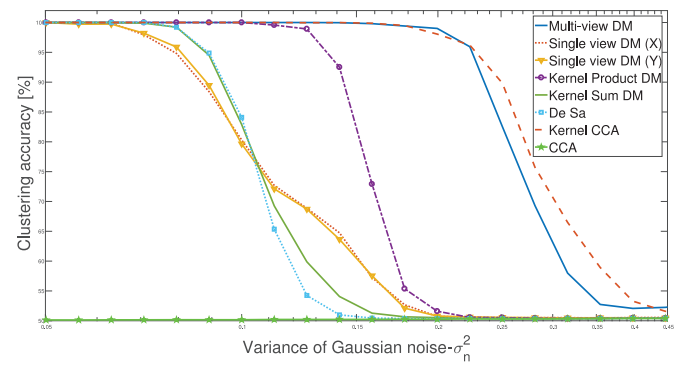


Fig. 22. Clustering results from averaging 200 trials vs. the variance of the Gaussian noise. The simulation performed on the Coupled Circles data (Eqs. (60)–(62)).

each view. To evaluate the clustering performance, we apply K-means 100 times to the reduced representations with dimensions 5,10,15 and 20 and report the top results for each method. In Table 1 we present the top Normalized Mutual Information (NMI) and clustering accuracy's the proposed multi-view and various alternative methods.

7.1.3. Isolet data set

The Isolet data set was constructed by recording 150 people pronouncing each letter twice for all 26 letters. The feature vector available is a concatenation of the following features: spectral coefficients, contour, sonorant, pre-sonorant and post-sonorant. The authors do not provide the feature's separation, therefore, the dimension of the feature vector is 617. We use a subset of the data with 1599 instances, thus the features space is $X \in \mathbb{R}^{1599 \times 617}$. To apply the multi-view approach we compute 3 different kernels and fuse them together. The first kernel K^1

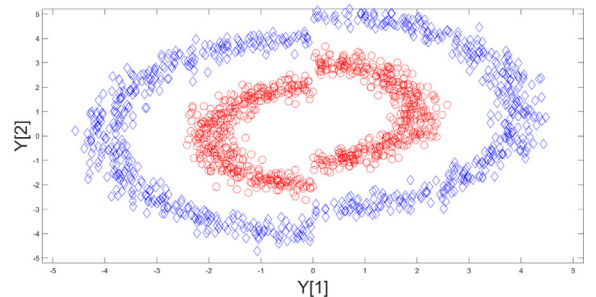
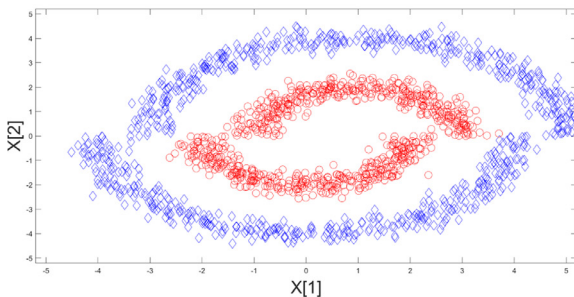


Fig. 21. Left: first view X . Right: second view Y . The ground truth clusters are represented by the marker's shape and color.

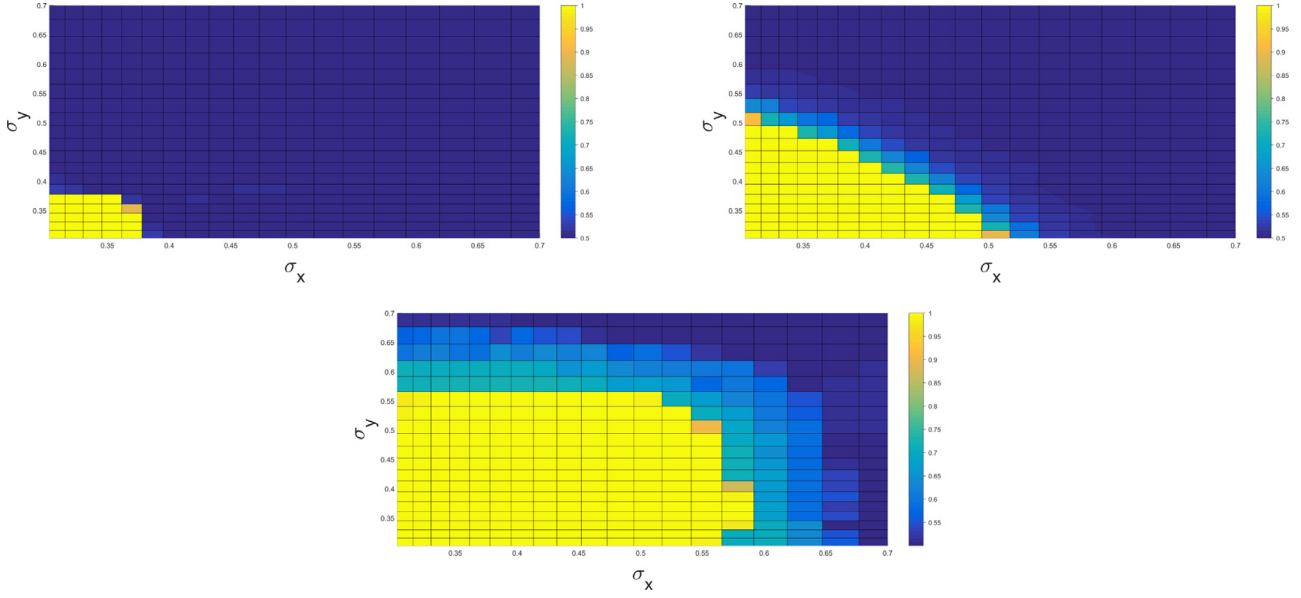


Fig. 23. Clustering results from averaging 20 trails using various values of σ_x , σ_y based on different mappings. The standard deviation of the noise is $\sigma_n = 0.16$. Top left- Kernel Sum DM, top right- Kernel Product DM, bottom- Multi View DM.

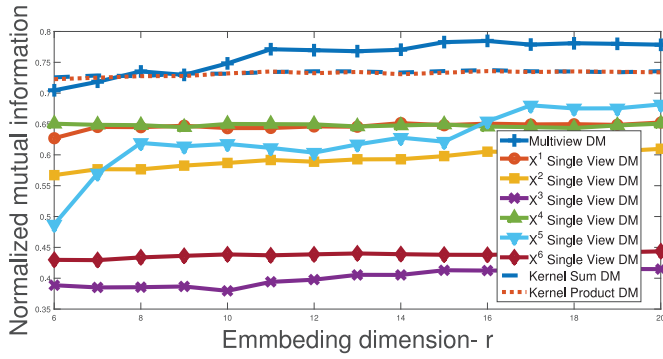


Fig. 24. Average clustering accuracy running 100 simulations on the Handwritten data set. Accuracy is measured using the Normalized Mutual Information (NMI).

is the standard Gaussian kernel defined in Eq. (12). K^2 is a Laplacian kernel defined by

$$K_{i,j}^2 \triangleq \exp\left(\frac{-|\mathbf{x}_i - \mathbf{x}_j|}{\sigma_2}\right). \quad (63)$$

The third kernel K^3 is an exponent with a correlation distance as the affinity measure, given by

$$K_{i,j}^3 \triangleq \exp\left(\frac{T_{i,j} - 1}{2\sigma^2}\right), i, j = 1, \dots, M, \quad (64)$$

where $T_{i,j}$ is the correlation coefficient between the i th and j th feature vectors, computed by

$$T_{i,j} \triangleq \frac{\tilde{\mathbf{x}}_i^T \cdot \tilde{\mathbf{x}}_j}{\sqrt{(\tilde{\mathbf{x}}_i^T \cdot \tilde{\mathbf{x}}_i)(\tilde{\mathbf{x}}_j^T \cdot \tilde{\mathbf{x}}_j)}}, i, j = 1, \dots, M. \quad (65)$$

The average subtracted features are $\tilde{\mathbf{x}}_i \triangleq \mathbf{x}_i - \eta_i \cdot \mathbf{1}$, where η_i is the average of the features for instance i . We fuse the kernels using multi-view, kernel product and kernel sum approach, we then apply K-Means to the extracted space. The average NMI for 26 classes is presented in Fig. 26.

Table 2

Clustering results on two subsets of Caltech-101 data set. Columns represent measures for the quality of clustering. First two rows are scores based on a single-view DM using different features. The last four rows present scores based on the proposed and alternative schemes for a Multiview DM based mapping.

Method	Silhouette(10)	Calinski(10)	Silhouette(15)	Calinski(15)
DM X^1 (SIFT)	-0.27	183	-0.27	293
DM X^2 (PHOG)	-0.23	169	-0.19	329
Kernel Prod	-0.08	348	-0.02	683
Kernel Sum	-0.11	322	-0.05	650
de Sa	-0.35	134	-0.34	227
Multiview	0.32	727	0.32	1275

7.1.4. Caltech 101

For this experiment we use an image dataset which consists of 101 categories collected in [61] for an object recognition problem. We use two subsets of the dataset, with 10 and 15 instances in each category. The views X^1 , X^2 are represented by the following features: Bag-of-words SIFT descriptors [62] and Pyramid Histogram of Oriented Gradients (PHOG) [63]. We use Kernel matrices computed by [64], therefore, we do not set the scale parameters. The quality of the clustering is again measured using Calinski-Harabasz Criterion [65] and the average silhouette width [66], for both metrics higher values indicates better separation. The silhouette is defined within the range of $[-1, 1]$. The average clustering results based on 10 to 15 coordinates are presented in Table 2.

7.2. Learning from multi sensor seismic data

Automatic detection and identification of seismic events is an important task, it is carried out constantly for seismic and nuclear monitoring. The monitoring process results in a seismic event bulletin that contains information about the detected events, their locations, magnitudes and type (natural or man made event). Seismic stations usually consist of multiple sensors recording continuously at a low frequency. The amount of available data is huge and only a fraction of the recordings contains the signal of interest. Thus, automatic tools for monitoring are of great interest. A suspect event is usually identified based on the energy of the signal, then, typical discrimination algorithms extract seismic parameters. A simple seismic parameter is the focal depth. Its drawback is

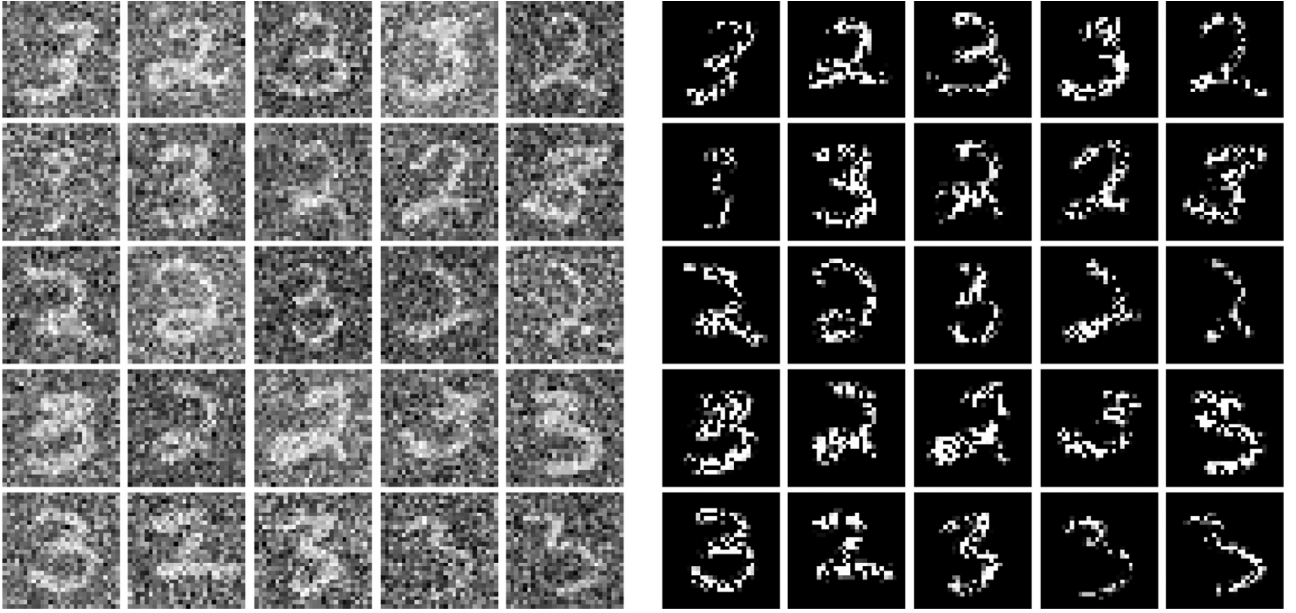


Fig. 25. Random samples from both views. Left- X^1 generated by adding Gaussian noise. Right- X^2 generated by randomly dropping out 50% of the pixels.

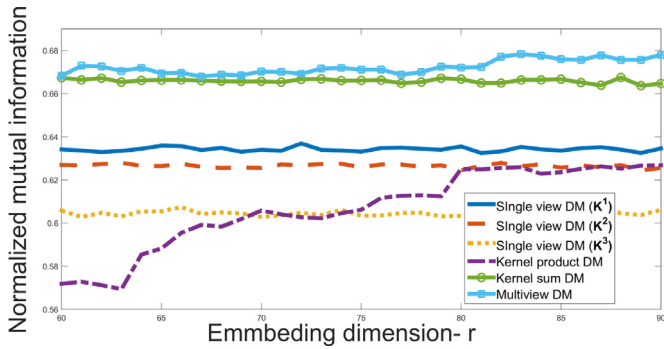


Fig. 26. Clustering accuracy measured with Normalized Mutual Information (NMI) on the Isolet data set by using 3 different kernel matrices. Clustering was performed in the r dimensional embedding space.

that its estimation is usually inaccurate without the depth phases. Other widely used seismic discrimination methods are Ms:mb (surface wave magnitude versus body wave magnitude) and spectral amplitude ratios of different seismic phases [67,68]. Current automatic seismic bulletins comprise a large number of false alarms, which have to be manually corrected by an analyst.

In this section we apply the proposed method to extract essential latent seismic parameters. A suspected event is identified based on a short and long time average ratio (STA/LTA) [69]. Then, a time-frequency representation is computed to which multi-view is applied to fuse the data from multiple seismic sensors. Using the multi-view low-dimensional embedding and simple classifiers, we demonstrate capabilities classification of event type (earthquakes vs. explosions) and quarry source for explosions.

7.2.1. Description of the data set

The dataset consists of recordings from two different broad band seismic stations MMLI (Malkishua) and HRFI (Harif). Both stations are operated by the Geophysical Institute of Israel (GII) and they are part of the Israel National Seismic Network [70]. MMLI and HRFI stations are located to the north and to the south of the analyzed region, respectively. Each station is equipped with a three component STS-2 seismometer, thus the total number of views is $L = 6$. All recording are sampled at

$F_s = 40$ Hz. The HRFI data includes 1654 explosions and 105 earthquakes, while MMLI includes a subset of 46 earthquakes, 62 explosions. The explosions occurred in the south of Israel between the years 2005 – 2015.

7.2.2. Feature extraction by normalized sonograms

A seismic event typically generates two underground traveling waves. A primary wave (P) and secondary wave (S). The two waves (P-S) arrive at the recording station with some time delay. A time frequency representation of the recording captures the spectral properties of the event while maintaining the P-S time gap. Here we use a time frequency representation termed sonogram [71] with some modifications. Each single-trace seismic waveform, denoted by $y[n] \in \mathbb{R}^N$ is a time series signal sampled at the rate of $F_s = 40$ Hz. The waveform $y[n]$ is decomposed into a set of overlapping windows of length $N_0 = 256$ using an overlap of $s = 0.9$. Thus, the shift between consecutive windows is $N_s = [0.1 \cdot 256] = 25$. A short-time Fourier transform (STFT) is applied to $y[n]$ in each time window. Then, power spectral densities are computed. The resulting spectrogram is denoted by $R(f, t)$, where f is a frequency bin and t is a number of time window (bin). Thus, it contains $T = 192$ time bins and $N_0/2 = 128$ frequency bins.

The sonogram is obtained by summing the spectrogram in equally tempered logarithmically scaled frequency bands, this is done for every time bin. Finally the sonogram is normalized such that the sum of energy in every frequency band is equal to 1. The result is a normalized sonogram and it is denoted by $S(k, t)$, where k is the frequency band number and t is the time window number. The resulting set of sonograms are denoted $X^1, X^2, X^3, X^4, X^5, X^6$. These are the input views for our framework.

7.2.3. Event classification

Each sensor records information from the seismic event as well as nuisance noise. We can assume that the noise at each station is independent. Thus, by fusing the measurements from different sensors we may be able to improve detection level. To evaluate how well the proposed approach fuses the information, we use a set which includes 46 earthquakes, 62 explosions recorded at MMLI (Malkishua) and HRFI (Harif). After we extract the sonogram, the proposed MVDm framework is applied, as well as the single-view DM, kernel product DM and kernel sum DM. Classification is performed by using K-NN ($K=1$), based on 3 or

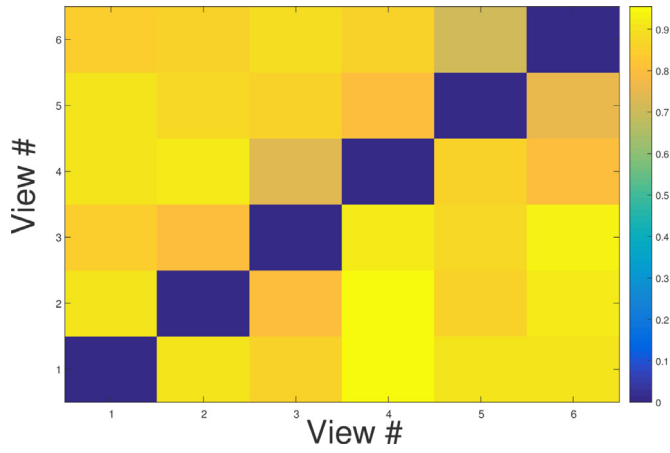


Fig. 27. Classification accuracy using K-nn (K=1) for all pairs of views X^l, X^m , $l \neq m$. The y-axis is the number of the first view used, while the x-axis is the number of the second view. Classification is performed in the multi-view low-dimensional embedding ($r = 4$). The diagonal terms are presented as zero since we did not simulate for $l = m$.

Table 3

Classification accuracy using 1-fold cross validation. r is the number of coordinates used in the embedding space.

Method	Accuracy [%] ($r = 3$)	Accuracy [%] ($r = 4$)
single-view DM (X^1)	89.9	93.0
single-view DM (X^2)	88.6	92.4
single-view DM (X^3)	89.3	91.1
single-view DM (X^4)	89.3	89.2
single-view DM (X^5)	89.3	90.5
single-view DM (X^6)	88.6	91.1
Kernel Sum DM	93.7	94.9
Kernel Product DM	94.3	93.0
Multi-view DM	97.5	98.1

4 coordinates from the reduced mapping. The results are presented in Table 3. This experiment demonstrates that applying multi-view DM to seismic recordings extracts a meaningful representation.

In the following test we check how each view affects the detection rate. We do this by applying the MV to subsets of the $L = 6$ views. We perform classification using representation computed based on all pairs of view $X^l, X^m, l, m = 1, \dots, 6, l \neq m$. The accuracy of classification based

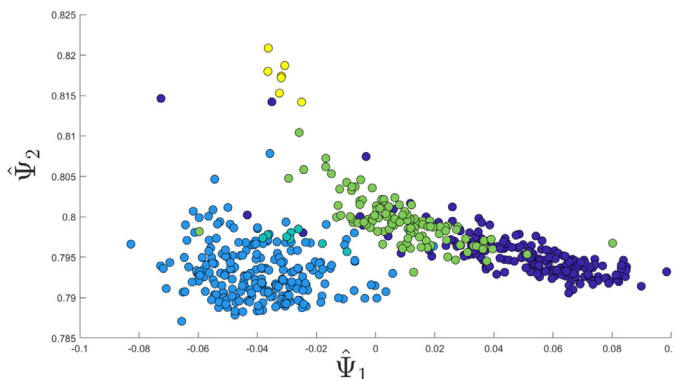


Table 4

Classification accuracy using 1-fold cross validation. r is the number of coordinates used in the embedding space.

Method	Accuracy [%] ($r = 3$)	Accuracy [%] ($r = 4$)
single-view DM (X^1)	80.3	80.6
single-view DM (X^2)	79.1	79.2
single-view DM (X^3)	76.4	77.8
Kernel Sum DM	82.2	82.8
Kernel Product DM	80.8	81.2
Multi-view DM	86.2	86.4

on the multi-view representation $\hat{\Psi}(X^l), l = 1, \dots, 6$ given that $X^m, m = 1, \dots, 6, m \neq l$ is presented in Fig. 27.

Identification and separation of quarries by attributing the explosions to the known sources is a challenging task [72,73]. Quarry blast have a similar spectral properties, they are usually classified by a triangulation process. Such a process requires to compare the arrival time between distinct seismic stations. Here we attempt to classify the source of the explosions, using 3-channels from the same station.

For this experiment 602 seismograms of explosions are used. The explosions occurred in 4 quarry clusters in Israel and 1 quarry in Jordan. All events were recorded in HRFI station, the distances from the event to the station vary between 50 and 130 Km. The association of each blast to quarry (labeling) is performed manually by an analyst from the GII. After extracting the sonograms, we apply multi-view DM and present the first 2 coordinates in Fig. 28. In Table 4 we summarize the classification results by applying K-NN ($k = 1$) in a leave one out procedure. Our method is compared to a single-view DM, kernel product and kernel sum.

8. Discussion

We presented a multi-view based framework for dimensionality reduction. The framework enables to extract simultaneous embeddings from coupled embeddings. Our approach is based on imposing an implied cross-domain transition in each single time-step. The transition probabilities depend on the connectivities in both views. We reviewed various theoretical aspects of the proposed method and demonstrated their applicability to both synthetic and real data. The experimental results demonstrate the strength of the proposed framework in cases where data is missing in each view or each of the manifolds is deformed by an unknown function. The framework is applicable to various real life machine learning tasks that consist of multiple views or multiple modalities.

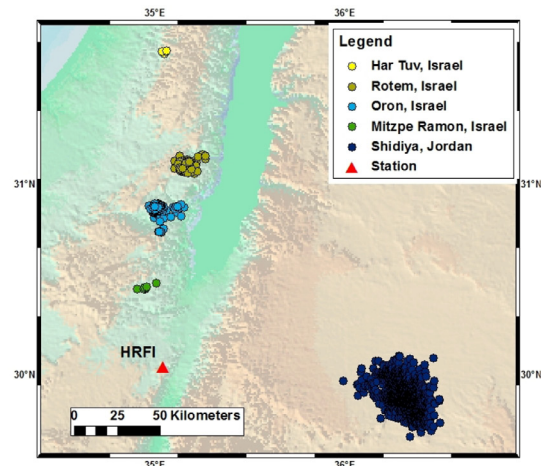


Fig. 28. Left- 2 dimensional multi-view DM of the 602 quarry blasts. Points are colored by quarry label. Right- a map with the approximated source location.

Appendix

We prove [Theorem 7](#) ([Section 5.6](#)), repeated here for convenience:

Theorem 7. The infinitesimal generator induced by the proposed kernel matrix $\hat{\mathbf{K}}$ (Eq. (13)) after row-normalization, denoted in here as $\hat{\mathbf{P}}$, converges when $M \rightarrow \infty$, $\epsilon \rightarrow 0$ (with $\epsilon = \sigma_x^2 = \sigma_y^2$) to a “cross domain Laplacian operator”. The functions $f(\mathbf{x})$ and $g(\mathbf{y})$ converge to eigenfunctions of $\hat{\mathbf{P}}$. These functions are the solutions of the following diffusion-like equations:

$$(\hat{\mathbf{P}}f)(\mathbf{x}_i) = g(\beta(\mathbf{x}_i)) + \epsilon \triangle \gamma(\beta(\mathbf{x}_i)) / \alpha(\beta(\mathbf{x}_i)) + \mathcal{O}(\epsilon^{3/2}), \quad (66)$$

$$(\hat{\mathbf{P}}g)(\mathbf{y}_i) = f(\beta^{-1}(\mathbf{y}_i)) + \epsilon \triangle \eta(\beta^{-1}(\mathbf{y}_i)) / \alpha(\beta(\mathbf{y}_i)) + \mathcal{O}(\epsilon^{3/2}), \quad (67)$$

where the functions γ , η are defined as $\gamma(\mathbf{z}) \triangleq g(\mathbf{z})\alpha(\mathbf{z})$, $\eta(\mathbf{z}) \triangleq f(\mathbf{z})\alpha(\mathbf{z})$.

Proof. By extending the single-view construction presented in [5], the eigenfunction of the limit-operator $\hat{\mathbf{P}}$ is defined using the functions $f(\mathbf{x})$ and $g(\mathbf{y})$ by concatenating the vectors such that

$$\mathbf{h} = [f(\mathbf{x}_1), f(\mathbf{x}_2), \dots, f(\mathbf{x}_M), g(\mathbf{y}_1), g(\mathbf{y}_2), \dots, g(\mathbf{y}_M)] \in \mathbb{R}^{2M}.$$

The limit of the top-half (Eq. (66)) of the characteristic equation is given by

$$\begin{aligned} \lim_{\substack{M \rightarrow \infty \\ \epsilon \rightarrow 0}} (\hat{\mathbf{P}}h_i) &= \lim_{\substack{M \rightarrow \infty \\ \epsilon \rightarrow 0}} h_i - \frac{\sum_{j=1}^{2M} \hat{\mathbf{K}}_{i,j} h_j}{\sum_{j=1}^{2M} \hat{\mathbf{K}}_{i,j}} \\ &= \lim_{\substack{M \rightarrow \infty \\ \epsilon \rightarrow 0}} f(\mathbf{x}_i) - \frac{\sum_{j=1}^M \sum_{\ell=1}^M K_{i,\ell}^x K_{\ell,j}^y g(\mathbf{y}_j)}{\sum_{j=1}^M \sum_{\ell=1}^M K_{i,\ell}^x K_{\ell,j}^y}, i = 1, \dots, M. \end{aligned} \quad (68)$$

We approximate the summations using a Riemann integral. Beginning with the denominator, we have

$$\begin{aligned} \frac{1}{M^2 \epsilon^d} \sum_{j=1}^M \sum_{\ell=1}^M K_{i,\ell}^x K_{\ell,j}^y &\xrightarrow{\substack{M \rightarrow \infty \\ \epsilon \rightarrow 0}} D(\mathbf{x}) \triangleq \frac{1}{\epsilon^d} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) K \\ &\left(\frac{\mathbf{y}-\beta(\mathbf{s})}{\sqrt{\epsilon}}\right) \alpha(\mathbf{y}) d\mathbf{s} d\mathbf{y}. \end{aligned}$$

Using a change of variables $\mathbf{z} = \frac{\mathbf{y}-\beta(\mathbf{s})}{\sqrt{\epsilon}}$, $\mathbf{y} = \beta(\mathbf{s}) + \sqrt{\epsilon}\mathbf{z}$, $d\mathbf{z} = d\mathbf{y}\epsilon^{d/2}$ we get

$$D(\mathbf{x}) = \frac{1}{\epsilon^{d/2}} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) K(\mathbf{z}) \alpha(\beta(\mathbf{s}) + \sqrt{\epsilon}\mathbf{z}) d\mathbf{s} d\mathbf{z}.$$

Using a first order Taylor expansion of $\alpha(\cdot)$ we get

$$\begin{aligned} D(\mathbf{x}) &\approx \frac{1}{\epsilon^{d/2}} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) K(\mathbf{z}) \left[\alpha(\beta(\mathbf{s})) + \frac{\sqrt{\epsilon}}{2} \mathbf{z}^T \nabla \alpha(\beta(\mathbf{s})) \right. \\ &\left. + \mathcal{O}(\epsilon) \right] d\mathbf{s} d\mathbf{z}, \end{aligned}$$

using the symmetry of the kernel $K(\mathbf{z})$ we have

$$\int_{\mathbb{R}^d} K(\mathbf{z}) \mathbf{z}^T d\mathbf{z} = \mathbf{0}^T,$$

therefore, applying another change of variables $\mathbf{t} = \frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}$, $\mathbf{s} = \sqrt{\epsilon}\mathbf{t} + \mathbf{x}$, $d\mathbf{t} = d\mathbf{s}\epsilon^{d/2}$ we get

$$D(\mathbf{x}) \approx \int_{\mathbb{R}^d} K(\mathbf{t}) [\alpha(\beta(\mathbf{x} + \sqrt{\epsilon}\mathbf{t})) + \mathcal{O}(\epsilon)] d\mathbf{t} \approx \alpha(\beta(\mathbf{x})) + \mathcal{O}(\epsilon)$$

(independent of $\mathbf{x}$), where the last transition is again based on a Taylor expansion (of $\beta(\cdot)$) and on zeroing out the odd (1st) moment of $K(\mathbf{t})$.

Turning to the numerator (of Eq. (68)),

$$\begin{aligned} \frac{1}{M^2 \epsilon^d} \sum_{j=1}^M \sum_{\ell=1}^M K_{i,\ell}^x K_{\ell,j}^y g(\mathbf{y}_j) &\xrightarrow[\epsilon \rightarrow 0]{M \rightarrow \infty} N(\mathbf{x}) \triangleq \frac{1}{\epsilon^d} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) \\ &K\left(\frac{\mathbf{y}-\beta(\mathbf{s})}{\sqrt{\epsilon}}\right) g(\mathbf{y}) \alpha(\mathbf{y}) d\mathbf{s} d\mathbf{y}. \end{aligned}$$

By applying a change of variables $\mathbf{z} = \frac{\mathbf{y}-\beta(\mathbf{s})}{\sqrt{\epsilon}}$, $\mathbf{y} = \beta(\mathbf{s}) + \sqrt{\epsilon}\mathbf{z}$, $d\mathbf{z} = d\mathbf{y}\epsilon^{d/2}$ we get

$$N(\mathbf{x}) = \frac{1}{\epsilon^{d/2}} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) K(\mathbf{z}) \gamma(\beta(\mathbf{s}) + \sqrt{\epsilon}\mathbf{z}) d\mathbf{s} d\mathbf{z}.$$

Using Taylor's expansion of $\gamma(\cdot)$ we get

$$\begin{aligned} N(\mathbf{x}) &\approx \frac{1}{\epsilon^{d/2}} \int_{\mathbb{R}^d} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) K(\mathbf{z}) \left[\gamma(\beta(\mathbf{s})) + \frac{\sqrt{\epsilon}}{2} \mathbf{z}^T \nabla \gamma(\beta(\mathbf{s})) \right. \\ &\left. + \frac{\epsilon}{2} \mathbf{z}^T \mathbf{H} \mathbf{z} + \mathcal{O}(\epsilon^{3/2}) \right] d\mathbf{s} d\mathbf{z} \end{aligned}$$

where $H_{i,j} \triangleq \frac{\partial^2 \gamma(\beta(\mathbf{s}))}{\partial s_i \partial s_j}$ is the Hessian. The first term yields the integral over $K(\mathbf{z})$, while the second term vanishes due to integration over an odd (1st) moment of the symmetric kernel $K(\mathbf{z})$. The last term yields

$$\begin{aligned} &\int_{\mathbb{R}^d} K(\mathbf{z}) \mathbf{z}^T \frac{\partial^2 \gamma(\beta(\mathbf{s}))}{\partial s_i \partial s_j} \mathbf{z} d\mathbf{z} \\ &= \sum_{i,j} \frac{\partial^2 \gamma(\beta(\mathbf{s}))}{\partial s_i \partial s_j} \int_{\mathbb{R}^d} z_i z_j K(\mathbf{z}) d\mathbf{z} \\ &= \sum_i \frac{\partial^2 \gamma(\beta(\mathbf{s}))}{\partial s_i^2} \int_{\mathbb{R}^d} z_i^2 K(\mathbf{z}) d\mathbf{z} = \triangle \gamma(\beta(\mathbf{s})), \end{aligned}$$

where $\triangle$ denotes the Laplacian operator, and where we assumed that $\int_{\mathbb{R}^d} z_i z_j K(\mathbf{z}) d\mathbf{z}$ vanishes for $i \neq j$ and equals 1 for $i = j$. Note that this is naturally satisfied, e.g., by the Gaussian kernel function $K(\mathbf{z}) = c \cdot \exp(-0.5\|\mathbf{z}\|^2)$ (with $c = (2\pi)^{-d/2}$, as per the scaling requirement). Substituting into $N(\mathbf{x})$ we get

$$N(\mathbf{x}) \approx \frac{1}{\epsilon^{d/2}} \int_{\mathbb{R}^d} K\left(\frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}\right) \left[\gamma(\beta(\mathbf{s})) + \frac{\epsilon}{2} \triangle \gamma(\beta(\mathbf{s})) + \mathcal{O}(\epsilon^{3/2}) \right] d\mathbf{s}.$$

Using a change of variables $\mathbf{t} = \frac{\mathbf{s}-\mathbf{x}}{\sqrt{\epsilon}}$, $\mathbf{s} = \sqrt{\epsilon}\mathbf{t} + \mathbf{x}$, $d\mathbf{t} = d\mathbf{s}\epsilon^{d/2}$ and $\gamma(\mathbf{y}) = g(\mathbf{y})\alpha(\mathbf{y})$ we get

$$N(\mathbf{x}) \approx \int_{\mathbb{R}^d} K(\mathbf{t}) \left[\gamma(\beta(\mathbf{x} + \sqrt{\epsilon}\mathbf{t})) + \frac{\epsilon}{2} \triangle \gamma(\beta(\mathbf{x} + \sqrt{\epsilon}\mathbf{t})) + \mathcal{O}(\epsilon^{3/2}) \right] d\mathbf{t}.$$

Using Taylor's expansion (of $\beta(\cdot)$) once again we get

$$N(\mathbf{x}) \approx \int_{\mathbb{R}^d} K(\mathbf{t}) \left[\gamma(\beta(\mathbf{x})) + \frac{\epsilon}{2} \triangle \gamma(\beta(\mathbf{x})) + \frac{\epsilon}{2} \mathbf{t}^T \mathbf{H} \mathbf{t} + \mathcal{O}(\epsilon^{3/2}) \right] d\mathbf{t},$$

here we neglected terms involving ϵ to a power higher than 3/2 and terms with odd order of $\mathbf{t}$ due to the symmetry of the kernel K . This leads to

$$N(\mathbf{x}) \approx \gamma(\beta(\mathbf{x})) + \epsilon \triangle \gamma(\beta(\mathbf{x})) + \mathcal{O}(\epsilon^{3/2}),$$

dividing by the denominator we get

$$(\hat{\mathbf{P}}f)(\mathbf{x}_i) \approx g(\beta(\mathbf{x}_i)) + \epsilon \triangle \gamma(\beta(\mathbf{x}_i)) / \alpha(\beta(\mathbf{x}_i)) + \mathcal{O}(\epsilon^{3/2})$$

In the same way, we compute the convergence on $g(\mathbf{y}_i)$

$$(\hat{\mathbf{P}}g)(\mathbf{y}_i) = f(\beta^{-1}(\mathbf{y}_i)) + \epsilon \triangle \eta(\beta^{-1}(\mathbf{y}_i)) / \alpha(\beta(\mathbf{y}_i)) + \mathcal{O}(\epsilon^{3/2}).$$

□

The following issues were ignored in the proof:

- Errors due to approximating the sum by an integral; An upper bound on the associated errors in the single-view DM is derived in [34].

- Deformation due to the fact that the data is sampled from a non uniform density. This changes the result by some constant.
- The data lies on some manifold. This could be dealt by changing the coordinate system and integrating on the manifold.
- When assuming that the data lies on some manifold, the Euclidean distance should be replaced by the geodesic distance along the manifold. As in the analysis of [5], this introduces a factor to the integral.

References

- [1] H.-P. Deutsch, Principle component analysis, in: *Derivatives and Internal Models*, Springer, 2002, pp. 539–547.
- [2] J.B. Kruskal, Nonmetric multidimensional scaling: a numerical method, *Psychometrika* 29 (2) (1964) 115–129.
- [3] S.T. Roweis, L.K. Saul, Nonlinear dimensionality reduction by local linear embedding, *Science* 290.5500 (2000) 2323–2326.
- [4] M. Belkin, P. Niyogi, Laplacian eigenmaps for dimensionality reduction and data representation, *Neural Comput.* 15 (6) (2003) 1373–1396.
- [5] R.R. Coifman, S. Lafon, Diffusion maps, *Appl. Comput. Harmon. Anal.* 21 (1) (2006) 5–30.
- [6] W. Liu, X. Ma, Y. Zhou, D. Tao, J. Cheng, p-laplacian regularization for scene recognition, *IEEE Trans. Cybern.* 49 (8) (2018) 5120–5129.
- [7] L.v.d. Maaten, G. Hinton, Visualizing data using t-sne, *J. Mach. Learn. Res.* 9 (Nov) (2008) 2579–2605.
- [8] L. McInnes, J. Healy, N. Saul, L. Großberger, Umap: uniform manifold approximation and projection, *J. Open Source Softw.* 3 (29) (2018) 861.
- [9] X. He, S. Yan, Y. Hu, P. Niyogi, H.-J. Zhang, Face recognition using laplacianfaces, *IEEE Trans. Pattern Anal. Mach. Intell.* 27 (3) (2005) 328–340.
- [10] A. Singer, R.R. Coifman, Non linear independent component analysis with diffusion maps, *Appl. Comput. Harmon. Anal.* 25(2) (2008) 226–239.
- [11] O. Lindenbaum, A. Yeredor, I. Cohen, Musical key extraction using diffusion maps, *Signal Process.* 117 (2015) 198–207.
- [12] W.T. Freeman, J.B. Tenenbaum, Learning bilinear models for two-factor problems in vision, in: *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition, IEEE*, 1997, pp. 554–560.
- [13] I.S. Helland, Partial least squares regression and statistical models, *Scand. J. Stat.* 17 (1990) 97–114.
- [14] K. Chaudhuri, S.M. Kakade, K. Livescu, K. Sridharan, Multi-view clustering via canonical correlation analysis, in: *Proceedings of the Twenty-Sixth Annual International Conference on Machine Learning, ACM*, 2009, pp. 129–136.
- [15] P.L. Lai, C. Fyfe, Kernel and nonlinear canonical correlation analysis, *Int. J. Neural Syst.* 10 (05) (2000) 365–377.
- [16] F.R. Bach, M.I. Jordan, Kernel independent component analysis, *J. Mach. Learn. Res.* 3 (Jul) (2002) 1–48.
- [17] A. Kumar, H. Daumé, A co-training approach for multi-view spectral clustering, in: *Proceedings of the Twenty-Eighth International Conference on Machine Learning (ICML-11)*, 2011, pp. 393–400.
- [18] B. Boots, G. Gordon, Two-manifold problems with applications to nonlinear system identification, *arXiv preprint arXiv:1206.4648*, 2012.
- [19] V.R. de Sa, Spectral clustering with two views, in: *Proceedings of the ICML Workshop on Learning with Multiple Views*, 2005, pp. 20–27.
- [20] S. Sun, X. Xie, M. Yang, Multiview uncorrelated discriminant analysis, *IEEE Trans. Cybern.* 46 (12) (2016) 3272–3284.
- [21] W. Liu, X. Yang, D. Tao, J. Cheng, Y. Tang, Multiview dimension reduction via hesian multitset canonical correlations, *Inf. Fusion* 41 (2018) 119–128.
- [22] B. Wang, J. Jiang, W. Wang, Z.-H. Zhou, Z. Tu, Unsupervised metric fusion by cross diffusion, in: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, IEEE, 2012, pp. 2997–3004.
- [23] D. Zhou, C. Burges, Spectral clustering and transductive learning with multiple views, in: *Proceedings of the Twenty-Fourth International Conference on Machine Learning*, 2007, pp. 1159–1166.
- [24] R.R. Lederman, R. Talmon, Learning the geometry of common latent variables using alternating-diffusion, *Appl. Comput. Harmon. Anal.* 44 (3) (2018) 509–536.
- [25] W. Wang, M.A. Carreira-Perpinán, The role of dimensionality reduction in classification., in: *Proceedings of the AAAI*, 2014, pp. 2128–2134.
- [26] H. Liu, L. Liu, T.D. Le, I. Lee, S. Sun, J. Li, et al., Non-parametric sparse matrix decomposition for cross-view dimensionality reduction, *IEEE Trans. Multimed.* 15 (2017) 138–145.
- [27] G. Andrew, R. Arora, J. Bilmes, K. Livescu, Deep canonical correlation analysis, in: *Proceedings of the International Conference on Machine Learning*, 2013, pp. 1247–1255.
- [28] J. Zhao, X. Xie, X. Xu, S. Sun, Multi-view learning overview: recent progress and new challenges, *Inf. Fusion* 38 (2017) 43–54.
- [29] O. Lindenbaum, A. Yeredor, M. Salhov, Learning coupled embedding using multiview diffusion maps, in: *Proceedings of the International Conference on Latent Variable Analysis and Signal Separation*, Springer, 2015, pp. 127–134.
- [30] D.S. Kershaw, The incomplete Cholesky conjugate gradient method for the iterative solution of systems of linear equations, *J. Comput. Phys.* 26 (1) (1978) 43–65.
- [31] S. Lafon, Diffusion maps and geometric harmonics, Yale, 2004 Ph.D dissertation.
- [32] Y. Keller, R.R. Coifman, S. Lafon, S.W. Zucker, Audio-visual group recognition using diffusion maps, *IEEE Trans. Signal Process.* 58 (1) (2010) 403–413.
- [33] D. Dov, R. Talmon, I. Cohen, Kernel-based sensor fusion with application to audio-visual voice activity detection, *IEEE Trans. Signal Process.* 64 (24) (2016) 6406–6416.
- [34] A. Singer, R. Erban, I.G. Kevrekidis, R.R. Coifman, Detecting intrinsic slow variables in stochastic dynamical systems by anisotropic diffusion maps, *Proc. Natl. Acad. Sci.* 106 (38) (2009) 16090–16095.
- [35] Y. Aizenbud, A. Averbuch, Matrix decompositions using sub-gaussian random matrices, *Inf. Inference: J. IMA* 8 (3) (2018) 445–469.
- [36] M. Holmes, A. Gray, C. Isbell, Fast SVD for large-scale matrices, in: *Proceedings of the Workshop on Efficient Machine Learning at NIPS*, 58, 2007, pp. 249–252.
- [37] M.P. Holmes, J. Isbell, C. Lee, A.G. Gray, QUIC-SVD: fast SVD using cosine trees, in: *Proceedings of the Advances in Neural Information Processing Systems*, 2009, pp. 673–680.
- [38] R.R. Coifman, M.J. Hirn, Diffusion maps for changing data, *Appl. Comput. Harmon. Anal.* 36 (1) (2014) 79–107.
- [39] T. Ando, Majorization relations for Hadamard products., *Linear Algebra Appl.* 223 (1995) 57–64.
- [40] G. Visick, A weak majorization involving the matrices A B and AB, *Linear Algebra Appl.* 224/224 (1995) 731–744.
- [41] A. Bermanis, A. Averbuch, R.R. Coifman, Multiscale data sampling and function extension, *Appl. Comput. Harmon. Anal.* 34 (1) (2013) 15–29.
- [42] K. Fukunaga, D.R. Olsen, An algorithm for finding intrinsic dimensionality of data, *IEEE Trans. Comput.* 100 (2) (1971) 176–183.
- [43] P.J. Verveer, R.P.W. Duin, An evaluation of intrinsic dimensionality estimators, *IEEE Trans. Pattern Anal. Mach. Intell.* 17 (1) (1995) 81–86.
- [44] C. Cerutti, S. Bassis, A. Rozza, G. Lombardi, E. Casiraghi, P. Campadelli, Danco: an intrinsic dimensionality estimator exploiting angle and norm concentration, *Pattern Recognit.* 47 (8) (2014) 2569–2581.
- [45] A. Bermanis, A. Averbuch, R.R. Coifman, Multiscale data sampling and function extension, *Appl. Comput. Harmon. Anal.* 34 (1) (2013) 15–29.
- [46] M. Belkin, P. Niyogi, Convergence of Laplacian eigenmaps, in: *Proceedings of the NIPS*, 2006, pp. 129–136.
- [47] M. Hein, J.-Y. Audibert, U. Von Luxburg, From graphs to manifolds—weak and strong pointwise consistency of graph laplacians, in: *Proceedings of the International Conference on Computational Learning Theory*, Springer, 2005, pp. 470–485.
- [48] A. Singer, From graph to manifold Laplacian: the convergence rate, *Appl. Comput. Harmon. Anal.* 21 (1) (2006) 128–134.
- [49] H. Hotelling, Relations between two sets of variates, *Biometrika* (1936) 321–377.
- [50] K.Q. Weinberger, F. Sha, L.K. Saul, Learning a kernel matrix for nonlinear dimensionality reduction, in: *Proceedings of the Twenty-First International Conference on Machine Learning, ACM*, 2004, p. 106.
- [51] Y. Shiohara, Y. Date, J. Kikuchi, Application of kernel principal component analysis and computational machine learning to exploration of metabolites strongly associated with diet, *Sci. Rep.* 8 (1) (2018) 3426.
- [52] Q. Zhang, S. Filippi, A. Gretton, D. Sejdinovic, Large-scale kernel methods for independence testing, *Stat. Comput.* 28 (1) (2018) 113–130.
- [53] B. Schölkopf, The kernel trick for distances, in: *Proceedings of the Advances in Neural Information Processing Systems*, 2001, pp. 301–307.
- [54] S. Lafon, Y. Keller, R. Coifman, Data fusion and multicue data matching by diffusion maps, *IEEE Trans. Pattern Anal. Mach. Intell.* 28 no. 11 (2006) 1784/1797.
- [55] G. Piella, Diffusion maps for multimodal registration, *Sensors* 14 (6) (2014) 10562–10577.
- [56] O. Lindenbaum, A. Yeredor, A. Averbuch, Clustering based on multiview diffusion maps, in: *Proceedings of the IEEE Sixteenth International Conference on Data Mining Workshops (ICDMW)*, IEEE, 2016, pp. 740–747.
- [57] A.Y. Ng, M.I. Jordan, Y. Weiss, On spectral clustering analysis and an algorithm, in: *Proceedings of Advances in Neural Information Processing Systems*, Cambridge, MA, 14, MIT Press, 2001, pp. 849–856.
- [58] M. Polito, P. Perona, Grouping and dimensionality reduction by locally linear embedding, in: *Proceedings of the NIPS*, 2001, pp. 1255–1262.
- [59] P.A. Estévez, M. Tesmer, C.A. Perez, J.M. Zurada, Normalized mutual information feature selection, *IEEE Trans. Neural Netw.* 20 (2) (2009) 189–201.
- [60] G. Loosli, S. Canu, L. Bottou, Training invariant support vector machines using selective sampling, *Large Scale Kernel Mach.* 2 (2007) 301–320.
- [61] L. Fei-Fei, R. Fergus, P. Perona, Learning generative visual models from few training examples: an incremental Bayesian approach tested on 101 object categories, *Comput. Vis. Image Underst.* 106 (1) (2007) 59–70.
- [62] Y.-G. Jiang, C.-W. Ngo, J. Yang, Towards optimal bag-of-features for object categorization and semantic video retrieval, in: *Proceedings of the Sixth ACM International Conference on Image and Video Retrieval*, ACM, 2007, pp. 494–501.
- [63] Y. Bai, L. Guo, L. Jin, Q. Huang, A novel feature extraction method using pyramid histogram of orientation gradients for smile recognition, in: *Proceedings of the Sixteenth IEEE International Conference on Image Processing (ICIP)*, IEEE, 2009, pp. 3305–3308.
- [64] P. Gehler, S. Nowozin, On feature combination for multiclass object detection, in: *Proceedings of the International Conference on Computer Vision*, 2, 2009, p. 6.
- [65] M.A. Mouchet, S. Villéger, N.W. Mason, D. Mouillot, Functional diversity measures: an overview of their redundancy and their ability to discriminate community assembly rules, *Funct. Ecol.* 24 (4) (2010) 867–876.
- [66] P.J. Rousseeuw, Silhouettes: a graphical aid to the interpretation and validation of cluster analysis, *J. Comput. Appl. Math.* 20 (1987) 53–65.
- [67] R. Blandford, Seismic event discrimination, *Bull. Seismol. Soc. Am.* 72 (1982) 569–587.
- [68] A.J. Rodgers, T. Lay, W.R. Walter, K.M. Mayeda, A comparison of regional-phase amplitude ratio measurement techniques, *Bull. Seismol. Soc. Am.* 87 (6) (1997) 1613–1621.
- [69] L. Küperkoch, T. Meier, T. Diehl, Automated event and phase identification, *New Man. Seismol. Obs. Pract.* 2 (2012) 1–52.

- [70] R. Hofstetter, Y. Gitterman, V. Pinsky, N. Kraeva, L. Feldman, Seismological observations of the northern dead sea basin earthquake on 11 february 2004 and its associated activity., *Israel J. Earth Sci.* 57 (2008) 101–124.
- [71] M. Joswig, Automated classification of local earthquake data in the bug small array, *Geophys. J. Int.* 120 (2) (1995) 262–286.
- [72] D.B. Harris, A waveform correlation method for identifying quarry explosions, *Bull. Seismol. Soc. Am.* 81 (6) (1991) 2395–2418.
- [73] D.B. Harris, T. Kvaerna, Superresolution with seismic arrays using empirical matched field processing, *Geophys. J. Int.* 182 (3) (2010) 1455–1477.