

Multiscale Laplacian Learning

Ekaterina Merkurjev Duc Duy Nguyen Guo-Wei Wei

Received: date / Accepted: date

Abstract Machine learning methods have greatly changed science, engineering, finance, business, and other fields. Despite the tremendous accomplishments of machine learning and deep learning methods, many challenges still remain. In particular, the performance of machine learning methods is often severely affected in case of diverse data, usually associated with smaller data sets or data related to areas of study where the size of the data sets is constrained by the complexity and/or high cost of experiments. Moreover, data with limited labeled samples is a challenge to most learning approaches. In this paper, the aforementioned challenges are addressed by integrating graph-based frameworks, multiscale structure, modified and adapted optimization procedures and semi-supervised techniques. This results in two innovative multiscale Laplacian learning (MLL) approaches for machine learning tasks, such as data classification, and for tackling diverse data, data with limited samples and smaller data sets. The first approach, called multikernel manifold learning (MML), integrates manifold learn-

ing with multikernel information and solves a regularization problem consisting of a loss function and a warped kernel regularizer using multiscale graph Laplacians. The second approach, called the multiscale MBO (MMBO) method, introduces multiscale Laplacians to a modification of the famous classical Merriman-Bence-Osher (MBO) scheme, and makes use of fast solvers for finding the approximations to the extremal eigenvectors of the graph Laplacian. We demonstrate the performance of our methods experimentally on a variety of data sets, such as biological, text and image data, and compare them favorably to existing approaches.

Keywords graph-based methods manifold learning
multiscale framework graph Laplacian

1 Introduction

Artificial intelligence, including machine learning, has irreversibly changed many fields including science and engineering [48, 53]. In fact, the combination of artificial intelligence (AI) and big data has been referred to as the “fourth industrial revolution” [89]. Nevertheless, machine learning tasks face several challenges.

First, while the big data challenge is well known, little attention is paid to the diverse data challenge. The success behind machine learning is that the behavior in unknown domains can be accurately estimated by quantitatively learning the pattern from sufficient training samples. However, while data sets in computer vision and image analysis often contain millions or billions of points, it is typically difficult to obtain large data sets in science [46]. We often deal with diverse data originating from a relatively small data set lying in a huge space. For example, due to the complexity, ethnicity, and high cost of scientific experiments [43, 86, 90, 91], it is extremely difficult to collect a relatively small set of drug candidates of the order of 10^6 for a

This work was supported in part by NIH grant GM126189, NSF Grants DMS-1721024, DMS-1761320, and IIS1900473, Michigan Economic Development Corporation, Bristol-Myers Squibb, Pfizer, and University of Kentucky Startup Fund. The authors thank the IBM TJ Watson Research Center, the COVID-19 High Performance Computing Consortium, and NVIDIA for computational assistance.

Ekaterina Merkurjev
Department of Mathematics,
Michigan State University, MI 48824, USA
E-mail: merkurje@msu.edu

Duc Duy Nguyen
Department of Mathematics, University of Kentucky, KY 40506, USA
E-mail: ducnguyen@uky.edu

Guo-Wei Wei
Department of Mathematics, Department of Biochemistry and Molecular Biology, Department of Electrical and Computer Engineering
Michigan State University, MI 48824, USA
E-mail: wei@math.msu.edu

therapeutic target, while the size of the underlying chemical space of potentially pharmacologically active molecules is about 10^{60} [11]. Therefore, researchers try to cover as many components as possible with a small number of sampling points. The diversity is created by deliberately sampling a wide distribution in the huge space to understand the landscape of potential drugs. This practice is very common in scientific explorations. Similar diverse data sets exist in materials design [30,110]. Overall, diverse data originated from a relatively small data set lying in a huge chemical space gives rise to a serve challenge for machine learning. Mathematically, diverse data involves disconnected submanifolds and/or nested submanifolds corresponding to multiphysics and multiscale natures of the diversity, respectively [23,70]. The multiphysics and multiscale representations of data have been addressed by the authors' earlier work on element-specific persistent homology [14, 16–18]. However, multiscale graph learning models have hardly been developed. The proposed algorithms of this paper fill the gap, addressing the multiphysics nature of data diversity through a multiphysics data representation, such as the element-specific feature extraction developed in [14, 16–18,69].

Second, the success of many existing approaches for machine learning tasks, such as data classification, is dependent on a sufficient amount of labeled samples. However, obtaining enough labeled data is difficult as it is time-consuming and expensive, especially in domains where only experts can determine experimental labels; thus, labeled data is scarce. As a result, the majority of the data embedded into a graph is unlabeled data, which is often much easier to obtain than labeled data but more challenging to predict. Overall, one of the key limitations of most existing approaches is their reliance on large labeled sets; in particular, deep learning approaches often require massive labeled sets to learn the patterns behind the data. These challenges call for innovative strategies to revolutionize the current state-of-the-art.

Recently, algorithms involving the graph-based framework, such as those described in Section 2.1, have recently become some of the most competitive approaches for applications ranging from image processing to the social sciences. Such methods have been successful in part due to the many advantages offered by using a graph-based approach. For example, a graph-based framework provides valuable information about the extent of similarity between elements of both labeled and unlabeled data via a weighted similarity graph and also yields information about the overall structure of the data. Moreover, in addition to handling nonlinear structure, a graph setting embeds the dimension of the features in a graph during weight computations, thus reducing the high-dimensionality of the problem. The graph framework is also able to incorporate diverse types of data, such as 3D point clouds, hyperspectral data, text, etc.

Inspired by the recent successes, we address the aforementioned challenges by integrating similarity graph-based frameworks, multiscale structure, modified and adapted optimization techniques and semi-supervised procedures, with both labeled and unlabeled data embedded into a graph. Overall, this paper formulates two multiscale Laplacian learning (MLL) approaches for machine learning tasks, such as data classification, and for dealing with diverse data, data with limited samples and smaller data sets. The first approach, the multikernel manifold learning (MML) method, introduces multiscale kernels to manifold regularization. This approach integrates new multiscale graph Laplacians into loss-function based minimization problems involving warped kernel regularizers. The second approach, the multiscale Merriman-Bence-Osher (MMBO) method, adapts and generalizes the classical Merriman-Bence-Osher (MBO) scheme [67] to a multiscale graph Laplacian setting for learning tasks. The MMBO approach also makes use of fast solvers, such as [10, 31, 32] and [6], for finding approximations of the extremal eigenvectors of the graph Laplacian. We validate the proposed MLL approaches using a variety of data sets.

There are several strengths of the proposed methods:

- The methods address the multiscale nature of data through a multiphysics data representation, allowing them to perform well in the case of diverse data, which often occurs in, e.g., scientific applications.
- The methods require less labeled training data to accurately classify a data set compared to most existing machine learning techniques, especially supervised approaches, and often in considerably smaller quantities. This is in part due to the usage of a similarity graph-based framework and the fact that the majority of the data embedded into the graph is unlabeled data. In fact, in most cases, a good accuracy can be obtained with at most 1%-5% of the data elements serving as labeled data. This is an important advantage due to the scarcity of labeled data for most applications.
- Although equally applicable and successful in the case of larger data, the new methods also perform well in the case of smaller data sets, which often result in unsatisfactory performances for existing machine learning techniques, due to an often insufficient number of labeled samples and a decreased ability for machine learning-based models to learn from the observed data.

The proposed MMBO method offers specific advantages:

- Although it can perform just as successfully on smaller data, the MMBO algorithm is equipped with a structure which allows it to be easily adapted and designed specifically for the use of large data. In particular, in the case of large data, one can use a slight modification of the fast Nystrom extension procedure [10, 31, 32] to compute an approximation to the extremal eigenvectors of

the multiscale graph Laplacian using a dense graph without the need to compute all the graph weights; in fact, only a small portion of the weights need to be calculated. Overall, the method uses a low-dimensional subspace spanned by only a small number of eigenfunctions.

- Once the N_e eigenvectors of the graph Laplacian are computed, the complexity of this algorithm is linear. Moreover, the Nystrom extension procedure allows the N_e eigenvectors of the graph Laplacian to be computed using only $O(NN_e)$ operations, where $N_e \ll N$ and N is the number of data elements.

The paper is organized as follows. In Section 2, we present background, previous work and an overview of graph learning methods. In Section 3, we derive the proposed MML and MMBO methods and provide details on the computation of eigenvectors of the graph Laplacian for the latter method. The results from experiments are described in Section 4, and we present a conclusion in Section 5.

2 Background

2.1 Previous work

In this section, we review recent graph-based methods for data classification and semi-supervised learning, including approaches related to convolutional neural networks, support vector machines, neural networks, label propagation, embedding methods, multi-view and multi-modal methods.

Convolutional neural networks have recently been extended to a graph-based framework for the purpose of semi-supervised learning. In particular, [52] presents a scalable approach using graph convolutional networks by integrating a convolutional architecture motivated by a localized first-order approximation of spectral graph convolutions. Deeper insights into the graph convolutional neural network model are discussed in [55]. Moreover, a dual graph-based convolutional network approach is described in [117], while a Bayesian graph convolutional network procedure is derived in [111]. In [4], a multi-scale graph convolution model is presented. In [13], generalizations of convolutional neural networks to signals defined on more general domains using two constructions are described; one of the generalizations is based on the spectrum of the graph Laplacian.

Neural networks have also been extended to a graph-based framework for the task of semi-supervised learning. For example, attention-based graph neural networks [97], graph partition neural networks [57], and graph Markov neural networks [85] have been proposed.

Moreover, support vector machines are also applied to semi-supervised learning using a graph-based framework. In [21], graph-based support vector machines are derived to

emphasize low density regions. Also, Laplacian support vector machines (LapSVM) [9, 60] and Laplacian twin support vector machines (Lap-TSVM) [84] have been formulated.

Label and measure propagation methods are discussed in, e.g., [44], where the authors derive a transductive label propagation method that is based on the manifold assumption. Label propagation techniques and the use of unlabeled data in classification are investigated in [115]. Dynamic label propagation is studied in [99], while semi-supervised learning with measure propagation is described in [95].

Embedding methods are also used for semi-supervised learning. Nonlinear embedding algorithms for use with shallow semi-supervised learning techniques, such as kernel methods, are applied to deep multi-layer architectures in [104]. Other graph embedding methods are presented in [108].

Multi-view and multi-modal methods include [77], which proposes a reformulation of a standard spectral learning model that can be used for multiview clustering and semi-supervised tasks. The work [76] proposes novel multi-view learning, while [38] describes multi-modal curriculum learning.

Other techniques for graph-based semi-supervised learning include fast anchor graph regularization [101], a Bayesian framework for learning hyperparameters [49], and random subspace dimensionality reduction. In [37], a classification method is proposed to learn from dissimilarity and similarity information on labeled and unlabeled data using a novel graph-based encoding of dissimilarity. Random graph walks are used in [58], and sampling theory for graph signals is utilized in [34]. In [100], a bivariate formulation for graph-based semi-supervised learning is shown to be equivalent to a linearly constrained max-cut problem. Lastly, reproducing kernel Hilbert spaces are used in [93].

Various approaches involving graph-based regularization terms include regularization frameworks [113, 114], regularization developments [20], anchor graph regularization [101], manifold regularization [9], measure propagation [95], approximate energy minimization [12], nonlocal discrete regularization [28], power watershed [27], spectral matting [54], Laplacian regularized least squares [94], locality and similarity preserving embedding [29], and clustering [79]. Examples for graph Laplacian regularization include label propagation [115] and deep semi-supervised embedding [104].

Merkurjev and co-authors have studied graph-based spectral approaches [35, 36, 61–66] using Ginzburg-Landau techniques and modifications of the MBO scheme [67], which is an efficient method for evolving an interface by mean curvature in a continuous setting and which can be linked to optimization problems involving the Ginzburg-Landau functional. Specifically, the MBO scheme can be derived from a Ginzburg-Landau functional minimization procedure, and can be modified and transferred to a graph setting using more general operators on graphs, as shown in Merkurjev's work on data classification [35, 61, 64–66].

Overall, Merkurjev and co-authors have shown that multiclass data classification can be achieved using techniques from topological spaces and the Gibbs simplex [35, 64]. In particular, MBO-like methods were developed for image processing applications [65], hyperspectral imaging [36, 66], Cheeger and ratio cut applications [63], heat kernel pagerank applications [62], and unsupervised learning [61]. The subject of this paper is to integrate elements of this prior work, prior work on manifold learning and novel graph-based formulations into a multiscale framework to develop new multiscale graph-based methods for machine learning tasks, such as data classification. Our methods will be able to deal with a variety of scales present in many data sets.

2.2 Graph-based framework

The methods presented in this paper use a similarity graph framework consisting of a graph $G = (V; E)$, where $V = \{x_1, \dots, x_N\}$ is a set of vertices associated with the elements of the data set, and E is a set of edges connecting some pairs of vertices. The edges are weighted by a weight function $w: V \times V \rightarrow \mathbb{R}$, where $w(x_i; x_j)$ measures the degree of similarity between x_i and x_j . Larger values indicate similar elements and smaller values indicate dissimilar elements. Naturally, the embedding of data into a graph depends greatly on the edge weights. This section provides more details about graph construction, but the exact manner of weight construction for particular data sets is described in Section 4.

The use of the graph-based framework offers many advantages. First, it provides valuable information about the extent of similarity between pairs of elements of both labeled and unlabeled data via a weighted similarity graph and also yields information about the overall structure of the data. This reduces the amount of labeled data needed for good accuracy. Moreover, a graph-based setting embeds the dimension of the features in the graph during weight computations, thus reducing the high-dimensionality of the problem. It also provides a way to handle nonlinearly separable classes and affords the flexibility to incorporate diverse types of data. In addition, in image processing, the graph setting allows one to capture texture more accurately.

The exact technique of computing the similarity value between two elements of data depends on the data set, but first involves feature (attribute) vector construction and a distance metric chosen specifically for the data and task at hand. For example, for hyperspectral data, one may choose the feature vector to be the vector of intensity values in its many bands and the distance measure to be the cosine distance. For 3D sensory data, one can take the feature vector to contain both geometric and color information; the weights can be calculated using a Gaussian function incorporating normal vectors, e.g., [7]. For text classification, popular feature extraction methods include term frequency- inverse doc-

ument frequency and bag-of-words, both described in [2]. For biological data tasks, such as protein classification, persistent homology [14] can be used for feature construction.

Once the features are constructed, the weights are computed. Popular weight functions include the Zelnik-Manor and Perona function [83] and the Gaussian function [98]:

$$w(x_i; x_j) = \exp \frac{d(x_i; x_j)^2}{s^2}; \quad (1)$$

where $d(x_i; x_j)$ represents a distance between feature vectors of data elements x_i and x_j , and $s > 0$. Using the weight function w , one can construct a weight matrix W defined as $W_{ij} = w(x_i; x_j)$, and define the degree of a vertex $x_i \in V$ as $d(x_i) = \sum_j w(x_i; x_j)$. If D is the diagonal matrix with elements $d(x_i)$, then the graph Laplacian is defined as

$$L = D - W; \quad (2)$$

It is sometimes beneficial to use normalized versions of the graph Laplacian, such as a symmetric graph Laplacian [98].

For some data, it is more desirable to compute the weights directly by calculating pairwise distances. In this case, the efficiency can be increased by using parallel computing or by reducing the dimension of data. Then, a graph is often made sparse using, e.g., thresholding or a nearest neighbors technique, resulting in graph where most of the edge weights are zero. Thus, the number of needed computations is reduced. Overall, a nearest neighbor graph can be computed efficiently using the kd-tree code of VLFeat library [3]. In particular, for the nearest neighbor technique, vertices x_i and x_j are connected only if x_i is among the N_n nearest neighbors of x_j or if x_j is among the N_n nearest neighbors of x_i . Otherwise, $w(x_i; x_j)$ is set to 0.

For very large data sets, one can efficiently construct an approximation to the full graph using e.g. sampling-based approaches, such as the fast Nyström Extension method [31].

2.3 Semi-supervised setting

Despite the tremendous accomplishments of machine learning, its success depends on a sufficient amount of labeled samples. However, obtaining enough labeled data is difficult as it is time-consuming and expensive. Therefore, labeled data is scarce for most applications.

However, unlabeled data is usually easier and less costly to obtain than labeled data. Therefore, it is advantageous to use a semi-supervised setting, which uses both labeled and unlabeled data to construct the graph in order to reduce the amount of labeled data needed for good accuracy. In fact, the use of unlabeled data for graph construction allows one to obtain structural information of the data. Overall, for most graph-based semi-supervised methods, the majority of data embedded into a graph is unlabeled data. This paper derives methods which use a semi-supervised setting of this kind.

3 Methods

3.1 Background and related graph Laplacian methods

3.1.1 Manifold learning

For the derivation of the MML method, let K be the number of classes, L be the set of labeled vertices, and U be the set of unlabeled vertices. We assume that L is drawn from the joining distribution P on $V \times R$, while U is drawn from the marginal distribution P_V of P . We also assume that the conditional distribution $P(y|x)$ varies smoothly in the intrinsic geometry generated by P_V , where $y \in [1; K]$ and $x \in V$.

In graph-based methods, information about labeled data and the geometric structure of the marginal distribution P_V of the unlabeled samples is incorporated into the problem:

$$f = \arg \min_{f \in \mathbb{R}^{|H_M|}} \frac{1}{|L|} \sum_{x_i \in L} J(f; x_i; y_i) + g_A \|f\|_{M^*}^2 + g_I \|f\|_I^2; \quad (3)$$

where the Mercer kernel $M : V \times V \rightarrow \mathbb{R}$ uniquely defines a reproducing kernel Hilbert space (RKHS) H_M with the corresponding norm $\|\cdot\|_{M^*}$, J is a loss function which gives rise to different types of regularization problems, $g_A > 0$, $g_I > 0$, and $\|f\|_I^2$ is an additional regularizer that reflects the intrinsic geometry of P_V . The solution f to (3) can be described using the classical representer theorem [1]:

$$f(x) = \sum_{x_i \in L} a_i M(x_i; x) + \sum_S a(z) M(x; z) dP_V(z); \quad (4)$$

where S is the support of the marginal distribution P_V [9].

In practice, that marginal distribution is unknown. In spite of that, one could empirically estimate $\|f\|_I^2$ by making use of the weighted graph as discussed in Section 2.2. With the pre-defined graph Laplacian matrix L , the manifold regularizer $\|f\|_I^2$ can be empirically estimated [9] as

$$\|f\|_I^2 = \sum_{i,j=1}^n (f(x_i) - f(x_j))^2 w_{ij} = f^T L f; \quad (5)$$

where $f = [f(x_1); f(x_2); \dots; f(x_n)]^T$.

The ambient norm $\|\cdot\|_{M^*}$ and the intrinsic norm $\|\cdot\|_I$ in (3) can be integrated in one term under the warped kernel $\tilde{M}$ [94]. This kernel defines an alternative reproducing kernel Hilbert space $\tilde{H}$ by considering a modified inner product:

$$\langle f; g \rangle_{\tilde{H}_M} = \langle f; g \rangle_{H_M} + f^T P g; \quad (6)$$

where P is a positive semi-definite matrix defined on labeled and unlabeled data, $f = [f(x_1); f(x_2); \dots; f(x_n)]^T$ and $g = [g(x_1); g(x_2); \dots; g(x_n)]^T$. With $h; : i_H \rightarrow \cdot$, the warped kernel $\tilde{M}$ is shown in [94] to have the following representation:

$$\tilde{M}(x; z) = M(x; z) - M_x^T (I + P M)^{-1} P M_z; \quad (7)$$

where $M = [m_{ij}]$ is the Gram matrix with $m_{ij} = M(x_i; x_j)$, M_x denotes the vector $(M(x_1; x); M(x_2; x); \dots; M(x_n; x))^T$, and M_z denotes the vector $(M(z_1; x); M(z_2; x); \dots; M(z_n; x))^T$.

The regularization problem for the warped kernel $\tilde{M}$ is:

$$f = \arg \min_{f \in \mathbb{R}^{|H_{\tilde{M}}|}} \frac{1}{|L|} \sum_{x_i \in L} J(f; x_i; y_i) + g_A \|f\|_{\tilde{M}^*}^2; \quad (8)$$

Problem (8) exploits the intrinsic geometry of P_V via the data-dependent kernel $\tilde{M}$ but still makes use of the classical regularization solvers. In fact, the classical representer theorem [1] allows f in (8) to be expressed as:

$$f(x) = \sum_{x_i \in L} a_i M(x; x_i); \quad (9)$$

In practice, $a_i; g$ are numerically determined by an appropriate optimization solver, e.g., [26].

3.1.2 MBO reduction

For the derivation of the MMBO method, we first note that a typical learning algorithm involves finding an optimal label matrix $U = (u_1; \dots; u_N)^T$ associated with data elements, where $u_i \in \mathbb{R}^K$ represents the probability distribution over the classes for data element x_i ; the i^{th} row of U is set to u_i . The vector u_i is an element of the Gibbs simplex:

$$S^K := \{f(z_1; \dots; z_K) \in [0; 1]^K \mid \sum_{k=1}^K z_k = 1\}; \quad (10)$$

where K is the number of classes. Moreover, the k^{th} vertex of the simplex is given by the unit vector e_k .

A general form of a graph-based problem for data classification is the minimization of $E(U) = R(U) + \text{Fid}(U)$, where U is the data classification function, $R(U)$ is a graph-based regularization term incorporating the graph weights, and $\text{Fid}(U)$ is a term incorporating labeled points.

Not surprisingly, the choice of the regularization term has non-trivial consequences in the final accuracy. In [35], Garcia et al. successfully take for the regularization term a multiclass graph-based Ginzburg-Landau (GL) functional:

$$GL(U) = \frac{e}{2} hU; LU + \frac{1}{2e} \sum_i \tilde{a}_i \sum_{k=1}^K \frac{1}{4} \|ku_i - e_k\|_{L_1}^2; \quad (11)$$

Here, $e > 0$, L is a normalized graph Laplacian, K is the number of classes, $hU; LU = \text{trace}(U^T L U)$, u_i is the i^{th} row of U , u_i is a vector indicating prior class knowledge of x_i , e_k is an indicator vector of size K with one in the k^{th} component and zero elsewhere, and m_k is a parameter that takes the value of $m > 0$ if x_i is a labeled data element and zero otherwise. The variable $\hat{u}_i = e_k$ if x_i is a labeled element of class k . The first (smoothing) term in (11) measures variations in the vector field, while the second (potential) term in (11) drives

the system closer to the vertices of the simplex. The third (fidelity) term enables the incorporation of labeled data.

While it is possible to develop a convex splitting scheme to minimize the graph-based multiclass GL energy [35], a more efficient technique involves MBO reduction. Specifically, if one considers the minimization of the GL functional plus a fidelity term (consisting of a fit to elements of known class) in the continuous case, one can apply L_2 gradient descent resulting in a modified Allen-Cahn equation. If a time-splitting scheme is then applied, one obtains a procedure where one alternates between propagation using the heat equation with a forcing term and thresholding. In such a state, the resulting procedure has similar elements to the MBO scheme [68], which evolves an interface by mean curvature, in a continuous, rather than graph-based, setting. The procedure can then be transferred to a graph-based setting using [35, 64, 65]. Moreover, in order for the scheme to be applicable to the multiclass case, one can convert the thresholding operation to the displacement of the vector field variable towards the closest vertex in (10) [35, 64, 65].

3.2 The derivation of the multiscale setting and the proposed methods

3.2.1 Multiscale graph Laplacian operator

The dominance of multiscale information over the single one has been proved in various biophysic-related works, such as those involving thermal fluctuation predictions [81, 105] and binding affinity predictions [72]. Therefore, it is promising to explore how the multiscale approach can improve the accuracy of graph-based data classification. We examine a novel multiscale graph Laplacian in the form of

$$L_{\text{multiscale}} = \sum_{t=0}^m c_t L_t^{p_t}; \quad (12)$$

where $p_t > 0$, $c_t > 0$, and L_t is an extended Laplacian matrix defined by $L_t = D_t - W_t$, where D_t is a degree matrix, and W_t is an extended adjacent graph edge matrix

$$[W_t]_{ij} = \frac{1}{s_t} H_t^{jjx_i} \frac{x_{jj}}{s_t} e^{-\frac{jjx_i x_{jj}^2}{s_t^2}} \quad (13)$$

where $s_t > 0$ and H_t is the t^{th} order Hermite polynomial. Usually, only two or three multiscale Laplacian terms in (12), i.e., $m = 1$ or $m = 2$, are needed to obtain a significant improvement in accuracy; by setting $m = 0$ and $c_0 = 1$, one can restore the regular graph Laplacian discussed in (2). In this formulation, s_t is automated scale filtration parameter that controls the shape of a submanifold for a data set, while c_t weighs contributions from different scales. The parameters c_t and s_t may vary for different Hermite polynomials.

In case of large data for which computing all the graph weights can be computationally expensive, one can use the Nyström extension method [10, 31, 32] to compute approximations to the few smallest eigenvalues and corresponding eigenvectors of the multiscale graph Laplacian while calculating only a small fraction of the graph weights. We will modify the Nystrom procedure to incorporate the new multiscale graph Laplacian $L_{\text{multiscale}}$. In this case, the weights in the procedure are computed using

$$W_{\text{multiscale}} = \sum_{t=0}^m c_t W_t^{p_t}; \quad (14)$$

where, in most cases, $m = 1$ or $m = 2$ is enough to obtain a significant accuracy improvement.

When the number of data elements is not too large, one can compute the eigenvectors via the Rayleigh-Chebyshev method [6] or the Shifted Block Lanczos algorithm [40].

3.2.2 Multikernel manifold learning (MML) scheme

In multikernel manifold learning (MML), the multiscale Laplacian matrices proposed in (12) is employed to form N_n -nearest neighbors subgraphs. By setting $P = \frac{g_L}{g_A} W_{\text{multiscale}}$ in (7), we attain an MML scheme enabling the reconstruction of the regularization problem presented in (3). Even with the integration of multiscale Laplacian operator into the data kernel, the manifold learning algorithms still retains its classical representation as presented in (8). One, therefore, could utilize traditional solvers to derive the multiscale manifold learning's optimizer [94]. The MML procedure is summarized as Algorithm 1.

Algorithm 1 MML Algorithm (multiscale)

Require: labeled data $L = f(x_i; y_i)g_i$, where y_i is the label of x_i , unlabeled data $U = f(x_j)g_j$, N_n (# of nearest neighbors), $m + 1$ (# of scales), where 2 or 3 scales is usually sufficient, $f_{c_t} g_{t=0}^m$ (Laplacian matrix coefficients), $f_{p_t} g_{t=0}^m$ (matrix powers), $f_{s_t} g_{t=0}^m$ (kernel scales), and $g_L > 0$, $g_A > 0$ (scalars).

Ensure: Estimated optimizer f , where $f(x)$ is the prediction for x .

- 1: Construct $m + 1$ multiscale subgraphs with N_n nearest neighbors with weights $[W_t]_{ij}$ for $t = 0; \dots; m$, where it is usually sufficient to use $m = 1$ or $m = 2$, i.e. two or three scales.
 - 2: Select the kernel $M(x; x_i)$, e.g., radial basis function kernel or a Gaussian kernel.
 - 3: Compute the Gram matrix $M = [m_{ij}]$ with $m_{ij} = M(x_i; x_j)$.
 - 4: Compute the multiscale Laplacian $L_{\text{multiscale}}$ using (12) and $f_{c_t} g_{t=0}^m$, $f_{p_t} g_{t=0}^m$ and $f_{s_t} g_{t=0}^m$.
 - 5: Formulate the warped kernel $M(x; x_i)$ using (7) and $P = L_{\text{multiscale}}$.
 - 6: Solve for optimizer of (8) using an SVM quadratic programming solver for soft margin loss, e.g., [26].
-

3.2.3 Multiscale MBO (MMBO) scheme

Our proposed MMBO scheme uses a semi-implicit approach where the multiscale Laplacian term is computed implicitly due to the stiffness of the operator which is caused by a wide range of its eigenvalues. An implicit term is needed since an explicit scheme requires all scales of eigenvalues to be resolved numerically.

To derive the MMBO scheme, let U represent a matrix where each row is a probability distribution of each data element over the classes and let u_i represent the i^{th} row of U . In addition, let N be the number of data set elements, K be the number of classes, $dt > 0$, and m be a vector which takes a value m_i in the i^{th} place if x_i is a labeled element and 0 otherwise. Moreover, let U_{labeled} be the following matrix: for rows corresponding to labeled points, the entry corresponding to the class of the labeled point is set to 1. All other entries of the matrix are set to 0. Lastly, let $m(U \cdot U_{\text{labeled}})$ indicate row-wise multiplication by a scalar.

As described in Section 3.1.2, if one considers the minimization of a GL functional plus a fit to elements of known class in the continuous case, an L_2 gradient descent results in a modified Allen-Cahn equation. If a time-splitting scheme is then applied, one obtains a procedure where one alternates between propagation using the heat equation with a forcing term and thresholding. The scheme can then be transferred to a graph-based setting and the Laplace operator can be replaced by a graph-based multiscale Laplacian. The thresholding can be changed to the displacement of the variable towards the closest vertex in (10). A projection to the simplex is then necessary before the displacement step.

Our proposed MMBO procedure thus consists of the following procedure. Starting with an initial guess for U , obtain the next iterate of U via the following three steps:

1. Multiscale heat equation with a forcing term: $U^{n+\frac{1}{2}} = U^n - dt f L_{\text{multiscale}} U^{n+\frac{1}{2}} + m(U^n \cdot U_{\text{labeled}})g$, where m is a vector which takes a value m_i in the i^{th} place if x_i is a labeled element and 0 otherwise, and $m(U^n \cdot U_{\text{labeled}})$ indicates row-wise multiplication by a scalar.
2. Projection to simplex: Each row of $U^{n+\frac{1}{2}}$ is projected onto the simplex using [24].
3. Displacement: $u_i^{n+1} = e_k$, where u_i^{n+1} is the i^{th} row of U^{n+1} , and e_k is the indicator vector (with a value of 1 in the k^{th} place and 0 elsewhere) associated with the vertex in the simplex closest to the i^{th} row of the projected $U^{n+\frac{1}{2}}$ from step 2.

This implicit scheme allows the evolution of U to be numerically stable regardless of the time step dt , in spite of the “stiffness” of the differential equations which could otherwise force dt to be impractically small.

One can compute $U^{n+\frac{1}{2}}$ very efficiently using spectral techniques and projections onto a low-dimensional subspace

spanned by a small number of eigenfunctions in the following manner, where I is the identity:

$$U^{n+\frac{1}{2}} = X_{\text{multiscale}} (I + dt L_{\text{multiscale}})^{-1} X_{\text{multiscale}}^T U_{\text{update}}; \quad (15)$$

where $U_{\text{update}} = U^n - dt m(U^n \cdot U_{\text{labeled}})$, $X_{\text{multiscale}}$ is an $N \times N_e$ truncated matrix retaining only $N_e \ll N$ smallest eigenvectors of the multiscale graph Laplacian $L_{\text{multiscale}}$, and $L_{\text{multiscale}}$ is a $N_e \times N_e$ diagonal matrix retaining the smallest eigenvalues of $L_{\text{multiscale}}$ along the diagonal.

The proposed MMBO procedure is detailed as Algorithm 2. It is important to note that in the MMBO method, the graph weights are only used to compute the few eigenvectors and eigenvalues of the multiscale graph Laplacian, and the multiscale MMBO procedure themselves do not involve graph weights. This crucial property allows one to use the Nyström extension procedure [10,31,32] to approximate the extremal eigenvectors of the Laplacian by only computing a small portion of the graph weights; this enables one to apply the multiscale models very efficiently on large data.

For initialization, the rows of U corresponding to labeled points are set to the vertices of the simplex corresponding to the known labels, while the rows of U corresponding to the rest of the points initially represent random probability distributions over the classes.

The energy minimization proceeds until a steady state condition is reached. The final classes are obtained by assigning class k to node i if u_i is closest to vertex e_k on the Gibbs simplex. Consequently, the calculation is stopped when, for a positive constant $h > 0$,

$$\frac{\max_i |u_i^{n+1} - u_i^n|}{\max_i |u_i^{n+1}|} < h; \quad (16)$$

In regards to computational complexity, in practice, once the N_e eigenvectors of the graph Laplacian are computed, the complexity of the MMBO scheme is linear in the number of data elements N . In particular, let K be the number of classes and $m+1$ be the number of terms in the multiscale Laplacian (12). Usually, $m=1$ or $m=2$ is enough to obtain a good accuracy. Then, one needs $O(NKN_e)$ operations for the multiscale heat equation with a forcing term, $O(NK \log K)$ operations for the projection to the simplex and $O(NK)$ operations for the displacement step. Moreover, [10,31,32] allows one to compute the N_e eigenvectors of the multiscale graph Laplacian using $O(NN_e m)$ operations. Since $N_e \ll N$ and $K \ll N$, in practice, the complexity of this method is linear.

3.3 Computation of eigenvalues and eigenvectors of the multiscale graph Laplacian

The MMBO method requires one to compute a few of the smallest eigenvalues and the corresponding eigenvectors of

the multiscale graph Laplacian to form $X_{\text{multiscale}}$. We examine and use three techniques for this task. Nyström extension [10, 31, 32] is the preferred method for very large data.

Algorithm 2 MMBO Algorithm (multiscale)

Require: labeled data $L = f(x_i; y_i)g_i$, where y_i is the label of x_i , unlabeled data $U = f(x_j; y_j)g_j$, N_n (# of nearest neighbors), $m + 1$ (# of scales), $fc_t g_{t=0}^m$ (Laplacian matrix coefficients), $fp_t g_{t=0}^m$ (matrix powers), $fs_t g_{t=0}^m$ (kernel scales), $dt > 0$, N (# of data set elements), $N_e < N$ (# of eigenvectors to be computed), N_t (maximum # of iterations), m (an $N \times 1$ vector which takes a value $\min$ in the i^{th} place if x_i is a labeled element and 0 otherwise).

Ensure: $out = U^{\text{end}}$; the i^{th} row of U^{end} is a probability distribution of data element x_i over the classes.

```

1. For larger data, go to Step 4. For smaller data, go to Step 2.
2: Construct  $m + 1$  multiscale subgraphs with  $N_n$  nearest neighbors
   with weights  $[W_t]_{ij}$  for  $t = 0, \dots, m$ , where it is usually sufficient to
   use  $m = 1$  or  $m = 2$ , i.e. two or three scales.
3: Compute the multiscale Laplacian  $L_{\text{multiscale}}$  using (12) and
    $fc_t g_{t=0}^m$ ,  $fp_t g_{t=0}^m$  and  $fs_t g_{t=0}^m$ .
4: Compute  $U_{\text{labeled}}$ ,  $L_{\text{multiscale}}$  and  $X_{\text{multiscale}}$  as described in Sec-
   tion 3.2.3 and using  $N_e < N$ . For smaller data, use methods such
   as [6]. For larger data, use Nystrom extension [10, 31, 32].
5: Complete the following steps: starting with  $n = 1$ .
for  $i = 1 \dots N$  do
   $U_{ik}^0 = \text{rand}((0; 1)); u_i^0 = \text{projectToSimplex}(u_i^0)$  using [24],
  where  $i^{\text{th}}$  row of  $U^0$ .
  If  $m_i > 0$ ;  $U_{ik}^0 = U_{\text{labeled}ik}$ 
end for
 $B = (I + dt L_{\text{multiscale}})^{-1} X_{\text{multiscale}}^T$ 
while Stop criterion not satisfied or  $n > N_t$  do
   $C = U^n dt m(U^n U_{\text{labeled}})$ 
   $A = BC$ 
   $U^{n+1} = X_{\text{multiscale}} A$ 
  for  $i = 1 \dots N$  do
     $u_i^{n+1} = \text{projectToSimplex}(u_i^{n+1})$  using [24]
     $u_i^{n+1} = e_k$ , where  $k$  is closest simplex vertex to  $u_i^{n+1}$ 
  end for
  The matrix  $U^{n+1}$  is such that its  $i^{\text{th}}$  row is  $u_i^{n+1}$ .
   $n = n + 1$ 
end while
```

3.3.1 Nyström extension for fully connected graphs

Nyström extension [10, 31, 32] is a matrix completion method, and it performs faster than many other techniques because it computes approximations to eigenvectors and eigenvalues using much smaller submatrices of the original matrix.

Note that if λ is an eigenvalue of $\hat{W} = D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$, then λ is an eigenvalue of the symmetric Laplacian $L_s = I - D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$, and the two matrices have the same eigenvectors. Thus, one can use Nystrom extension to calculate approximations to the eigenvectors of $\hat{W}$ and thus of L_s .

Now, consider the problem of approximating the extremal N_e eigenvalues and eigenvectors of a full graph $\hat{W}$ and let $\hat{w}(x_i; x_j) = \hat{W}_{ij}$. Nyström extension [10, 31, 32] approximates

the eigenvalue equation using a quadrature rule and $N_e \ll N$ randomly chosen interpolation points from V , which represents data elements. Denote the set of N_e randomly chosen points by $X = f(x_i)g_{i=1}^{N_e}$ and its complement by Y . Partitioning V into $V = X \cup Y$ and letting $f_k(x)$ be the k^{th} eigenvector of W and λ_k be its associated eigenvalue, we obtain:

$$\hat{w}(y_i; x_j) f_k(x_j) = \lambda_k f_k(y_i) \quad \forall y_i \in Y; \quad \forall k = 1, \dots, N_e; \quad (17)$$

This system cannot be solved directly since the eigenvectors are unknown; thus, the N_e eigenvectors of $\hat{W}$ are approximated using much smaller submatrices of $\hat{W}$.

The efficiency of Nyström extension lies with the following fact: when computing the N_e eigenvalues and eigenvectors of an $N \times N$ matrix, where N is large, the algorithm approximates them using much smaller matrices, the largest of which has dimension $N \times N_e$, where $N_e \ll N$. In particular, when the method is applied to W or $\hat{W}$, only a small portion of the weight matrix W or $\hat{W}$ needs to be computed. In our experience, $N_e = 100$ or $N_e = 200$ were good choices.

If the number of scales is $m + 1$, the complexity of the Nystrom procedure is $O(N N_e (m + 1))$, which is linear in N .

3.3.2 Rayleigh-Chebyshev method

The Rayleigh-Chebyshev method [6] is a fast algorithm for finding a small subset of eigenvalues and eigenvectors of sparse symmetric matrices, such as a symmetric graph Laplacian which can be made sparse using techniques such as N_n -nearest neighbors. The method is a modification of an inverse subspace iteration procedure and uses adaptively determined Chebyshev polynomials.

3.3.3 A shifted block Lanczos algorithm

A shifted block Lanczos algorithm [40], as well as other variations of the Lanczos method [82] that is an adaptation of power methods, are efficient techniques for solving sparse symmetric eigenproblems and for finding a few of the extremal eigenvalues. They can be used to find a subset of the eigenvalues and eigenvectors of the symmetric graph Laplacian which can be made sparse using N_n -nearest neighbors.

4 Results and discussion

4.1 Data sets

In this work, we validate the proposed MML and MMBO methods against three common data sets:

- G50C is an artificial data set inspired by [39] and generated from two normal covariance Gaussian distributions. This data set has 550 data points located in R^{50} and two labels $f = -1; +1g$.

- USPST data set includes images of handwritten digits taken from the USPS test data set. This data has 2007 images to be classified into ten labels corresponding to ten numbers from 0 to 9.
- Mac-Win data set categorizes documents, taken from 20-Newsgroups data, into 2 classes: mac or windows [96]. This set has 1946 elements and each element is represented by a vector in $\mathbb{R}^{7511}$.
- WebKB data set is taken from the web documents of the CS department of four universities and has been used extensively. It has 1051 data samples and two labels: course and non-course. There are two ways to describe each web document: the textual content of the webpage (called page representation), and the anchor text on hyperlinks pointing from other webpages to the current one. The data points with page representation are in $\mathbb{R}^{3000}$, while the ones with link representation belong to $\mathbb{R}^{1840}$. When we combine two different kinds of representations, we achieve the data points in $\mathbb{R}^{4840}$.
- a;b-protein data set consists of three different protein domains, namely alpha proteins, beta proteins, and mixed alpha and beta proteins, classified based on protein secondary structures [15]. This data has 900 biomolecules, and each family has 300 instances.

The details of the data sets are outlined in Table 1.

4.2 Hyperparameters selection

In the MMBO setting, for each data point, we do not compute the complete graph but instead construct a N_n -nearest neighbor graph for the calculation efficiency. The parameter l is one of the hyperparameters and is selected on a case by case basis. Moreover, as discussed in Section 2.2, the weight function used is the Gaussian kernel $w(x_i; x_j) = \exp(-d(x_i; x_j)^2/s^2)$. Here, the scalar s is optimized so that it perfectly fits the labeled set information. In the multiscale approach, each kernel is assigned different s values depending on the outcome of hyperparameter selection. Overall, due to the random initialization of the non-labeled points, we use the same random seed for all the experiments in this work for reproducible purposes.

The Nyström extension method [10, 31, 32] allows for fast computations even in case of larger data since this approach approximates the eigenvalues and eigenvectors of the origin matrix using much smaller matrices randomly selected from the bigger ones. Thus, only a small portion of the graph weights need to be computed. However, in case of smaller data, it is often more advantageous to use methods such as [6] which can directly compute the eigenvalues and eigenvectors. Therefore, to obtain optimal results, we employ the Rayleigh-Chebyshev procedure [6] (see Section 3.3.2) for our experiments. This method is well-known for

efficiently calculating the smallest eigenvectors of a sparse symmetric matrix. The hyperparameters of the MMBO models are the number of leading eigenvalues (N_e), the time step for solving heat equation (dt), the constraint constant on fidelity term (m), and the number of iterations (N_t).

The hyperparameter selection for MML model is carried out in a similar fashion as that of the MMBO algorithm. The tuning parameters are: the number of nearest neighbors (N_n), the scalar factor (s), the penalty coefficient (g_A), the manifold regularizer constraint (g_l), and the Laplacian degree (p). The optimizer is solved using the primal SVM solver [60]. The optimal hyperparameters of the proposed methods are documented in the Supporting Information.

4.3 Performance and discussion

4.3.1 Non-biological data sets

The non-biological data sets we used for our experiments are the G50C, USPST, Mac-Win, and WebKB data sets. In the experiments involving these data sets, we utilize the original representations without carrying out any feature generation procedures. In addition, following the previous work [21, 93], we only consider accuracy as the main evaluation metric for the G50C, USPST, and Mac-Win data sets, and compute the Precision/Recall Breakeven Point (PRBEP) for the WebKB data set due to its imbalanced labeling.

In all cases, the results of the proposed MML and MMBO methods show promising improvements from non-multiscale frameworks. Specifically, the best performances of the algorithms are achieved with three kernels. In particular, there is a significant accuracy improvement from single kernel to two kernel architectures on the USPST data (from 86.11% to 90.57% for the MML model, and from 86.55% to 88.65% for the MMBO model) and Mac-Win data (from 89.98% to 90.01% for the MML model, and from 92.06% to 93.49% for the MMBO model). The improvement from single kernel to multi-kernel learning is less for the G50C and WebKB data, but that is to be expected since G50C is a small data set consisting of 550 samples. Furthermore, it is an artificial data set drawn from two unit covariance normal distributions. As a result, a single kernel is enough to capture the crucial structure of data. Moreover, the WebKB data poses a challenge for multiscale learning due to its imbalanced data.

In almost all experiments, the proposed models obtain the best results. In particular, in most experiments, the proposed MMBO model obtains the best results, all with three-kernel learning, while the proposed MML model obtains the best result for USPST. In particular, for G50C, the MMBO method achieves the best accuracy (95.06%), but the MML method is still comparable with its accuracy being 94.56%. Moreover, the superior performance of our proposed algorithms over the state-of-the-art models is also displayed in

Table 1: Data sets used in the experiments.

Data set	No. of classes	Sample dim.	No. of data elements	No. of labeled data
G50C	2	50	550	50
USPST	20	256	2007	50
Mac-Win	2	7511	1946	50
WebKB (page)	2	3000	1051	12
WebKB (link)	2	1840	1051	12
WebKB (page+link)	2	4840	1051	12
a; b-protein	3	50	900	720

the case of the more complex USPST data, a set of hand-written digit images with 1440 samples. While the proposed MML algorithm obtains the best accuracy at 90.57%, the MMBO method with three-kernel information still obtains a good accuracy of 88.73%. The other published approaches, such as LapRLS, obtain lower accuracies. For Mac-Win, our multi-scale models perform slightly lower than $\tilde{N}$ TSVM (94.3%) [21] and LDS (94.9%) [21]; the fact that there are only 1966 samples but the dimension of each sample is very high, i.e., 7511, might indicate noisy information which can reduce the performances of graph-based kernel models. For WebKB, our proposed methods perform extremely well. WebKB is the last data set in this category and has three different feature representations, namely, link, page and page+link. The overall performance of our proposed models is very encouraging. We see that using only one kernel already produces great results, with a little improvement in using multiple kernels. The best model is the MMBO method with 3 kernels which obtains a PRBEP at 96.22%, 97.93%, and 98.87% for the link, page, and page+link experiments, respectively. The MML method obtains the next best result with a PRBEP of 95.75%, 95.81%, and 95.84% for the link, page, and page+link experiments, respectively. After the proposed MMBO and MML methods, the next best result is obtained by LapSVM: 94.3%, 93.4%, and 94.9%.

The results for non-biological data sets are shown in Figure 1. In most experiments with non-biological data, the proposed MMBO method is clearly the most dominant. The other proposed model, the MML method, is the second best model with promising performances.

4.3.2 Alpha and beta protein classification

We also tested the proposed multiscale learning models using biological data, such as data involving protein classification. In this data, based on the secondary structure, proteins are typically grouped into three classes, namely alpha helices, beta sheets, and mixed alpha and beta domains. Figure 2 plots the secondary-structure representations of 3 types of protein structures. The data, which consists of 900 structures equally distributed into three classes, was collected by Cang et al [15] and taken from SCOPe (Structural Classification of Proteins-extended), an online database [33].

Five-fold cross validation is conducted to examine the performance of the proposed models. To preserve the unbiased information, in each fold, the test set consisted of 180 instances with 60 samples from each group. Overall, the protein data sets originally provide the coordinates and atom types for each structure. However, feature generation is needed to translate such information to a vector format suitable for machine learning algorithms. Moreover, for this data, the feature generation has to sustain crucial physical and chemical interactions such as covalent and non-covalent bonds, electrostatic, hydrogen bonds, etc. In the past few years, we have developed numerous mathematical-based feature engineering models including geometric and algebraic graph [69, 73], differential geometry [71], persistent homology [14], and persistent graph [102] for representing 3D molecular information in low dimensional representations.

We employ our geometric graph representation in [73]. In order to represent the physical and chemical properties of a biomolecule, we consider four atom types, namely C_a, C, N, and O. In particular, the protein structures are described by vectors of 50 components. Overall, the details of the parameters for the feature generated approach is provided in the Supporting Information.

Both MMBO and MML models perform well. Moreover, similarly to previous experiments, multiscale information strengthens the accuracy of both MML and MMBO approaches. In fact, there is an encouraged improvement from the one kernel model to the two kernel model, i.e., 84% to 85% accuracy for the MMBO model. There is also an improvement in the MMBO method results using three kernels, i.e. 85.11%. For the MML method, there is a slight improvement by using multiple kernels. For this data, the MMBO method outperforms its counterpart, which indicates the versatility of the MMBO algorithm when dealing with a variety of data. All results are presented in Figure 3.

4.4 Comparison Algorithms

We compare our algorithms to many recent methods, most of which are from 2015 and later.

For WebKB data, we compare classification accuracy against recent methods such as semi-supervised multi-view

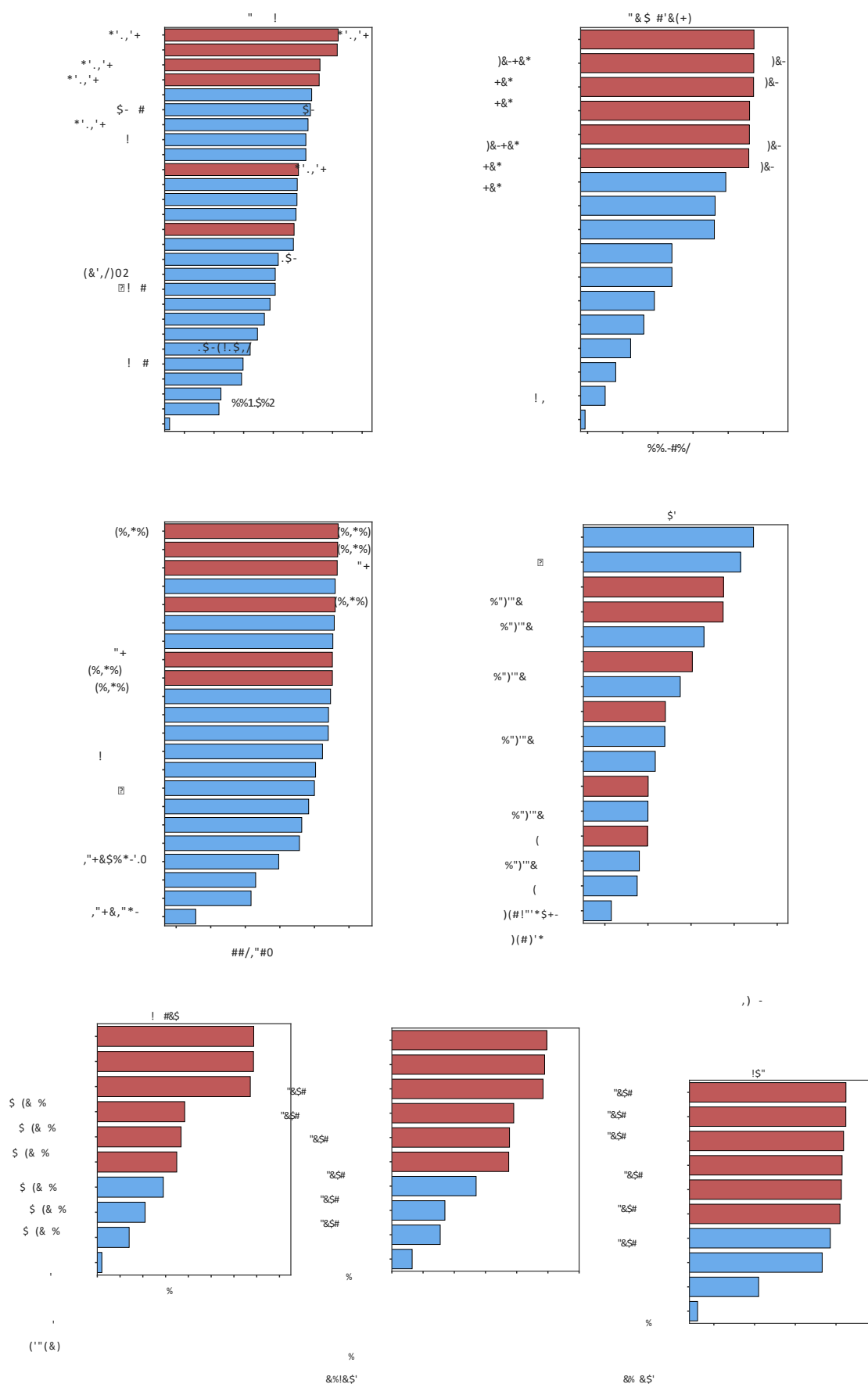


Fig. 1: Comparison of MML and MMBO with other methods on non-biological data. The proposed methods are in red, and other methods are in blue. We note that some of the comparison methods for USPST use more labeled samples than the proposed methods. Please refer to Section 4.4 for more details.

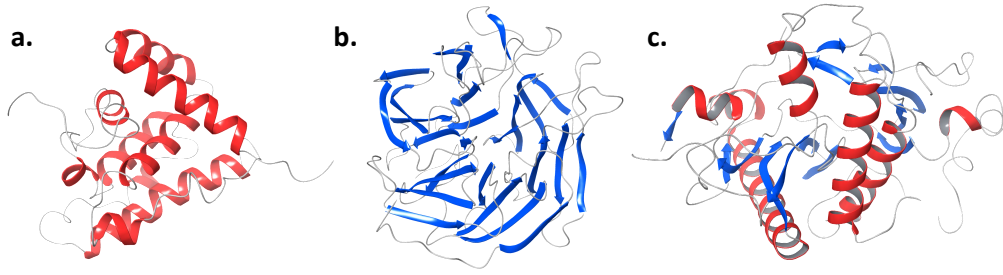


Fig. 2: Secondary-structure representations of proteins taken from a; b-protein data. Here, alpha helix is colored in red, beta sheet is colored in blue. a) Alpha protein (PDBID: 1WIX), b) Beta protein (PDBID: 3O4P), c) Mixed-alpha and beta protein (PDBID: 2CNQ). PDBID stands for protein data bank ID with experimental structures available at <https://www.rcsb.org/>.

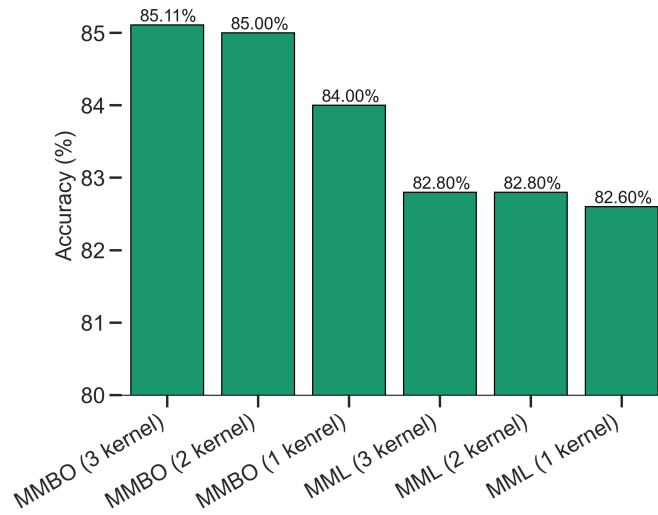


Fig. 3: The performances of MMBO and MML models on the protein classification data set.

deep discriminant representation learning (SMDDRL) [45], vertical ensemble co-training (VE-CoT) [51], auto-weighted multiple graph learning (AMGL) [78], multi-view learning with adaptive neighbors (MLAN) [75], deep canonically correlated autoencoder (DCCAE) [103], multi-view discriminative neural network (MDNN) [80], semi-supervised learning for multiple graphs by gradient flow (MGSC) [59], multi-domain classification w/ domain selection (MCS) [22], multi-view semi-supervised learning (FMSSL, FMSSL-K) [109], and semi-supervised multimodal deep learning framework (SMDLF) [25]. Our results are obtained using 105 labels, and using the classification accuracy metric. Results for SDMDRL, VE-CoT, AMGL, MLAN, SMDLF, DCCAE and MDNN are from [45], the results for MGSC and MCS are from [59], and the results for FMSSL and FMSSL-K are from [109]. All methods use 105 labels.

For USPST, we compare against recent methods such as transductive minimax probability machines (TMPM) [42], semi-supervised extreme learning machines (SS-ELM) [41], graph embedding-based dimension reduction with extreme

learning machines (GDR-ELM) [106], extreme learning machine auto-encoder (ELM-AE) [50], extreme learning machine auto-encoder with invertible functions (ELM-AEIF) [107] and extreme learning machines for dimensionality reduction (SR-ELM) [5]. Our results are obtained using only 50 labels. The results for TMPM (with 50 labels) are from [42], the results for GDR-ELM, ELM-AE, ELM-AEIF and SR-ELM (with 150 labels) are from [106], and the result for SS-ELM (with 100 labels) are from [41].

For G50C, we compare against recent methods such as clustering (CLSST) [88], semi-supervised broad learning system (SS-BLS) [112], classification from positive and unlabeled data (PNU) [87], classification from unlabeled positive and negative data (PUNU) [87], semi-supervised extreme learning machines (SS-ELM) [41], semi-supervised hierarchical extreme learning machine (SS-HELM) [92], safe semi-supervised support vector machines (S4VM) [56], robust and fast transductive support vector machines (RTSVM, RTSVM-LDS) [19]. Our results are obtained using 50 labels. The result for CLSST is from [88], the results for SS-

BLS, SS-ELM and SS-HELM are obtained from [112], the results for PNU, PUNU and S4VM are obtained from [87], and the results for RTSVM and RTSVM-LDS are obtained from [19]. All comparison methods use 50 labels.

For Mac-Win, we compare against recent methods such as support vector machines with manifold regularization and partially labeling privacy protection (SVM-MR&PLPP) [74] and a scalable version (SSVM-MR&PLPP) [74]. These results are obtained from [74]. All comparison methods and the proposed algorithms use 50 labels, the same number of labeled samples as the proposed methods.

We also compare results for all data sets with slightly older methods such as transductive graph methods (Graph-Trans), closely related to [8, 113, 116], transductive support vector machines (TSVM) [47], support vector machines on a graph-distance derived kernel (Graph-density) [21], TSVM by gradient descent ($\tilde{\text{N}}\text{TSVM}$) [21], low density separation (LDS) [21], Laplacian support vector machines (LapSVM) [93] and Laplacian regularized least squares (LapRLS) [93]. For WebKB, we use the PRBEP metric when comparing against these methods. The results for all older methods, except LapSVM and LapRLF, are obtained from [21], the results for LapSVM and LapRLF are from [93]. All comparisons with older methods use the same number of labeled samples as the proposed methods: 12 labels for WebKB and the PRBEP metric, and 50 labels for the rest of the data.

4.5 Efficiency

The proposed MML and MMBO procedures are very efficient. The timing results are listed for all data sets in Table 2 (for MMBO) and Table 3 (for MML).

The timing of the proposed MMBO method is divided into two parts: (1) the timing for the construction of the graph weights and the calculation of the extremal eigenvectors of the multiscale graph Laplacian, and (2) the timing of the MMBO procedure. From Table 2, one can see that the proposed MMBO procedure takes under 2 seconds for all data sets, and the graph construction and computation of the eigenvectors takes little time as well.

The timing of the proposed MML method consists of two categories: (1) the timing for the construction of the warped kernels, and (2) the timing of the optimizer. One can see from Table 3 that the procedure of generating the multiscale graph and the warped kernel is the most time-consuming step of the MML algorithm, but it is still under 5 seconds when handling the Mac-Win data set having 1946 samples with a feature dimension of 7511. For other data sets, the MML method takes under 0.3 seconds to formulate the multiscale graph and the warped kernel. Due to the simplified version of the optimizer of the MML method, one can

directly use the standard solver of SVM for the MML algorithm. This procedure is extremely fast and needs no more than 0.03 seconds to complete the task for all experiments. The computations were performed on a personal laptop 2.4 GHz 8-Core Intel Core i9.

5 Conclusion

This work presents several methods for machine learning tasks and for dealing with some of the challenges of machine learning, such as data with limited samples, smaller data sets, and diverse data, usually associated with small data sets or data related to areas of study where the size of the data sets is constrained by the complexity and/or high cost of experiments. In particular, we integrate graph-based techniques, multiscale structure, adapted and modified optimization procedures and semi-supervised frameworks to derive two multiscale Laplacian learning (MLL) approaches for machine learning tasks, such as data classification.

The first approach introduces a multiscale kernel representation to a manifold learning technique and is called the multikernel manifold learning (MML) algorithm.

The second approach combines multiscale analysis with an interesting adaptation and modification of the famous classical Merriman-Bence-Osher (MBO) scheme, originally intended to approximate motion by mean curvature, and is called the multiscale MBO (MMBO) algorithm.

The performance of the proposed MLL approaches is favorably compared to existing recent and related approaches through experiments on a variety of data sets. The two new MML methods form powerful techniques for dealing with some of the most important challenges and tasks in machine learning and data science.

Supporting Information

We present the optimal hyperparameters of the proposed MMBO and MML methods for all experiments conducted in this work in Online Resource: Supporting Information.

Availability

The source code for the proposed MMBO and MML methods is available at Github: <https://github.com/ddnguyenmath/Multiscale-Laplacian-Learning>.

Conflict of interest

The authors declare that they have no conflict of interest.

Table 2: The timing of the proposed MMBO method

Data set	Size of data set	Sample dimension	Timing (Construction of graph and eigenvectors)	Timing (MMBO procedure)
G50C	550	50	0.02 seconds	0.31 seconds
USPST	1440	1024	1.41 seconds	1.52 seconds
Mac-Win	1946	7511	9.8 seconds	1.17 seconds
WebKB (page)	1051	3000	1.04 seconds	0.60 seconds
WebKB (link)	1051	1840	0.67 seconds	0.60 seconds
WebKB (page+link)	1051	4840	1.58 seconds	0.60 seconds
a;b-protein	900	50	0.18 seconds	1.96 seconds

Table 3: The timing of the proposed MML method

Data set	Size of data set	Sample dimension	Timing (Deformed Kernel)	Timing (Optimization)
G50C	550	50	0.039 seconds	0.001 seconds
USPST	1440	1024	0.24 seconds	0.003 seconds
Mac-Win	1946	7511	4.51 seconds	0.002 seconds
WebKB (page)	1051	3000	0.21 seconds	0.02 seconds
WebKB (link)	1051	1840	0.15 seconds	0.01 seconds
WebKB (page+link)	1051	4840	0.28 seconds	0.02 seconds
a;b-protein	900	50	0.05 seconds	0.02 seconds

References

1. A Generalized Representer Theorem. <https://alex.smola.org/papers/2001/SchHerSmo01.pdf>
2. Quick Introduction to Bag-of-Words (BoW) and TF-IDF for Creating Features from Text. <https://www.analyticsvidhya.com/blog/2020/02/quick-introduction-bag-of-words-bow-tf-idf/>
3. VLFeat Library. <https://www.vlfeat.org>
4. Abu-El-Haija, S., Kapoor, A., Perozzi, B., Lee, J.: N-GCN: Multi-scale graph convolution for semi-supervised node classification. arXiv preprint arXiv:1802.08888 (2018)
5. Anam, K., Al-Jumaily, A.: A novel extreme learning machine for dimensionality reduction on finger movement classification using sEMG. In: International IEEE/EMBS Conference on Neural Engineering, pp. 824–827. IEEE (2015)
6. Anderson, C.: A Rayleigh-Chebyshev procedure for finding the smallest eigenvalues and associated eigenvectors of large sparse Hermitian matrices. Journal of Computational Physics 229, 7477–7487 (2010)
7. Bae, E., Merkurjev, E.: Convex variational methods on graphs for multiclass segmentation of high-dimensional data and point clouds. Journal of Mathematical Imaging and Vision 58(3), 468–493 (2017)
8. Belkin, M., Matveeva, I., Niyogi, P.: Regularization and semi-supervised learning on large graphs. In: International Conference on Computational Learning Theory, pp. 624–638. Springer (2004)
9. Belkin, M., Niyogi, P., Sindhvani, V.: Manifold regularization: A geometric framework for learning from labeled and unlabeled examples. Journal of Machine Learning Research 7(Nov), 2399–2434 (2006)
10. Belongie, S., Fowlkes, C., Chung, F., Malik, J.: Spectral partitioning with indefinite kernels using the Nystrom extension. European Conference on Computer Vision pp. 531–542 (2002)
11. Bohacek, R.S., McMartin, C., Guida, W.C.: The art and practice of structure-based drug design: a molecular modeling perspective. Medicinal Research Reviews 16(1), 3–50 (1996)
12. Boykov, Y., Veksler, O., Zabih, R.: Fast approximate energy minimization via graph cuts. In: ICCV (1), pp. 377–384 (1999). URL citeseer.ist.psu.edu/boykov99fast.html
13. Bruna, J., Zaremba, W., Szlam, A., LeCun, Y.: Spectral networks and locally connected networks on graphs. arXiv preprint arXiv:1312.6203 (2013)
14. Cang, Z., Mu, L., Wei, G.W.: Representability of algebraic topology for biomolecules in machine learning based scoring and virtual screening. PLoS Computational Biology 14(1), e1005929 (2018)
15. Cang, Z.X., Mu, L., Wu, K., Opron, K., Xia, K., Wei, G.W.: A topological approach to protein classification. Molecular Based Mathematical Biology 3, 140–162 (2015)
16. Cang, Z.X., Wei, G.W.: Analysis and prediction of protein folding energy changes upon mutation by element specific persistent homology. Bioinformatics 33, 3549–3557 (2017)
17. Cang, Z.X., Wei, G.W.: TopologyNet: Topology based deep convolutional and multi-task neural networks for biomolecular property predictions. PLoS Computational Biology 13(7), e1005690, <https://doi.org/10.1371/journal.pcbi.1005690> (2017)
18. Cang, Z.X., Wei, G.W.: Integration of element specific persistent homology and machine learning for protein-ligand binding affinity prediction. International Journal of Numerical Methods in Biomedical Engineering 34(2), DOI: 10.1002/cnm.2914 (2018)
19. Cevikalp, H., Franc, V.: Large-scale robust transductive support vector machines. Neurocomputing 235, 199–209 (2017)
20. Chapelle, O., Schölkopf, B., Zien, A.: Semi-Supervised Learning. MIT Press, Cambridge, MA (2006)
21. Chapelle, O., Zien, A.: Semi-supervised classification by low density separation. In: International Conference on Artificial Intelligence and Statistics, vol. 2005, pp. 57–64. Citeseer (2005)
22. Chen, C., Xin, J., Wang, Y., Chen, L., Ng, M.K.: A semisupervised classification approach for multidomain networks with domain selection. IEEE Transactions on Neural Networks and Learning Systems 30(1), 269–283 (2018)
23. Chen, J., Zhao, R., Tong, Y., Wei, G.W.: Evolutionary de rham-hodge method. Discrete & Continuous Dynamical Systems - B (In press, 2020)
24. Chen, Y., Ye, X.: Projection onto a simplex. arXiv preprint arXiv:1101.6081 (2011)
25. Cheng, Y., Zhao, X., Cai, R., Li, Z., Huang, K., Rui, Y.: Semi-supervised multimodal deep learning for RGB-D object recog-

- dition. In: International Joint Conferences on Artificial Intelligence, pp. 3345–3351 (2016)
26. Cortes, C., Vapnik, V.: Support-vector networks. *Machine Learning* 20(3), 273–297 (1995)
 27. Couprie, C., Grady, L., Najman, L., Talbot, H.: Power watershed: A unifying graph-based optimization framework. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 33(7), 1384–1399 (2011)
 28. Elmoataz, A., Lezoray, O., Boughleux, S.: Nonlocal discrete regularization on weighted graphs: a framework for image and manifold processing. *IEEE Transactions on Image Processing* 17(7), 1047–1060 (2008)
 29. Fang, X., Xu, Y., Li, X., Fan, Z., Liu, H., Chen, Y.: Locality and similarity preserving embedding for feature selection. *Neurocomputing* 128, 304–315 (2014)
 30. Feng, S., Zhou, H., Dong, H.: Using deep neural network with small dataset to predict material defects. *Materials & Design* 162, 300–310 (2019)
 31. Fowlkes, C., Belongie, S., Chung, F., Malik, J.: Spectral grouping using the Nyström method. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 26(2), 214–225 (2004)
 32. Fowlkes, C., Belongie, S., Malik, J.: Efficient spatiotemporal grouping using the Nyström method. *Proceedings of the 2001 IEEE Computer Society Conference on Computer Vision and Pattern Recognition. CVPR 2001* 1, I–I (2001)
 33. Fox, N.K., Brenner, S.E., Chandonia, J.M.: SCOPe: Structural classification of proteins-extended, integrating SCOP and ASTRAL data and classification of new structures. *Nucleic Acids Research* 42(D1), D304–D309 (2014)
 34. Gadde, A., Anis, A., Ortega, A.: Active semi-supervised learning using sampling theory for graph signals. In: *Proceedings of the 20th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 492–501 (2014)
 35. Garcia-Cardona, C., Merkurjev, E., Bertozzi, A.L., Flenner, A., Percus, A.: Fast multiclass segmentation using diffuse interface methods on graphs. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2014)
 36. Gerhart, T., Sunu, J., Lieu, L., Merkurjev, E., Chang, J.M., Gilles, J., Bertozzi, A.L.: Detection and tracking of gas plumes in LWIR hyperspectral video sequence data. In: *SPIE Conference on Defense, Security, and Sensing*, pp. 87430J–87430J (2013)
 37. Goldberg, A.B., Zhu, X., Wright, S.: Dissimilarity in graph-based semi-supervised classification. In: *Artificial Intelligence and Statistics*, pp. 155–162 (2007)
 38. Gong, C., Tao, D., Maybank, S., Liu, W., Kang, G., Yang, J.: Multi-modal curriculum learning for semi-supervised image classification. *IEEE Transactions on Image Processing* 25(7), 3249–3260 (2016)
 39. Grandvalet, Y., Bengio, Y.: Semi-supervised learning by entropy minimization. In: *Advances in Neural Information Processing Systems*, pp. 529–536 (2005)
 40. Grimes, R.G., Lewis, J.G., Simon, H.D.: A shifted block Lanczos algorithm for solving sparse symmetric generalized eigenproblems. *SIAM Journal on Matrix Analysis and Applications* 15(1), 228–272 (1994)
 41. Huang, G., Song, S., Gupta, J., Wu, C.: Semi-supervised and unsupervised extreme learning machines. *IEEE Transactions on Cybernetics* 44(12), 2405–2417 (2014)
 42. Huang, G., Song, S., Xu, Z.E., Weinberger, K.: Transductive min-max probability machine. In: *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pp. 579–594. Springer (2014)
 43. Hudson, D.L., Cohen, M.E.: *Neural networks and artificial intelligence for biomedical engineering*. Institute of Electrical and Electronics Engineers (2000)
 44. Iscen, A., Tolias, G., Avrithis, Y., Chum, O.: Label propagation for deep semi-supervised learning. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 5070–5079 (2019)
 45. Jia, X., Jing, X.Y., Zhu, X., Chen, S., Du, B., Cai, Z., He, Z., Yue, D.: Semi-supervised multi-view deep discriminant representation learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2020)
 46. Jiang, J., Wang, R., Wang, M., Gao, K., Nguyen, D.D., Wei, G.W.: Boosting tree-assisted multitask deep learning for small scientific datasets. *Journal of Chemical Information and Modeling* 60(3), 1235–1244 (2020)
 47. Joachims, T., et al.: Transductive inference for text classification using support vector machines. In: *International Conference on Machine Learning*, vol. 99, pp. 200–209 (1999)
 48. Jordan, M.I., Mitchell, T.M.: Machine learning: Trends, perspectives, and prospects. *Science* 349(6245), 255–260 (2015)
 49. Kapoor, A., Ahn, H., Qi, Y., Picard, R.W.: Hyperparameter and kernel learning for graph based semi-supervised classification. In: *Advances in Neural Information Processing Systems*, pp. 627–634 (2006)
 50. Kasun, L., Yang, Y., Huang, G.B., Zhang, Z.: Dimension reduction with extreme learning machine. *IEEE Transactions on Image Processing* 25(8), 3906–3918 (2016)
 51. Katz, G. and Caragea, C., Shabtai, A.: Vertical ensemble co-training for text classification. *ACM Transactions on Intelligent Systems and Technology* 9(2), 1–23 (2017)
 52. Kipf, T.N., Welling, M.: Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:1609.02907* (2016)
 53. Kotsiantis, S.B., Zaharakis, I., Pintelas, P.: Supervised machine learning: A review of classification techniques. *Emerging Artificial Intelligence Applications in Computer Engineering* 160(1), 3–24 (2007)
 54. Levin, A., Rav-Acha, A., Lischinski, D.: Spectral matting. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 30(10), 1699–1712 (2008)
 55. Li, Q., Han, Z., Wu, X.M.: Deeper insights into graph convolutional networks for semi-supervised learning. In: *Thirty-Second AAAI Conference on Artificial Intelligence* (2018)
 56. Li, Y.F., Zhou, Z.H.: Towards making unlabeled data never hurt. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 37(1), 175–188 (2014)
 57. Liao, R., Brockschmidt, M., Tarlow, D., Gaunt, A., Urtasun, R., Zemel, R.S.: Graph partition neural networks for semi-supervised classification (2018). URL <https://openreview.net/forum?id=rk4Fz2e0b>
 58. Lin, F., Cohen, W.W.: Semi-supervised classification of network data using very few labels. In: *2010 International Conference on Advances in Social Networks Analysis and Mining*, pp. 192–199. IEEE (2010)
 59. Liu, Y., Ng, M.K., Zhu, H.: Multiple graph semi-supervised clustering with automatic calculation of graph associations. *Neurocomputing* 429, 33–46 (2021)
 60. Melacci, S., Belkin, M.: Laplacian support vector machines trained in the primal. *Journal of Machine Learning Research* 12(3) (2011)
 61. Meng, G., Merkurjev, E., Koniges, A., Bertozzi, A.L.: Hyperspectral video analysis using graph clustering methods. *Image Processing On Line* 7, 218–245 (2017)
 62. Merkurjev, E., Bertozzi, A.L., Chung, F.: A semi-supervised heat kernel pagerank mbo algorithm for data classification. *Communications in Mathematical Sciences* 16(5), 1241–1265 (2018)
 63. Merkurjev, E., Bertozzi, A.L., Lerman, K., Yan, X.: Modified Cheeger and ratio cut methods using the Ginzburg-Landau functional for classification of high-dimensional data. *Inverse Problems* 33(7), 074003 (2017)

64. Merkurjev, E., Garcia-Cardona, C., Bertozzi, A.L., Flenner, A., Percus, A.: Diffuse interface methods for multiclass segmentation of high-dimensional data. *Applied Mathematics Letters* 33, 29–34 (2014)
65. Merkurjev, E., Kostic, T., Bertozzi, A.L.: An MBO scheme on graphs for segmentation and image processing. *SIAM Journal of Imaging Sciences* 6(4), 1903–1930 (2013)
66. Merkurjev, E., Sunu, J., Bertozzi, A.L.: Graph MBO method for multiclass segmentation of hyperspectral stand-off detection video. *Proceedings of IEEE International Conference on Image Processing* (2014)
67. Merriman, B., Bence, J.K., Osher, S.: Diffusion generated motion by mean curvature. *AMS Selected Lectures in Mathematics Series: Computational Crystal Growers Workshop 8966*, 73–83 (1992)
68. Merriman, B., Bence, J.K., Osher, S.J.: Motion of multiple functions: a level set approach. *Journal of Computational Physics* 112(2), 334–363 (1994). DOI 10.1006/jcph.1994.1105. URL <http://dx.doi.org/10.1006/jcph.1994.1105>
69. Nguyen, D., Wei, G.W.: Agl-score: Algebraic graph learning score for protein-ligand binding scoring, ranking, docking, and screening. *Journal of Chemical Information and Modeling* (2019)
70. Nguyen, D.D., Cang, Z., Wei, G.W.: A review of mathematical representations of biomolecular data. *Physical Chemistry Chemical Physics* 22(8), 4343–4367 (2020)
71. Nguyen, D.D., Wei, G.W.: DG-GL: Differential geometry-based geometric learning of molecular datasets. *International Journal for Numerical Methods in Biomedical Engineering* 35(3), e3179 (2019)
72. Nguyen, D.D., Xia, K., Wei, G.W.: Generalized flexibility-rigidity index. *The Journal of Chemical Physics* 144(23), 234106 (2016)
73. Nguyen, D.D., Xiao, T., Wang, M.L., Wei, G.W.: Rigidity strengthening: A mechanism for protein-ligand binding. *Journal of Chemical Information and Modeling* 57, 1715–1721 (2017)
74. Ni, T., Chung, F.L., Wang, S.: Support vector machine with manifold regularization and partially labeling privacy protection. *Information Sciences* 294, 390–407 (2015)
75. Nie, F., Cai, G., Li, J., Li, X.: Auto-weighted multi-view learning for image clustering and semi-supervised classification. *IEEE Transactions on Image Processing* 27(3), 1501–1511 (2017)
76. Nie, F., Cai, G., Li, X.: Multi-view clustering and semi-supervised classification with adaptive neighbours. In: *Thirty-First AAAI Conference on Artificial Intelligence* (2017)
77. Nie, F., Li, J., Li, X.: Parameter-free auto-weighted multiple graph learning: a framework for multiview clustering and semi-supervised classification. In: *IJCAI*, pp. 1881–1887 (2016)
78. Nie, F., Li, J., Li, X.: Parameter-free auto-weighted multiple graph learning: a framework for multiview clustering and semi-supervised classification. In: *International Joint Conferences on Artificial Intelligence*, pp. 1881–1887 (2016)
79. Nie, F., Wang, X., Huang, H.: Clustering and projected clustering with adaptive neighbors. In: *Proceedings of the 20th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 977–986 (2014)
80. Noroozi, V., Bahaadini, S., Zheng, L., Xie, S., Shao, W., Philip, S.Y.: Semi-supervised deep representation learning for multi-view problems. In: *IEEE International Conference on Big Data*, pp. 56–64. *IEEE* (2018)
81. Opron, K., Xia, K., Wei, G.W.: Communication: Capturing protein multiscale thermal fluctuations (2015)
82. Paige, C.C.: Computational variants of the Lanczos method for the eigenproblem. *IMA Journal of Applied Mathematics* 10(3), 373–381 (1972)
83. Perona, P., Zelnik-Manor, L.: Self-tuning spectral clustering. *Advances in Neural Information Processing Systems* (2004)
84. Qi, Z., Tian, Y., Shi, Y.: Laplacian twin support vector machine for semi-supervised classification. *Neural Networks* 35, 46–53 (2012)
85. Qu, M., Bengio, Y., Tang, J.: GMNN: Graph Markov neural networks. *arXiv preprint arXiv:1905.06214* (2019)
86. Saha, B., Gupta, S., Phung, D., Venkatesh, S.: Multiple task transfer learning with small sample sizes. *Knowledge and Information Systems* 46(2), 315–342 (2016)
87. Sakai, T., Plessis, M.C., Niu, G., Sugiyama, M.: Semi-supervised classification based on classification from positive and unlabeled data. In: *International Conference on Machine Learning*, pp. 2998–3006. *PMLR* (2017)
88. Sansone, E., Passerini, A., De Natale, F.: Clustering: Joint classification and clustering with mixture of factor analysers. In: *Proceedings of the Twenty-second European Conference on Artificial Intelligence*, pp. 1089–1095 (2016)
89. Schwab, K.: *The Fourth Industrial Revolution*. Currency (2017)
90. Shaikhina, T., Khovanova, N.A.: Handling limited datasets with neural networks in medical applications: A small-data approach. *Artificial Intelligence in Medicine* 75, 51–63 (2017)
91. Shaikhina, T., Lowe, D., Daga, S., Briggs, D., Higgins, R., Khovanova, N.: Machine learning for predictive modelling based on small data in biomedical engineering. *IFAC-PapersOnLine* 48(20), 469–474 (2015)
92. She, Q., Hu, B., Luo, Z., Nguyen, T., Zhang, Y.: A hierarchical semi-supervised extreme learning machine method for EEG recognition. *Medical & Biological Engineering & Computing* 57(1), 147–157 (2019)
93. Sindhwani, V., Niyogi, P., Belkin, M.: Beyond the point cloud: from transductive to semi-supervised learning. In: *International Conference on Machine Learning*, pp. 824–831 (2005)
94. Sindhwani, V., Niyogi, P., Belkin, M.: Beyond the point cloud: from transductive to semi-supervised learning. In: *International Conference on Machine Learning*, pp. 824–831. *ACM* (2005)
95. Subramanya, A., Bilmes, J.: Semi-supervised learning with measure propagation. *Journal of Machine Learning Research* 12, 3311–3370 (2011)
96. Szummer, M., Jaakkola, T.: Partially labeled classification with markov random walks. In: *Advances in Neural Information Processing Systems*, pp. 945–952 (2002)
97. Thekumparampil, K.K., Wang, C., Oh, S., Li, L.J.: Attention-based graph neural network for semi-supervised learning. *arXiv preprint arXiv:1803.03735* (2018)
98. Von Luxburg, U.: A tutorial on spectral clustering. *Statistics and Computing* 17(4), 395–416 (2007)
99. Wang, B., Tu, Z., Tsotsos, J.K.: Dynamic label propagation for semi-supervised multi-class multi-label classification. In: *Proceedings of the IEEE International Conference on Computer Vision*, pp. 425–432 (2013)
100. Wang, J., Jebara, T., Chang, S.F.: Semi-supervised learning using greedy max-cut. *Journal of Machine Learning Research* 14(Mar), 771–800 (2013)
101. Wang, M., Fu, W., Hao, S., Tao, D., Wu, X.: Scalable semi-supervised learning by efficient anchor graph regularization. *IEEE Transactions on Knowledge and Data Engineering* 28(7), 1864–1877 (2016)
102. Wang, R., Nguyen, D.D., Wei, G.W.: Persistent spectral graph. *International Journal for Numerical Methods in Biomedical Engineering* p. e3376 (2020)
103. Wang, W., Arora, R., Livescu, K., Bilmes, J.: On deep multi-view representation learning. In: *International Conference on Machine Learning*, pp. 1083–1092. *PMLR* (2015)
104. Weston, J., Ratle, F., Mobahi, H., Collobert, R.: Deep learning via semi-supervised embedding. In: *Neural networks: Tricks of the trade*, pp. 639–655. Springer (2012)

105. Xia, K., Opron, K., Wei, G.W.: Multiscale Gaussian network model (mGNM) and multiscale anisotropic network model (mANM). *The Journal of Chemical Physics* 143(20), 11B616_1 (2015)
106. Yang, L., Song, S., Li, S., Chen, Y., Huang, G.: Graph embedding-based dimension reduction with extreme learning machine. *IEEE Transactions on Systems, Man, and Cybernetics: Systems* (2019)
107. Yang, Y., Wu, Q.J., Wang, Y.: Autoencoder with invertible functions for dimension reduction and image reconstruction. *IEEE Transactions on Systems, Man, and Cybernetics: Systems* 48(7), 1065–1079 (2016)
108. Yang, Z., Cohen, W., Salakhudinov, R.: Revisiting semi-supervised learning with graph embeddings. In: *International Conference on Machine Learning*, pp. 40–48 (2016)
109. Zhang, B., Qiang, Q., Wang, F., Nie, F.: Fast multi-view semi-supervised learning with learned graph. *IEEE Transactions on Knowledge and Data Engineering* (2020)
110. Zhang, Y., Ling, C.: A strategy to apply machine learning to small datasets in materials science. *Npj Computational Materials* 4(1), 1–8 (2018)
111. Zhang, Y., Pal, S., Coates, M., Ustebay, D.: Bayesian graph convolutional neural networks for semi-supervised classification. In: *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 33, pp. 5829–5836 (2019)
112. Zhao, H., Zheng, J., Deng, W., Song, Y.: Semi-supervised broad learning system based on manifold regularization and broad network. *IEEE Transactions on Circuits and Systems I: Regular Papers* 67(3), 983–994 (2020)
113. Zhou, D., Bousquet, O., Lal, T.N., Weston, J., Schölkopf, B.: Learning with local and global consistency. In: S. Thrun, L.K. Saul, B. Schölkopf (eds.) *Advances in Neural Information Processing Systems* 16, pp. 321–328. MIT Press, Cambridge, MA (2004)
114. Zhou, D., Schölkopf, B.: A regularization framework for learning from graph data. In: *Workshop on Statistical Relational Learning*. *International Conference on Machine Learning*, Banff, Canada (2004)
115. Zhu, X., Ghahramani, Z.: Learning from labeled and unlabeled data with label propagation. *CMU CALD Tech Report CMU-CALD-02-107* (2002)
116. Zhu, X., Ghahramani, Z., Lafferty, J.: Semi-supervised learning using Gaussian fields and harmonic functions. In: *Proceedings of the 20th International conference on Machine learning*, pp. 912–919 (2003)
117. Zhuang, C., Ma, Q.: Dual graph convolutional networks for graph-based semi-supervised classification. In: *Proceedings of the 2018 World Wide Web Conference*, pp. 499–508 (2018)