



A nonlinear orthogonal non-negative matrix factorization approach to subspace clustering



Dijana Tolić^{a,*}, Nino Antulov-Fantulin^b, Ivica Kopriva^a

^a Laboratory for Machine Learning and Knowledge Representations, Ruder Bošković Institute, Zagreb, Croatia, Bijenicka cesta, 54, Zagreb 10000, Croatia

^b ETH Zürich, Swiss Federal Institute of Technology, Clausiusstrasse 50, COSS, CLU, Zürich 8092, Switzerland

ARTICLE INFO

Article history:

Received 28 March 2017

Revised 22 March 2018

Accepted 27 April 2018

Available online 2 May 2018

Keywords:

Subspace clustering

Non-negative matrix factorization

Orthogonality

Kernels

Graph regularization

ABSTRACT

A recent theoretical analysis shows the equivalence between non-negative matrix factorization (NMF) and spectral clustering based approach to subspace clustering. As NMF and many of its variants are essentially linear, we introduce a nonlinear NMF with explicit orthogonality and derive general kernel-based orthogonal multiplicative update rules to solve the subspace clustering problem. In nonlinear orthogonal NMF framework, we propose two subspace clustering algorithms, named kernel-based non-negative subspace clustering KNSC-Ncut and KNSC-Rcut and establish their connection with spectral normalized cut and ratio cut clustering. We further extend the nonlinear orthogonal NMF framework and introduce a graph regularization to obtain a factorization that respects a local geometric structure of the data after the nonlinear mapping. The proposed NMF-based approach to subspace clustering takes into account the nonlinear nature of the manifold, as well as its intrinsic local geometry, which considerably improves the clustering performance when compared to the several recently proposed state-of-the-art methods.

© 2018 Elsevier Ltd. All rights reserved.

1. Introduction

Introduced in [1] as a parts-based low-rank representation of the original data matrix, non-negative matrix factorization (NMF) has shown to be a useful decomposition of multivariate data [2–4]. The most important feature of NMF is the non-negativity of all elements of the matrices involved, which allows an additive parts-based decomposition of the data. This non-negativity is often encountered in real world data, providing a natural interpretation in contrast to other decomposition techniques that allow negative combinations (such as SVD). Related NMF factorizations include convex NMF, orthogonal NMF and kernel NMF [5–10].

The key idea in subspace clustering is to construct a weighted affinity graph from the initial data set, such that each node represents a data point and each weighted edge represents the similarity based on distance between each pair of points (e.g. the Euclidean distance). Spectral clustering then finds the cluster membership of the data points by using the spectrum of an affinity graph.

State-of-the-art methods in single view subspace clustering learn affinity graph matrix by imposing sparseness [11], low-rank [12] or jointly sparseness and low-rank constraints [13] on

representation matrix. In multi-view subspace clustering representation matrices across views can be learnt by utilization of independence criterion which decreases redundancy between representations [14]. Joint low-rank sparseness constrained approach can be extended to multi-view clustering [15]. The NMF methods proposed herein to handle single view subspace clustering problem can be extended to NMF-based multi-view subspace clustering [16]. Furthermore, the methods proposed by us could possibly improve performance further through post-processing step that re-assigns samples to more suitable clusters [17].

Spectral clustering can be seen as a graph partition problem and solved by the eigenvalue decomposition of the graph Laplacian matrix [18–22]. In particular, there is a close relationship between the eigenvector corresponding to the second eigenvalue of the Laplacian and the graph cut problem [23,24]. However, the complexity of optimizing graph cut objective function is high, e.g. the optimization of the normalized cut (Ncut) is known to be an NP-hard problem [5,25–27]. Spectral clustering seeks to get the relaxed solution, which is an approximate solution for the graph partition. Compared with conventional clustering algorithms, spectral clustering has advantages to converge to global optimum and performs well for the sample space of arbitrary shape [18,19,26,28].

Despite empirical success of spectral clustering, one drawback is that a mixed-signed result given by the eigenvalue decomposition of the Laplacian may lack clustering interpretability or degrade the clustering performance [2]. The computational complexity of

* Corresponding author.

E-mail addresses: dijana.tolic@irb.hr, dtolic@irb.hr (D. Tolić).

the eigenvalue decomposition is $\mathcal{O}(n^3)$, where n denotes the number of points. To avoid the computation of eigenvalues and eigenvectors, a recently established connection of the spectral clustering and non-negative matrix factorization (NMF) was utilized in [29–31]. As pointed out in [30], the formulation of non-negative spectral clustering is motivated by practical reasons: (i) one can use the update algorithms of NMF to solve spectral clustering, and (ii) NMF framework can easily incorporate additional constraints to spectral clustering algorithms.

It was shown in [30] that spectral clustering Ncut can be treated as a symmetric NMF problem of the graph affinity matrix constructed from the data matrix. Similarly, it was also proven that the Rcut spectral clustering is equivalent to the symmetric NMF of the graph affinity matrix, introducing the non-negative Laplacian embedding (NLE) [31]. Both results [30,31] only factorize the graph affinity matrix, imposing the assumption that the input data comes in as a matrix of pairwise similarities. The factorization of the graph affinity matrix was replaced with the factorization of the data matrix itself in [29], and including an additional global discriminative regularization term in [32]. However, both NMF-based NSC methods [29,32], minimize data fidelity term in the linear input space.

In this paper we propose a nonlinear orthogonal NMF approach to subspace clustering. We establish an equivalence with spectral clustering and propose two non-negative spectral clustering algorithms, named *kernel-based non-negative spectral clustering* KNSC-Ncut and KNSC-Rcut. To further explore the nonlinear orthogonal NMF framework, we also introduce a graph regularization term [4] which captures the intrinsic local geometric structure in the nonlinear feature space. By preserving the geometric structure, the graph regularization term allows the factorization method to have more discriminating power for clustering data points sampled from a submanifold which lies in a higher dimensional ambient space [4].

Recently, a similar connection between kernel PCA and spectral methods has been shown in [18,28,33,34]. Our method gives an insight into the connection between kernel NMF and spectral methods, where the kernel matrix from multiplicative updates corresponds to the nonlinear graph affinity matrix in spectral clustering. Different from [29–32], our equivalence is established by directly factorizing the nonlinearly mapped input data matrix. To the best of our knowledge, this is the first approach to non-negative spectral clustering that uses kernel orthogonal NMF.

By constraining the orthogonality of the clustering matrix during the nonlinear NMF updates, the cluster membership can be obtained directly from the orthogonal clustering matrix, avoiding the need of usual k -means clustering [29–32]. The proposed approach has a total run-time complexity of $\mathcal{O}(kn^2)$ for clustering n data points to k clusters, which is less than standard spectral clustering methods $\mathcal{O}(n^3)$ and the same complexity as the state-of-the-art methods [29,32,35].

We perform a comprehensive analysis of our approach, including the convergence proofs for the kernel-based graph regularized orthogonal multiplicative update rules. We conduct extensive experiments to compare our methods with other non-negative spectral clustering methods and further perform the sensitivity analysis of the parameters used in our approach. We highlight here the main contributions of the paper:

1. We formulate a nonlinear NMF with explicitly enforced orthogonality to address the subspace clustering problem.
2. We derive kernel-based orthogonal multiplicative updates to solve the constrained non-convex nonlinear NMF problem. We perform the convergence analysis for the multiplicative updates and give the convergence proofs using an auxiliary function approach [36].

Table 1
Notations.

Notation	Definition
m	the dimensionality of a data set
n	the number of data points
k	the number of clusters
$\mathcal{L}$	the Lagrangian
$\mathbf{K} \in \mathbb{R}^{n \times n}$	the kernel matrix
$\mathbf{X} \in \mathbb{R}^{m \times n}$	the input data matrix
$\mathbf{A} \in \mathbb{R}^{n \times n}$	the graph affinity matrix
$\mathbf{D} \in \mathbb{R}^{n \times n}$	the degree matrix based on $\mathbf{A}$
$\mathbf{L} \in \mathbb{R}^{n \times n}$	the graph Laplacian
$\mathbf{L}_{\text{sym}} \in \mathbb{R}^{n \times n}$	the normalized graph Laplacian
$\Phi(\mathbf{X}) \in \mathbb{R}^{D \times n}$	the nonlinear mapping
$\mathbf{H}, \mathbf{Z} \in \mathbb{R}^{k \times n}$	the cluster indicator matrices
$\mathbf{V} \in \mathbb{R}^{m \times k}$	the basis matrix in input space
$\mathbf{F} \in \mathbb{R}^{n \times k}$	the basis matrix in mapped space

3. We formulate a nonlinear (kernel-based) orthogonal graph regularized NMF approach to subspace clustering. The ability of the proposed method to exploit both the nonlinear nature of the manifold as well as its local geometric structure considerably improves the clustering performance.

4. The proposed clustering algorithms provide an insight into the connection between the spectral clustering methods and kernel NMF, where the kernel matrix in the kernel-based NMF multiplicative updates corresponds to the nonlinear graph affinity matrix in Ncut and Rcut spectral clustering.

The rest of the paper is organized as follows: In Section 2 we present a brief overview of the NMF-based spectral clustering. In Section 3, we propose our framework and present three non-negative spectral clustering algorithms, along with the theoretical results on the equivalence of our approach and non-negative spectral clustering. In Section 4, we compare our methods to the 9 recently proposed non-negative spectral clustering methods on 6 data sets. Lastly, we give the conclusions in Section 5.

2. Related work

We denote all matrices with bold upper case letters, all vectors with bold lower case letters. $\mathbf{A}^T$ denotes the transpose of the matrix $\mathbf{A}$, and $\mathbf{A}^{-1}$ denotes the inverse of the matrix $\mathbf{A}$. $\mathbf{I}$ denotes the identity matrix. The Frobenius norm is denoted as $\|\cdot\|_F$. The trace of the matrix is denoted with $\text{Tr}(\cdot)$. In Table 1 we summarize the rest of the notation.

2.1. Definitions

The task of subspace clustering is to find a low-dimensional subspace to fit each group of data points [37–40]. Let $\mathbf{X} \in \mathbb{R}^{m \times n}$ denote the data matrix $m \times n$ which is comprised of n data points $\mathbf{x}_i \in \mathbb{R}^m$, drawn from a union of k linear subspaces $S_1 \cup S_2 \cup \dots \cup S_k$ of dimensions $\{m_i\}_{i=1}^k$. Let $\mathbf{X}_i \in \mathbb{R}^{m \times n_i}$ be a submatrix of $\mathbf{X}$ of rank m_i with $\sum_{i=1}^k n_i = n$. Given the input matrix $\mathbf{X}$, subspace clustering assigns data points according to their subspaces. The first step is to construct a weighted similarity graph $G(V, E)$ from $\mathbf{X}$, such that each node from the node set $V = \{1, 2, \dots, n\}$ represents a data point $\mathbf{x}_i \in \mathbb{R}^m$ and each weighted edge represents a similarity based on distance (e.g. the Euclidean distance) between the corresponding pair of nodes. Typical methods to construct the similarity graph are ϵ -neighbourhood graphs, k -nearest neighbour graphs and fully connected graphs with Gaussian similarity function [4,41]. Spectral clustering then finds the cluster membership of data points by using the spectrum of the graph Laplacian matrix. Let $\mathbf{A} \in \mathbb{R}^{n \times n}$ be a symmetric affinity matrix of the graph and $A_{ij} \geq 0$ be the pairwise similarity between the nodes. The degree matrix $\mathbf{D}$ based on $\mathbf{A}$ is defined as the diagonal matrix with the

degrees $d_1, \dots, d_n$ on the diagonal, where the degree d_i of a node i is

$$d_i = \sum_{j=1} A_{ij} \quad (1)$$

Given a weighted graph $G(V, E)$ its unnormalized graph Laplacian matrix $\mathbf{L}$ is given as [42]

$$\mathbf{L} = \mathbf{D} - \mathbf{A} \quad (2)$$

The symmetric normalized graph Laplacian matrix $\mathbf{L}_{sym}$ is defined as

$$\mathbf{L}_{sym} = \mathbf{D}^{-1/2} \mathbf{L} \mathbf{D}^{-1/2} = \mathbf{I} - \mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \quad (3)$$

where $\mathbf{I}$ is the identity matrix.

2.2. Graph cuts

The spectral clustering can be seen as partitioning a similarity graph $G(V, E)$ into a set of nodes $S \subset V$ separated from the complementary set $\bar{S} = V \setminus S$. Depending on the choice of the function to optimize, the graph partition problem can be defined in different ways. The simplest choice of the function is the *cut* $s(S, \bar{S})$ defined as $s(S, \bar{S}) = \sum_{v_i \in S, v_j \in \bar{S}} A_{ij}$. To achieve a better balance in the cardinality of S and $\bar{S}$, the Ncut and Rcut optimization functions are proposed [42–44]. Let $\mathbf{h}_i$ be the indicator vector for cluster C_i , i.e. $\mathbf{h}_i(i) = 1$ if $\mathbf{x}_i \in C_i$, otherwise $\mathbf{h}_i(i) = 0$, then $|C_i| = \mathbf{h}_i^T \mathbf{h}_i$. The cluster indicator matrix $\mathbf{H} \in \mathbb{R}^{k \times n}$ can be defined as

$$\mathbf{H}^T = \left(\frac{\mathbf{h}_1}{\|\mathbf{h}_1\|}, \frac{\mathbf{h}_2}{\|\mathbf{h}_2\|}, \dots, \frac{\mathbf{h}_k}{\|\mathbf{h}_k\|} \right) \quad (4)$$

Evidently, $\mathbf{H}\mathbf{H}^T = \mathbf{I}$. Rcut spectral clustering can be formulated as the following optimization problem

$$\min_{\mathbf{H}} \text{Tr}(\mathbf{H}\mathbf{L}\mathbf{H}^T) \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^T = \mathbf{I} \quad (5)$$

where $\text{Tr}(\cdot)$ denotes the trace of a matrix and $\mathbf{L}$ is the graph Laplacian. Similarly, define the cluster indicator vector as $\mathbf{z}_k = \mathbf{D}^{1/2} \mathbf{h}_k / \|\mathbf{D}^{1/2} \mathbf{h}_k\|$ and the cluster indicator matrix as $\mathbf{Z}^T = (\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k)$ where $\mathbf{Z} \in \mathbb{R}^{k \times n}$. Then Ncut is formulated as the minimization problem

$$\min_{\mathbf{Z}} \text{Tr}(\mathbf{Z}\mathbf{L}_{sym}\mathbf{Z}^T) \quad \text{s.t.} \quad \mathbf{Z}\mathbf{Z}^T = \mathbf{I} \quad (6)$$

By allowing the cluster indicator matrices $(\mathbf{H}, \mathbf{Z})$ to be continuous valued the problem is solved by eigenvalue decomposition of the graph Laplacian matrix given in Eqs. (2) and (3) [18,19,28].

2.3. NMF approach to non-negative spectral clustering

The connection between the Ncut spectral clustering and symmetric NMF has been established in [30]

$$\mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} = \mathbf{H}^T \mathbf{H}, \quad \text{s.t.} \quad \mathbf{H} \geq 0. \quad (7)$$

According to the Theorem 2 from [30], enforcing symmetric factorization approximately retains the orthogonality of $\mathbf{H}$. Similarly, according to the Theorem 5 from [31] the Rcut spectral clustering has been proved to be equivalent to the following symmetric NMF problem

$$\mathbf{A} - \mathbf{D} + \sigma \mathbf{I} = \mathbf{H}^T \mathbf{H}, \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^T = \mathbf{I}, \quad \mathbf{H} \geq 0 \quad (8)$$

where σ is the largest eigenvalue of the graph Laplacian matrix $\mathbf{L}$ and the matrix $\mathbf{H} \in \mathbb{R}^{k \times n}$ contains cluster membership information that data point $\mathbf{x}_i$ belongs to the cluster c_i

$$c_i = \underset{1 \leq j \leq k}{\text{argmax}} \mathbf{H}_{ji}. \quad (9)$$

In Eqs. (7) and (8) a factorization of $n \times n$ symmetric similarity matrix $\mathbf{A}$ has a complexity $\mathcal{O}(kn^2)$ for k clusters.

Based on the results [30,31], in [29] it is proved that for non-negative input data matrix $\mathbf{X}$, and fully connected graph affinity matrix $\mathbf{A}$ given as the standard inner product $\mathbf{A} = \mathbf{X}^T \mathbf{X}$. Ncut spectral clustering is equivalent to the NMF of the scaled input data matrix (NSC-Ncut)

$$\mathbf{D}^{-1/2} \mathbf{X}^T \approx \mathbf{Z}^T \mathbf{Y} \quad \text{s.t.} \quad \mathbf{Z}\mathbf{Z}^T = \mathbf{I}, \mathbf{Z} \geq 0 \quad (10)$$

with cluster indicator matrix $\mathbf{Z} \in \mathbb{R}^{k \times n}$. Similarly, the Theorem 2 of [29] establishes the connection of Rcut non-negative spectral clustering (NSC-Rcut) and NMF problem

$$\mathbf{X}^T \approx \mathbf{H}^T \mathbf{Y} \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^T = \mathbf{I}, \mathbf{H} \geq 0 \quad (11)$$

with cluster indicator matrix $\mathbf{H} \in \mathbb{R}^{k \times n}$. Both NMF-based approaches to non-negative spectral clustering (10) and (11) are formulated in the input data space as a factorization of an input data matrix $\mathbf{X} \in \mathbb{R}^{m \times n}$ with the complexity $\mathcal{O}(nmk)$ [29]. The matrix factorization in Eqs. (10) and (11) is limited to the graph affinity matrix defined as an inner product of the input data matrix.

Furthermore, the global discriminative NMF-based NSC model introduced in [32], includes an additional nonlinear discriminative regularization term to the NMF optimization function proposed in [29]. As shown in [32], the global discriminant information greatly improves the accuracy of NSC-Ncut and NSC-Rcut [29]. Although in [32] the nonlinear character of the manifold is taken into account through the nonlinear discriminative matrix, the NMF data fidelity terms are still defined in the input data space.

3. Nonlinear orthogonal NMF approach to subspace clustering

In this section we develop a nonlinear orthogonal NMF approach to subspace clustering and establish its equivalence with Ncut and Rcut spectral clustering algorithms. We generalize the NMF objective function to a nonlinear transformation of the input data and derive kernel-based NMF update rules with explicitly imposed orthogonality constraints on the clustering matrix $\mathbf{H}$ (or $\mathbf{Z}$). Enforcing the explicit orthogonality into the multiplicative rules allows obtaining the cluster membership directly from the cluster indicator matrix. In this way, we obtain a formulation of the nonlinear NMF that explicitly addresses the subspace clustering problem.

3.1. Kernel-based orthogonal NMF multiplicative updates

In this paper we emphasize the orthogonality of the nonlinear NMF to keep the clustering interpretation while taking into account the nonlinearity of the space data are drawn from. We enforce rigorous orthogonality constraint into the NMF optimization problem and seek to obtain kernel-based orthogonal multiplicative update rules to solve it.

Let $\mathbf{X} = (\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n) \in \mathbb{R}^{m \times n}$ be the data matrix of non-negative elements. The NMF factorizes $\mathbf{X}$ into two low-rank non-negative matrices

$$\mathbf{X} \approx \mathbf{V}\mathbf{H} \quad (12)$$

where $\mathbf{V} = (\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k) \in \mathbb{R}^{m \times k}$ and $\mathbf{H}^T = (\mathbf{h}_1, \mathbf{h}_2, \dots, \mathbf{h}_k) \in \mathbb{R}^{n \times k}$ and k is a prespecified rank parameter. Generally, the rank of matrices $\mathbf{V}$ and $\mathbf{H}$ is much lower than the rank of $\mathbf{X}$ (i.e., $k \ll \min(m, n)$). The non-negative matrices $\mathbf{V}$ and $\mathbf{H}$ are obtained by solving the following minimization problem

$$\min_{\mathbf{V}, \mathbf{H} \geq 0} \|\mathbf{X} - \mathbf{V}\mathbf{H}\|_F^2 \quad (13)$$

Consider now a nonlinear transformation (a mapping) to the higher D -dimensional (or infinite) space $\mathbf{x}_i \rightarrow \Phi(\mathbf{x}_i)$ or $\mathbf{X} \rightarrow \Phi(\mathbf{X}) = (\Phi(\mathbf{x}_1), \Phi(\mathbf{x}_2), \dots, \Phi(\mathbf{x}_n)) \in \mathbb{R}^{D \times n}$. The nonlinear NMF problem aims to find two non-negative matrices $\mathbf{W}$ and $\mathbf{H}$ whose

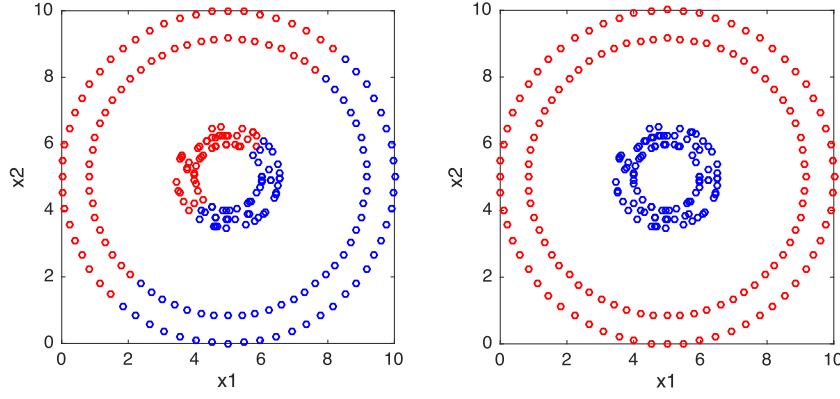


Fig. 1. Clustering with NMF (left) and nonlinear NMF (right). We apply the nonlinear NMF (KNSC-Ncut) (35) with Gaussian kernel (right) and linear NMF introduced in [1] to the synthetic data set composed of two rings and denote the cluster memberships with different colors. The nonlinear NMF is able to produce the nonlinear separating hypersurfaces between the two rings.

product can approximate the mapping of the original matrix $\Phi(\mathbf{X})$

$$\Phi(\mathbf{X}) \approx \mathbf{W}\mathbf{H} \quad (14)$$

For instance, we can consider nonlinear data set composed of two rings as in Fig. 1. The standard linear NMF (13) [45] is not able to separate the two nonlinear clusters. Compared with the solution of Eq. (17), the nonlinear NMF is able to produce the nonlinear separating hypersurfaces between the clusters. We formulate the objective function for the nonlinear orthogonal NMF as

$$\min_{\mathbf{H}, \mathbf{F} \geq 0} \|\Phi(\mathbf{X}) - \mathbf{W}\mathbf{H}\|_F^2 \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^T = \mathbf{I} \quad (15)$$

Here, $\mathbf{W}$ is the basis in feature space and $\mathbf{H}$ is the clustering matrix. It is worth noting that since Φ can be infinite dimensional, it is impossible to directly factorize $\Phi(\mathbf{X})$ [7,21,22]. In what follows we will derive a practical method to solve this problem, and keep the rigorous orthogonality imposed on the clustering matrix. Following [7] we restrict $\mathbf{W}$ to be a linear combination of transformed input data points, i.e., assume that $\mathbf{W}$ lies in the column space of $\Phi(\mathbf{X})$

$$\mathbf{W} = \Phi(\mathbf{X})\mathbf{F} \quad (16)$$

The Eq. (16) can be interpreted as a simple transformation to the new basis, leading to the following minimization problem

$$\min_{\mathbf{H}, \mathbf{F} \geq 0} \|\Phi(\mathbf{X}) - \Phi(\mathbf{X})\mathbf{F}\mathbf{H}\|_F^2, \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^T = \mathbf{I} \quad (17)$$

The optimization problem of Eq. (17) is convex in either $\mathbf{F}$ or $\mathbf{H}$, but not in both, meaning that the algorithm can only guarantee convergence to a local minimum [46]. The standard way to optimize (17) is to adopt an iterative, two-step strategy to alternatively optimize $(\mathbf{F}, \mathbf{H})$. At each iteration, one of the matrices $(\mathbf{F}, \mathbf{H})$ is optimized while the other one is fixed. The resulting multiplicative update rules with explicitly included orthogonality constraints are obtained as

$$H_{ij} \leftarrow H_{ij} \frac{(\alpha \mathbf{F}^T \mathbf{K} + 2\mu \mathbf{H})_{ij}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^T \mathbf{H})_{ij}} \quad (18)$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{H}^T)_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T)_{jl}} \quad (19)$$

where $\mathbf{K} \in \mathbb{R}^{n \times n}$ is the kernel matrix [47,48] defined as $\mathbf{K} \equiv \Phi^T(\mathbf{X})\Phi(\mathbf{X})$, where $\Phi(\mathbf{X})$ is a feature matrix in a nonlinear infinite-dimensional feature space.

We discuss two issues: (i) Convergence of the algorithm, (ii) correctness of the converged solution. **Correctness.** The correctness of the solution is assured by the fact that the solution at convergence will satisfy the Karush–Kahn–Tucker (KKT) conditions for

(17). The Lagrangian $\mathcal{L}$ of the above optimization problem (17) is

$$\mathcal{L} = \alpha \text{Tr}[\Phi(\mathbf{X})\Phi^T(\mathbf{X})] - 2\alpha \text{Tr}[\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\Phi^T(\mathbf{X})] + \alpha \text{Tr}[\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\mathbf{H}^T\mathbf{F}^T\Phi^T(\mathbf{X})] + \mu \|\mathbf{H}\mathbf{H}^T - \mathbf{I}_k\|_F^2 \quad (20)$$

By computing the partial derivatives of (20) with respect to $\mathbf{H}$ and $\mathbf{F}$, we obtain

$$\frac{\partial \mathcal{L}}{\partial \mathbf{H}} = -2\alpha \mathbf{F}^T \Phi^T(\mathbf{X})\Phi(\mathbf{X}) + 2\alpha \mathbf{F}^T \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{F}\mathbf{H} + 4\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n}) \quad (21)$$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{F}} = -\alpha \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{H}^T + \alpha \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\mathbf{H}^T \quad (22)$$

Substituting the quadratic terms with the kernel matrix $\mathbf{K} = \Phi^T(\mathbf{X})\Phi(\mathbf{X})$ yields

$$\alpha(\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} - \mathbf{F}^T \mathbf{K}) + 2\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n}) = 0 \quad (23)$$

$$-2\alpha \mathbf{K} \mathbf{H}^T + 2\alpha \mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T = 0 \quad (24)$$

Defining the Lagrange multiplier matrix for constraint $\mathbf{H} \geq 0$ as $\Psi = [\psi_{ij}]$ gives the KKT condition $\psi_{ij}H_{ij} = 0$. Similarly, the Lagrange multiplier matrix for constraint $\mathbf{F} \geq 0$ is given by $\Xi = [\xi_{jl}]$ and $\xi_{jl}F_{jl} = 0$. We obtain

$$[\alpha(\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} - \mathbf{F}^T \mathbf{K}) + 2\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n})]_{ij} H_{ij} = 0 \quad (25)$$

$$[2\alpha \mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T - 2\alpha \mathbf{K} \mathbf{H}^T]_{jl} F_{jl} = 0 \quad (26)$$

Separating positive and negative parts of the gradient leads to the multiplicative update rules (33) and (32).

Convergence. The convergence is proved by following the auxiliary function method in [7,31]. As shown in [7], these update rules guarantee the decrease of the error and eventual convergence to local minima. Note that in [7] a more general proof of the convergence can be obtained, for semi-nonnegative matrix factorization, where input data matrix is negative $\mathbf{X} < 0$. We provide the proof for the convergence in the Appendix B.

3.2. Kernel-based orthogonal NMF and spectral clustering

A connection between spectral clustering and factorization of the graph affinity matrix $\mathbf{A}$ was demonstrated in [30] for Ncut spectral clustering, and for Rcut spectral clustering in [31]. It was also shown that the spectral clustering can be viewed as a factorization of the (scaled) data matrix itself [29]. Our question is whether the spectral clustering can be viewed as a non-negative

factorization of the input data matrix mapped to a nonlinear feature space. From Eq. (12) it can be seen that the Ncut spectral clustering is equivalent to the optimization problem

$$\max_{\mathbf{Z} \geq 0} \text{Tr}(\mathbf{Z}^{\mathbf{T}} \mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \mathbf{Z}) \quad \text{s.t.} \quad \mathbf{Z} \mathbf{Z}^{\mathbf{T}} = \mathbf{I} \quad (27)$$

Theorem 1. Let $\mathbf{X} \geq 0$ denote the input data matrix. Let the similarity between the data points be defined as the inner product in the nonlinear feature space, i.e. the graph affinity matrix $\mathbf{A} = \Phi^{\mathbf{T}}(\mathbf{X})\Phi(\mathbf{X})$. Then the k -way Ncut spectral clustering (27) is equivalent to the non-negative matrix factorization of the scaled input data matrix mapped to the nonlinear feature space $\Phi(\mathbf{X})\mathbf{D}^{-1/2} = \mathbf{W}\mathbf{Z}$ subject to $\mathbf{Z}\mathbf{Z}^{\mathbf{T}} = \mathbf{I}$, where $\mathbf{W} = \Phi(\mathbf{X})\mathbf{F}$ and $\mathbf{Z}$ and $\mathbf{F}$ are two non-negative matrices, and the columns of $\mathbf{Z}$ serve as a clustering indicator vector of each data point.

The proof of the Theorem 1 is given in the Appendix A. Theorem 1 shows that Ncut spectral clustering can be viewed as a nonlinear orthogonal NMF problem with the scaling factor $\mathbf{D}^{-1/2}$. For the Rcut spectral clustering we cannot obtain an exact equivalence. However, we can relax the Rcut spectral clustering and get an equivalence between the relaxed Rcut spectral clustering and nonlinear orthonormal NMF.

Theorem 2. Let $\mathbf{X} \geq 0$ denote the input data matrix. Let the similarity between the data points be defined by inner product in nonlinear feature space i.e. the affinity matrix $\mathbf{A} = \Phi^{\mathbf{T}}(\mathbf{X})\Phi(\mathbf{X})$. Then the k -way relaxed Rcut spectral clustering (11) is equivalent to the non-negative matrix factorization of the data matrix $\Phi(\mathbf{X}) = \mathbf{W}\mathbf{H}$ subject to $\mathbf{H}\mathbf{H}^{\mathbf{T}} = \mathbf{I}$, where $\mathbf{W} = \Phi(\mathbf{X})\mathbf{F}$ and $\mathbf{H}$ and $\mathbf{F}$ are two non-negative matrices, and the columns of $\mathbf{H}$ serve as a clustering indicator vector of each data point.

The proof of the Theorem 2 is given in the Appendix A. Theorems 1 and 2 establish the nonlinear orthogonal NMF approach to non-negative spectral clustering. Our assumptions include that the similarity graph is fully connected and the similarity matrix $\mathbf{A}$ is given by the kernel $\mathbf{K} = \Phi^{\mathbf{T}}(\mathbf{X})\Phi(\mathbf{X})$. Similarly to this result, it was shown in [30] that the standard inner-product matrix $\mathbf{A} = \mathbf{X}^{\mathbf{T}}\mathbf{X}$ can be extended to any other kernel by a nonlinear transformation to a higher dimensional space.

To solve Ncut and Rcut spectral clustering we employ the kernel-based multiplicative update rules with orthonormal constraints. Considering the equivalence and solving the two optimization problems we obtain *kernel-based non-negative spectral clustering* for Ncut (KNSC-Ncut)

$$\min_{\mathbf{Z}, \mathbf{F} \geq 0} \|\Phi(\mathbf{X})\mathbf{D}^{-1/2} - \Phi(\mathbf{X})\mathbf{F}\mathbf{Z}\|_F^2, \quad \text{s.t.} \quad \mathbf{Z}\mathbf{Z}^{\mathbf{T}} = \mathbf{I} \quad (28)$$

with the following multiplicative update rule

$$Z_{ij} \leftarrow Z_{ij} \frac{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{D}^{-1/2} + 2\mu \mathbf{Z})_{ij}}{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{F} \mathbf{Z} + 2\mu \mathbf{Z} \mathbf{Z}^{\mathbf{T}} \mathbf{Z})_{ij}} \quad (29)$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{Z}^{\mathbf{T}})_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{Z} \mathbf{Z}^{\mathbf{T}})_{jl}} \quad (30)$$

The parameter μ can be set so that the orthogonality of the matrix $\mathbf{Z}$ is preserved during the updates. An exact orthogonality of the clustering matrix $\mathbf{Z}$ implies each column of $\mathbf{Z}$ can have only one non-zero element, which implies that each data object belongs only to one cluster. This is hard clustering, such as in k -means [5,30]. Furthermore, KNSC-Ncut has a soft clustering interpretation [1,30,31] where a data point could belong fractionally to more than one cluster. The soft clustering membership of data point $\mathbf{x}_i$ to cluster j can be defined as a probability distribution $c_{i,j} = \mathbf{Z}_{ji} / \sum_k \mathbf{Z}_{ki}$. We summarize the KNSC-Ncut algorithm in the Algorithm 1. Similarly, the optimization problem for *kernel-based*

non-negative spectral clustering for Rcut (KNSC-Rcut)

$$\min_{\mathbf{H}, \mathbf{F} \geq 0} \|\Phi(\mathbf{X}) - \Phi(\mathbf{X})\mathbf{F}\mathbf{H}\|_F^2, \quad \text{s.t.} \quad \mathbf{H}\mathbf{H}^{\mathbf{T}} = \mathbf{I} \quad (31)$$

gives the multiplicative update rule for KNSC-Rcut

$$H_{ij} \leftarrow H_{ij} \frac{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} + 2\mu \mathbf{H})_{ij}}{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^{\mathbf{T}} \mathbf{H})_{ij}} \quad (32)$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{H}^{\mathbf{T}})_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^{\mathbf{T}})_{jl}} \quad (33)$$

and summarize the KNSC-Rcut algorithm in Algorithm 2.

The convergence of the multiplicative update rules (29)–(30), and (32)–(33), has been proved in Appendix B by the auxiliary function method. These update rules guarantee the decrease of error and eventually converge to a local minima [7]. In our experiments, we have set the maximum amount of iterations to 300 (usually 100 iterations are enough) and we use the convergence rule $E_{i-1} - E_i \leq \kappa \max(1, E_{i-1})$ in order to stop the updates when the reconstruction error (E_i) between the current and previous update is small enough. We have set the $\kappa = 10^{-3}$.

The two proposed algorithms have a run-time complexity of $\mathcal{O}(kn^2)$ for clustering n data points to k clusters, which is less than

Algorithm 1 Kernel-based non-negative spectral clustering for Ncut (KNSC-Ncut).

Input: $\mathbf{X} \in \mathbb{R}^{m \times n}$, $\mathbf{K} \in \mathbb{R}^{n \times n}$, $\mathbf{A} \in \mathbb{R}^{n \times n}$, number of clusters k
Output: clustering matrix $\mathbf{Z} \in \mathbb{R}^{k \times n}$, vector of cluster memberships $c_i = \arg\max_{1 \leq j \leq k} \mathbf{Z}_{ji}$

Initialize two non-negative matrices $\mathbf{Z} \in \mathbb{R}^{k \times n}$ and $\mathbf{F} \in \mathbb{R}^{n \times k}$ with random numbers generated in the range [0,1].

Calculate the degree matrix $\mathbf{D} = \text{diag}(d_1, \dots, d_n)$

$$d_i = \sum_{j=1}^n A_{ij} \quad (34)$$

repeat

$$Z_{ij} \leftarrow Z_{ij} \frac{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{D}^{-1/2} + 2\mu \mathbf{Z})_{ij}}{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{F} \mathbf{Z} + 2\mu \mathbf{Z} \mathbf{Z}^{\mathbf{T}} \mathbf{Z})_{ij}}$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{Z}^{\mathbf{T}})_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{Z} \mathbf{Z}^{\mathbf{T}})_{jl}}$$

until Stopping criterion is reached

Algorithm 2 Kernel-based non-negative spectral clustering for Rcut (KNSC-Rcut).

Input: $\mathbf{X} \in \mathbb{R}^{m \times n}$, $\mathbf{K} \in \mathbb{R}^{n \times n}$, number of clusters k
Output: clustering matrix $\mathbf{H} \in \mathbb{R}^{k \times n}$, vector of cluster memberships $c_i = \arg\max_{1 \leq j \leq k} \mathbf{H}_{ji}$

Initialize two non-negative matrices $\mathbf{H} \in \mathbb{R}^{k \times n}$ and $\mathbf{F} \in \mathbb{R}^{n \times k}$ with random numbers generated in the range [0,1].

repeat

$$H_{ij} \leftarrow H_{ij} \frac{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} + 2\mu \mathbf{H})_{ij}}{(\alpha \mathbf{F}^{\mathbf{T}} \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^{\mathbf{T}} \mathbf{H})_{ij}}$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{H}^{\mathbf{T}})_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^{\mathbf{T}})_{jl}}$$

until Stopping criterion is reached

standard spectral clustering methods $\mathcal{O}(n^3)$ and the same complexity as the state-of-the-art methods [29,32,35]. The main advantage of the kernel-based NMF approach is that it can be easily optimized to achieve higher clustering accuracy for the data drawn from nonlinear manifolds, avoiding the computation of eigenvalues and eigenvectors.

3.3. Graph regularized kernel-based orthogonal NMF

A non-negative matrix factorization that respects the geometric structure of the data in the nonlinear feature space can be constructed by introducing an additional graph regularization term into the objective function (17). Recall that our nonlinear NMF tries to find a set of basis vectors that can be used to best approximate the data $\Phi(\mathbf{X}) = \mathbf{WH}$. Let $\mathbf{h}_j$ denote the j th column of $\mathbf{H}$, $\mathbf{h}_j = [h_{j1}, \dots, h_{jk}]$, then $\mathbf{h}_j$ can be regarded as the new representation of the j th data point with respect to the new basis $\mathbf{W} = \Phi(\mathbf{X})\mathbf{F}$. The graph regularization term can be viewed as a local invariance assumption [41,49,50], which states that if two data points $\Phi(\mathbf{x}_i)$ and $\Phi(\mathbf{x}_j)$ are close to each other in the original geometry of the data distribution, then $\mathbf{h}_j$ and $\mathbf{h}_i$, the low dimensional representations of these two points, are also close to each other. This can be written as

$$\mathcal{R} = \frac{1}{2} \sum_{j,l=1}^n \|\mathbf{h}_j - \mathbf{h}_l\|_F^2 \mathbf{A}_{jl} = \sum_{j=1}^n \mathbf{h}_j \mathbf{h}_j^T \mathbf{D}_{j,j} - \sum_{j,l=1}^n \mathbf{h}_j \mathbf{h}_l^T \mathbf{A}_{j,l} = \text{Tr}(\mathbf{H} \mathbf{L} \mathbf{H}^T) \quad (35)$$

By minimizing the regularization term $\mathcal{R}$ with respect to $\mathbf{H}$, we expect that when $\Phi(\mathbf{x}_i)$ and $\Phi(\mathbf{x}_j)$ are close (i.e. when $\mathbf{A}_{jl}$ is large) the points $\mathbf{h}_j$ and $\mathbf{h}_i$ are also close with respect to the new basis. The objective function for nonlinear orthogonal graph regularized NMF is given as

$$\min_{\mathbf{H}, \mathbf{F} \geq 0} \alpha \|\Phi(\mathbf{X}) - \Phi(\mathbf{X})\mathbf{F}\mathbf{H}\|_F^2 + \lambda \text{Tr}(\mathbf{H} \mathbf{L} \mathbf{H}^T), \text{ s.t. } \mathbf{H} \mathbf{H}^T = \mathbf{I} \quad (36)$$

By adopting the same iterative procedure to alternatively fix one of the matrices $\mathbf{F}$ and $\mathbf{H}$, we solve the minimization problem (36) and obtain the multiplicative update rules

$$H_{ij} \leftarrow H_{ij} \frac{(\alpha \mathbf{F}^T \mathbf{K} + 2\mu \mathbf{H} + \lambda \mathbf{H} \mathbf{A})_{ij}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^T \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ij}} \quad (37)$$

$$F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{H}^T)_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T)_{jl}} \quad (38)$$

where $\mathbf{K}$ is the kernel matrix. There are many choices to define the weight matrix $\mathbf{A}$ of the graph. For example, the scalar product weighting and the cosine similarity are most suitable for processing documents, while for image data the heat kernel is commonly used [4,41,51]. We will use the fully connected affinity graph with the Gauss kernel weighting, as we do not treat different weighting schemes separately.

Correctness. The correctness of the solution is assured by the fact that the solution at convergence will satisfy the KKT conditions for the optimization problem (36). The Lagrangian $\mathcal{L}$ of the optimization problem (36) can be written as

$$\begin{aligned} \mathcal{L} = & \alpha \text{Tr}[\Phi(\mathbf{X})\Phi^T(\mathbf{X})] - 2\alpha \text{Tr}[\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\Phi^T(\mathbf{X})] \\ & + \alpha \text{Tr}[\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\mathbf{H}^T\mathbf{F}^T\Phi^T(\mathbf{X})] \\ & + \mu \|\mathbf{H}\mathbf{H}^T - \mathbf{I}_k\|_F^2 + \lambda \text{Tr}[\mathbf{H}\mathbf{D}\mathbf{H}^T] - \lambda \text{Tr}[\mathbf{H}\mathbf{A}\mathbf{H}^T] \end{aligned} \quad (39)$$

Algorithm 3 Kernel-based orthogonal graph regularized NMF (KOGNMF).

Input: $\mathbf{X} \in \mathbb{R}^{m \times n}$, number of clusters k , $\mathbf{K} \in \mathbb{R}^{n \times n}$, $\mathbf{A} \in \mathbb{R}^{n \times n}$
Output: clustering matrix $\mathbf{H}$, vector of cluster memberships $c_i = \arg\max_{1 \leq j \leq k} H_{ji}$
Initialize two non-negative matrices $\mathbf{H} \in \mathbb{R}^{k \times n}$ and $\mathbf{F} \in \mathbb{R}^{n \times k}$ with random numbers generated in the range $[0,1]$.
Calculate the degree matrix $\mathbf{D} = \text{diag}(d_1, \dots, d_n)$
 $d_i = \sum_{j=1}^k A_{ij}$
repeat
 $H_{ij} \leftarrow H_{ij} \frac{(\alpha \mathbf{F}^T \mathbf{K} + 2\mu \mathbf{H} + \lambda \mathbf{H} \mathbf{A})_{ij}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^T \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ij}}$
 $F_{jl} \leftarrow F_{jl} \frac{(\mathbf{K} \mathbf{H}^T)_{jl}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T)_{jl}}$
until Stopping criterion is reached

We calculate the partial derivatives of (39) with respect to $\mathbf{H}$ and $\mathbf{F}$

$$\begin{aligned} \frac{\partial \mathcal{L}}{\partial \mathbf{H}} = & -2\alpha \mathbf{F}^T \Phi^T(\mathbf{X})\Phi(\mathbf{X}) + 2\alpha \mathbf{F}^T \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{F}\mathbf{H} \\ & + 4\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n}) + 2\lambda \mathbf{H} \mathbf{D} - 2\lambda \mathbf{H} \mathbf{A} \end{aligned} \quad (40)$$

$$\frac{\partial \mathcal{L}}{\partial \mathbf{F}} = -\alpha \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{H}^T + \alpha \Phi^T(\mathbf{X})\Phi(\mathbf{X})\mathbf{F}\mathbf{H}\mathbf{H}^T \quad (41)$$

Substituting the quadratic terms with kernel matrix gives

$$\alpha (\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} - \mathbf{F}^T \mathbf{K}) + 2\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n}) + \lambda \mathbf{H} \mathbf{L} = 0 \quad (42)$$

$$-2\alpha \mathbf{K} \mathbf{H}^T + 2\alpha \mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T = 0 \quad (43)$$

Defining the Lagrange multiplier matrix for constraint $\mathbf{H} \geq 0$ as $\Psi = [\psi_{ij}]$, the KKT condition is $\psi_{ij} H_{ij} = 0$. Similarly, the Lagrange multiplier matrix for constraint $\mathbf{F} \geq 0$ is given by $\Xi = [\xi_{jl}]$ and we obtain

$$[\alpha (\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} - \mathbf{F}^T \mathbf{K}) + 2\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}_{n \times n}) + \lambda \mathbf{H} \mathbf{L}]_{ij} H_{ij} = 0$$

$$[2\alpha \mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T - 2\alpha \mathbf{K} \mathbf{H}^T]_{jl} F_{jl} = 0 \quad (44)$$

We separate positive and negative parts of the gradient and obtain multiplicative update rules (37) and (38). By setting $\lambda = 0$ the update rules in Eqs. (37) and (38) reduce to the update rules of the KONMF. We summarize the graph regularized kernel-based orthogonal NMF in the Algorithm 3.

The proposed algorithm has two additional matrix multiplications $\mathbf{H} \mathbf{A}$ and $\mathbf{H} \mathbf{D}$ with complexity of $\mathcal{O}(kn^2)$. Therefore, the total run-time complexity is unchanged and equal to $\mathcal{O}(kn^2)$ for clustering n data points to k clusters. The convergence proof for the multiplicative updates (37) and (38) can be found in the Appendix B.

4. Experiments

In this section we carry out extensive experiments on synthetic and real world data sets to illustrate the effectiveness of the three proposed algorithms: KNSC-Ncut, KNSC-Rcut and KOGNMF. We compare nine recently proposed non-negative spectral clustering algorithms [29,31,32] and traditional Ncut and Rcut spectral clustering methods [19,26]. Our experimental setting is similar to [29,32]. For the purpose of reproducibility we provide the code and data sets (see supplementary files).

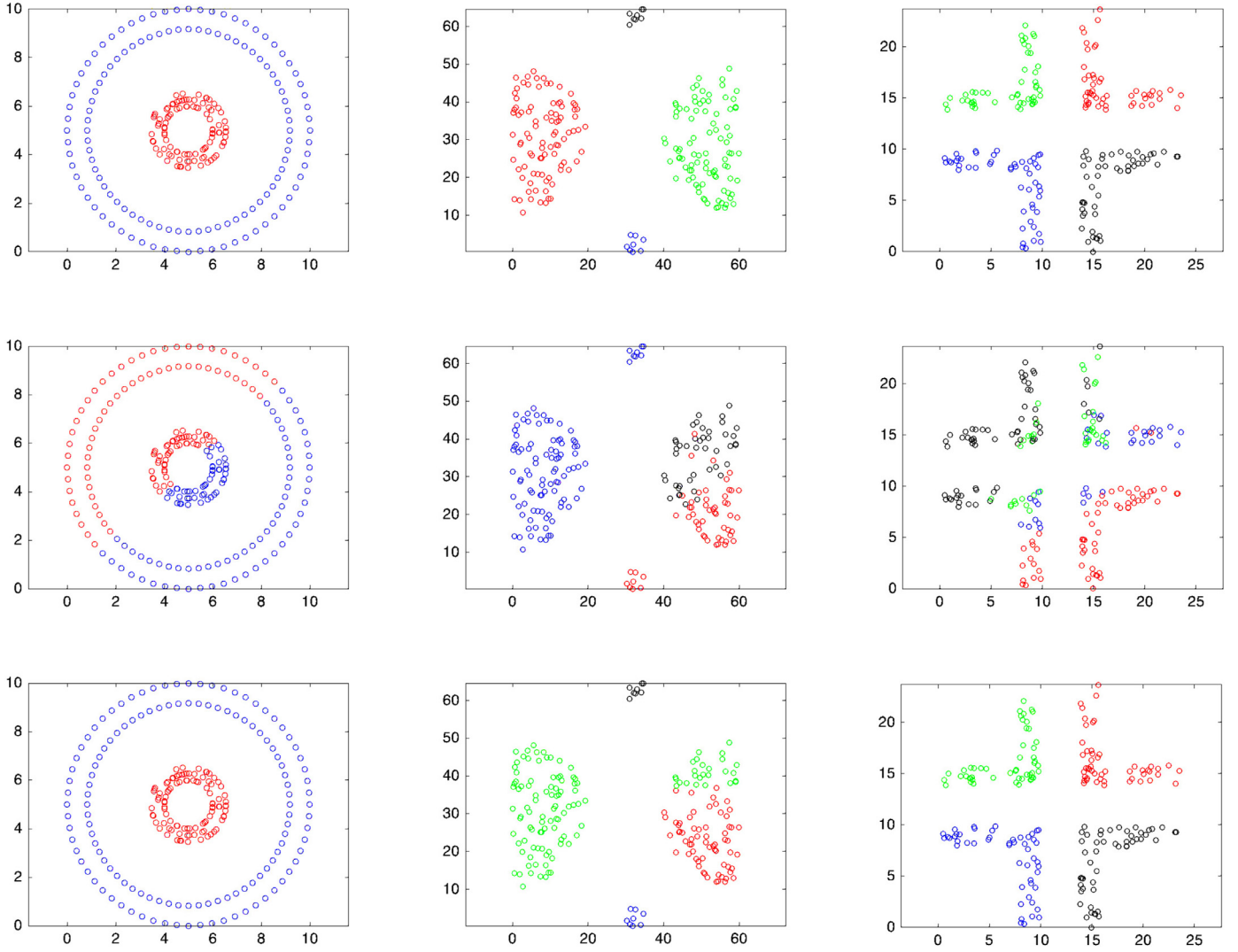


Fig. 2. The optimized clustering results of the KNSC-Ncut algorithm compared with the optimized clustering results of the NSC-Ncut [29]. The two-dimensional data sets with 2 and 4 clusters are plotted in the first row, different clusters are represented with different colors. In the second row we plot the clustering results of the NSC-Ncut. The clustering results of the KNSC-Ncut algorithm are plotted in the third row. The clustering accuracy over 256 independent runs is 0.5, 0.7 and 0.62 for NSC-Ncut, and 0.90, 0.85 and 0.82 for the KNSC-Ncut, for the three data sets respectively. The KNSC-Ncut is able to separate the nonlinear data set composed of two rings of points with high clustering accuracy.

Table 2
Features of the UCI and AT&T data sets.

Datasets	Samples	Dimension	Clusters
Soybean	47	35	4
Zoo	101	16	7
AT&T	400	10,304	40
Glass	214	9	6
Dermatology	366	33	6
Vehicle	846	18	4

4.1. Data sets and the evaluation metric

We have used the same data sets as in [29,30,32]: Five UCI [52] data sets and AT&T face database [53]. The UCI datasets are Soybean, Zoo, Glass, Dermatology and Vehicle. The AT&T face database consists of gray scale face images of 40 persons. Each person has 10 facial images under different light and illumination conditions and the images from the same person belong to the same cluster. The important statistics of these data sets are summarized

in the Table 2, including the number of samples, dimensions and the number of clusters.

The clustering accuracy is evaluated by the common clustering accuracy measure [29,31,32], which computes the percentage of data points that are correctly clustered with respect to the external ground truth labels. For each data point $\mathbf{x}_i$ its label is denoted with c_i and the ground truth cluster index with g_i . In order to calculate the optimal assignment of labels to cluster indices $f(c_i)$, the Hungarian bipartite matching algorithm [52] is used, with the complexity $O(k^3)$ for k clusters. The clustering accuracy can be expressed as:

$$ACC = \frac{\sum_{i=1}^n \delta(g_i, f(c_i))}{n}, \quad (45)$$

where n denotes the total number of data points and the δ function is defined as

$$\delta(g_i, c_i) = \begin{cases} 1 & : g_i = f(c_i), \\ 0 & : g_i \neq f(c_i). \end{cases}$$

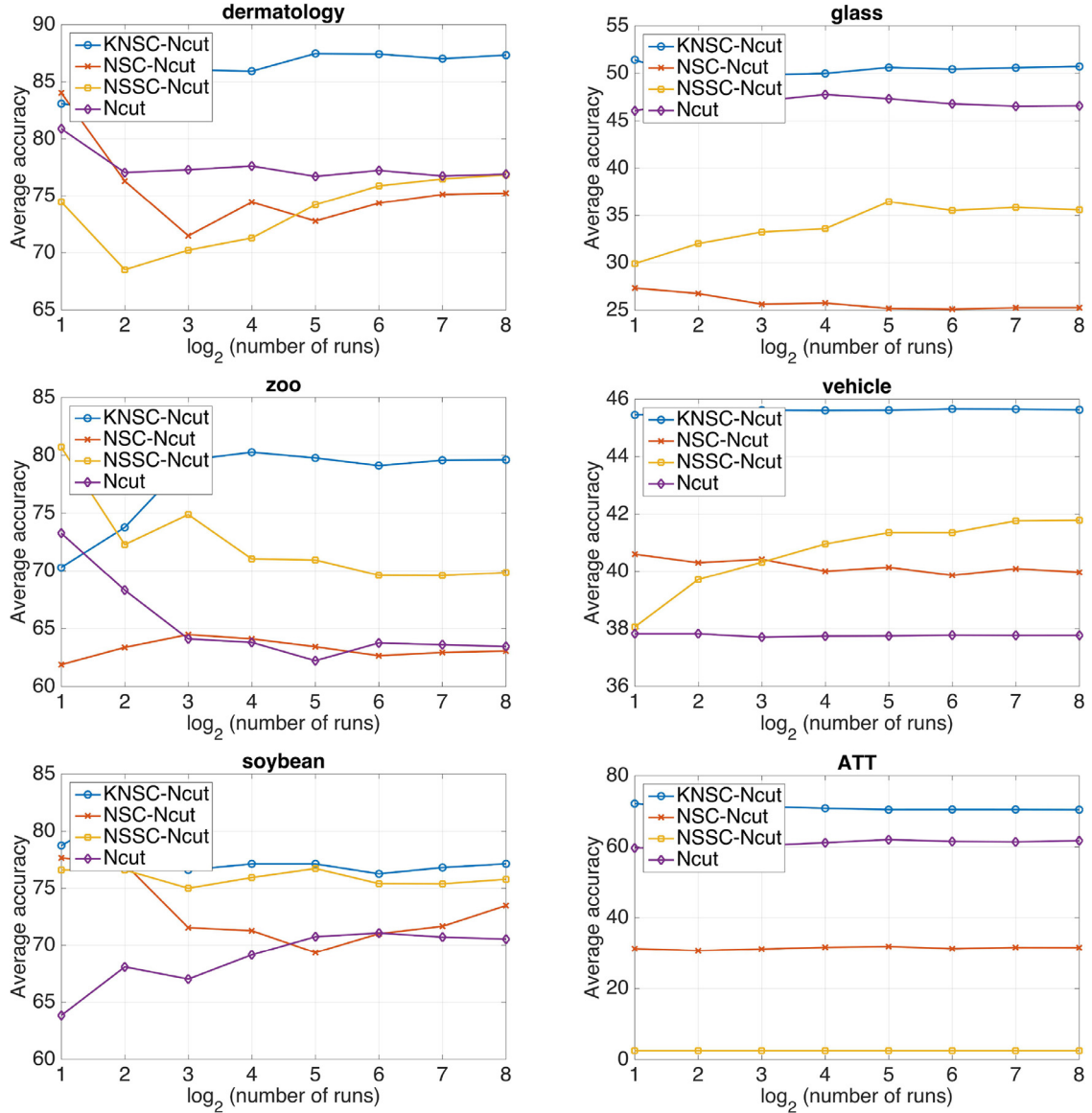


Fig. 3. The average clustering accuracy of KNSC-Ncut algorithm compared with Ncut, NSC-Ncut and NSSC-Ncut algorithms on five UCI [52] data sets and AT&T face database [53]. The average clustering accuracy is plotted for the independent number of runs $2^i = \{2, 4, \dots, 256\}$. The clustering accuracy of KNSC-Ncut is higher on the majority of data sets. The clustering accuracy for the AT&T face database is considerably improved when compared with the state-of-the-art non-negative spectral clustering methods.

4.2. Compared algorithms

We compare our methods to nine recently proposed non-negative spectral clustering approaches and traditional spectral clustering Ncut and Rcut methods:

- Normalized cut (Ncut) and ratio cut (Rcut) spectral clustering. Ncut spectral clustering exists in different normalizations [19,28]. Our implementation is according to Ncut from [19], where eigenvectors of normalized Laplacian matrix $\mathbf{Z}$ are normalized such that the L_2 norm of each row equals 1.
- Non-negative spectral clustering methods NSC-Ncut, NSC-Rcut, and non-negative sparse spectral clustering methods NSSC-Ncut and NSSC-Rcut from [29].
- Global discriminative-based nonnegative spectral clustering methods [32] GDBNSC-Ncut and GDBNSC-Rcut.
- Symmetric NMF for spectral clustering [31] (NLE). This is the symmetric NMF of the pairwise affinity matrix, which is originally implemented as the standard inner product linear kernel matrix $\mathbf{A} = \mathbf{X}^T \mathbf{X}$.

4.3. Clustering results

We perform $n = 256$ independent runs with random initializations for each of the proposed methods KNSC-Ncut, KNSC-Rcut and KOGNMF. In each run, we randomly initialize matrices $(\mathbf{H}, \mathbf{Z}, \mathbf{F})$ and then iterate multiplicative update rules to achieve convergence and obtain cluster indicator matrix. In all experiments we have used 300 iterations and the convergence occurred after approximately 100 iterations. The cluster memberships for each data point i is obtained by taking the index of the maximal value of i th column in the orthogonal clustering matrix $\mathbf{H}$ (or $\mathbf{Z}$). For the Rcut and Ncut, the first k eigenvectors are computed once and then 256 runs of k -means are performed.

In Fig. 2 we plot the clustering performance of the NSC-Ncut and KNSC-Ncut on two-dimensional synthetic examples. The synthetic example demonstrates the ability of KNSC-Ncut to separate the nonlinear clusters with high clustering accuracy. In Figs. 3–5 we plot the average clustering accuracy over 256 runs on the six data sets. The average clustering accuracy is reported for independent number of runs 2^i , where $i = 1, 2, \dots, 8$. The average

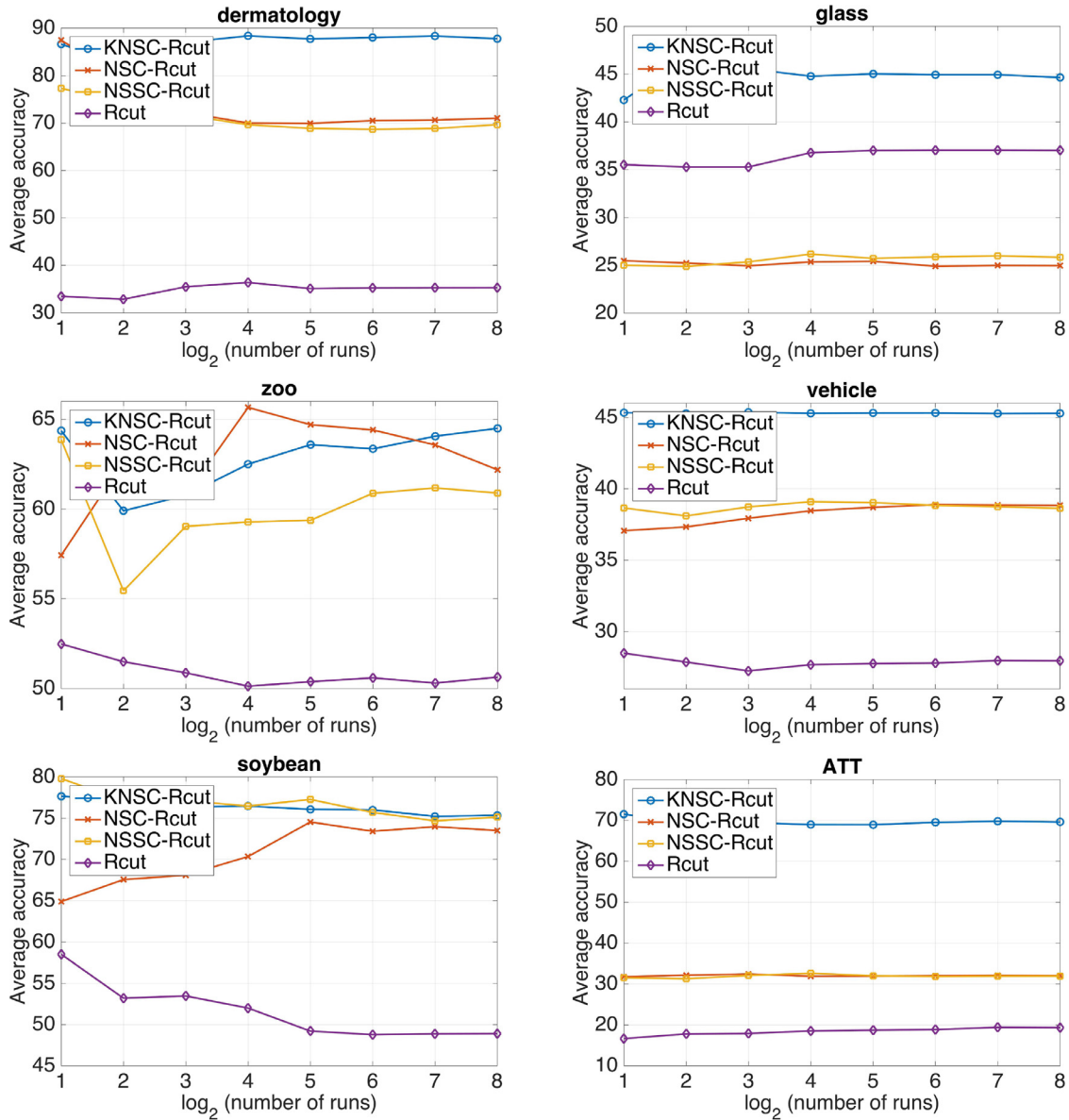


Fig. 4. The average clustering accuracy of KNSC-Rcut algorithm compared with Rcut, NSC-Rcut and NSSC-Rcut algorithms on five UCI [52] data sets and AT&T face database [53]. The average clustering accuracy is plotted for the independent number of runs $2^i = \{2, 4, \dots, 256\}$. The KNSC-Rcut algorithm outperforms NSC algorithms on the majority of data sets.

clustering accuracy for the Ncut group of algorithms is plotted in the Fig. 3. In the Fig. 4 the average clustering accuracy is plotted for the Rcut group.¹ The average clustering accuracy of KOGNMF is shown in Fig. 5. We summarize the average clustering accuracy results for the Ncut and the Rcut group of algorithms in Table 3. On data sets Dermatology, Glass, Zoo and AT&T, the KNSC-Ncut clustering accuracy is improved and KNSC-Ncut outperforms Ncut, NSC-Ncut, NSSC-Ncut and GDBNSC-Ncut. On the high dimensional AT&T face database the clustering accuracy of the KNSC-Ncut algorithm shows considerable improvement. On the Soybean and Vehicle data sets the KNSC-Ncut is comparable with the GDBNSC-Ncut. Similarly, on Dermatology, Glass, Zoo, Vehicle and AT&T data set, KNSC-Rcut outperforms Rcut, NSC-Rcut, NSSC-Rcut and GDBNSC-Rcut. In Fig. 5 we plot the average clustering accuracy for the

KOGNMF algorithm. The KOGNMF considerably outperforms all algorithms on every data set (Table 3).

4.4. The parameter selection

The kernel-based orthogonal NMF multiplicative rules have in total four parameters: α , μ and λ and the Gaussian kernel width σ . The three parameters α , μ and λ are a trade-off parameters which balance the reconstruction error, orthogonality regularization and the graph regularization, respectively. In all the experiments and data sets we have fixed the three trade-off parameters to the same constant values $\alpha = 10$, $\mu = 100$ and $\lambda = 10$. Furthermore, the three trade-off parameters can be reduced to two, as the NMF objective functions given in the Eqs. (17) and (36) can be divided by α . By fixing the trade-off parameters throughout all of the experiments we effectively need to optimize only one parameter, which is the kernel width. For the trade-off parameters we perform sensitivity analysis to demonstrate that the constant

¹ The results in Table 3 for the GDBNSC method are reported from original work [32], however in Figs. 3–5 the results for this method were omitted due to the numerical instabilities in reproduction of this method with reported parameters [32].

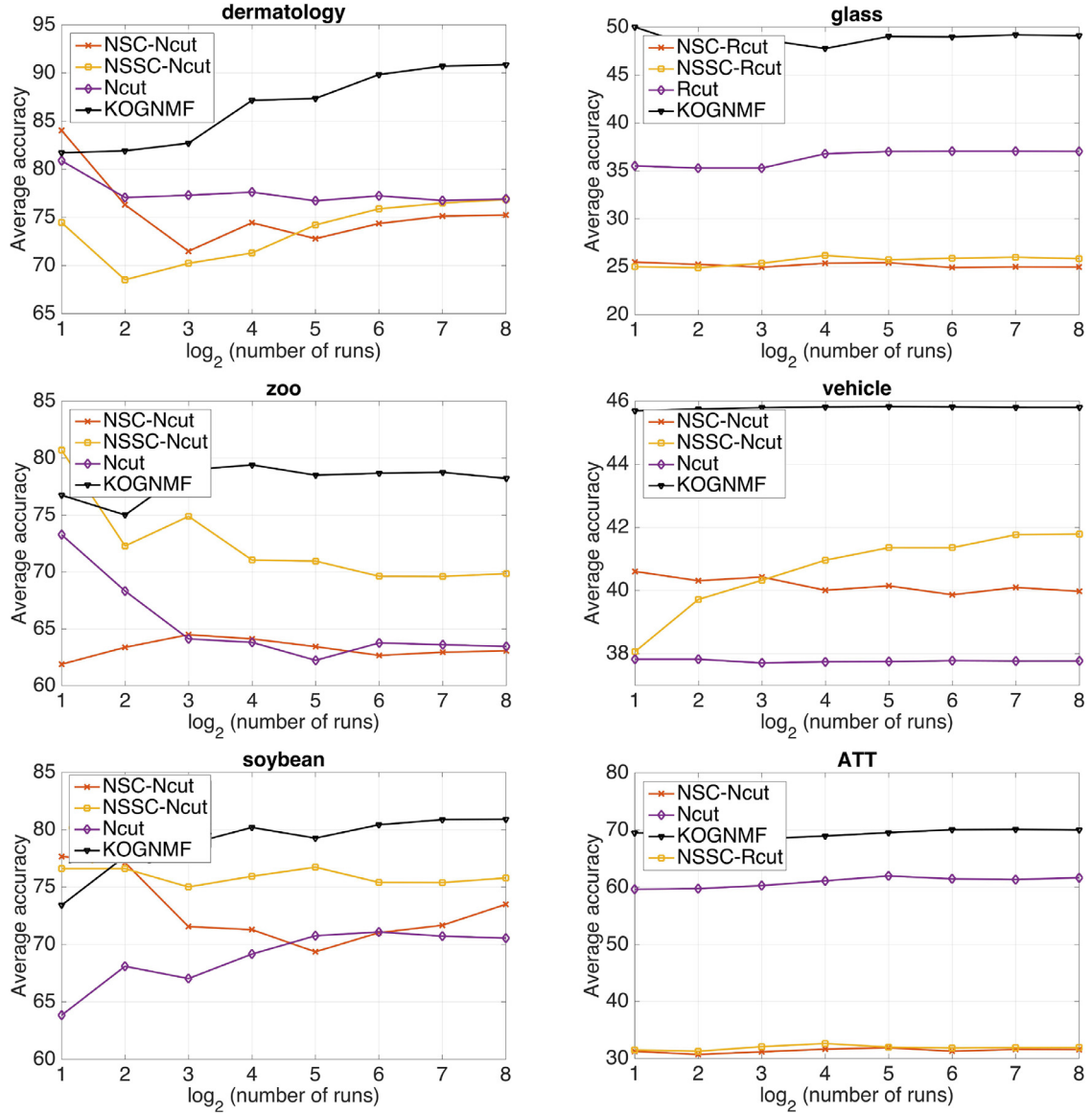


Fig. 5. The average clustering accuracy of KOGNMF algorithm on 5 UCI [52] data sets and AT&T face database. The average clustering accuracy is plotted for the independent number of runs $2^i = \{2, 4, \dots, 256\}$. The KOGNMF algorithm outperforms all non-negative spectral clustering methods on every data set, including the difficult AT&T face database [53].

values of the trade-off parameters can be chosen in a wide range of values (a few orders of magnitude), as shown in Figs. 6 and 7.

In the experiments we use the Gaussian kernel defined as $\mathbf{K}(\mathbf{x}_i, \mathbf{x}_j) = \exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2 / \sigma^2)$, where σ is the kernel width. For the graph regularization term we use a fully connected affinity graph with the Gaussian kernel weighting on the edges. To choose the parameter σ we perform a simple grid search for the 40 values of σ in the range of $[0.1, 4]$ with the step size $\Delta\sigma = 0.1$ for data sets Dermatology, Glass, Soybean and Zoo. For the AT&T face database we perform the grid search in the range $\sigma = [1000, 10000]$ with the step size $\Delta\sigma = 250$. For the Vehicle data set we perform the grid search in the range $\sigma = [10, 100]$ with the step size $\Delta\sigma = 10$. At the boundary values of the σ intervals the clustering accuracy saturates. For small values of σ the similarity of the data points with large distance $\|\mathbf{x}_i - \mathbf{x}_j\|$ goes to zero as $\exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2 / \sigma^2) \rightarrow 0$ when $\|\mathbf{x}_i - \mathbf{x}_j\|^2 / \sigma^2$ is large. Therefore, for small distances, the affinity graph captures the local Euclidean distance and gives a good representation of the manifold structure.

For KNSC-Ncut algorithm we used the same grid search to obtain a degree matrix $\mathbf{D}^{-1/2}$.

For each data set, we measure the average clustering accuracy out of 256 independent runs. We perform a hold-out validation for the parameter σ , as shown in the Table 4. The hold-out validation consists of randomly splitting each data set into two equally sized parts with the equally distributed cluster membership. The grid search optimization is performed on the first half of the data set, while the second half is used as a hold-out validation where optimized parameters are used. The results of the hold-out validation show robust average clustering accuracy for all three algorithms on all six data sets.

The sensitivity analysis of the algorithms is performed for the three trade-off parameters α , μ and λ , plotted in Figs. 6 and 7. The ratio of the parameters μ and α is fixed to a constant value in all experiments. The near-orthogonality of the clustering indicator matrix $\mathbf{H}(\mathbf{Z})$ is preserved during the multiplicative updates, as shown in Figs. 6 and 7. The near-orthogonality of columns is important for data clustering interpretation. An orthogonal clustering

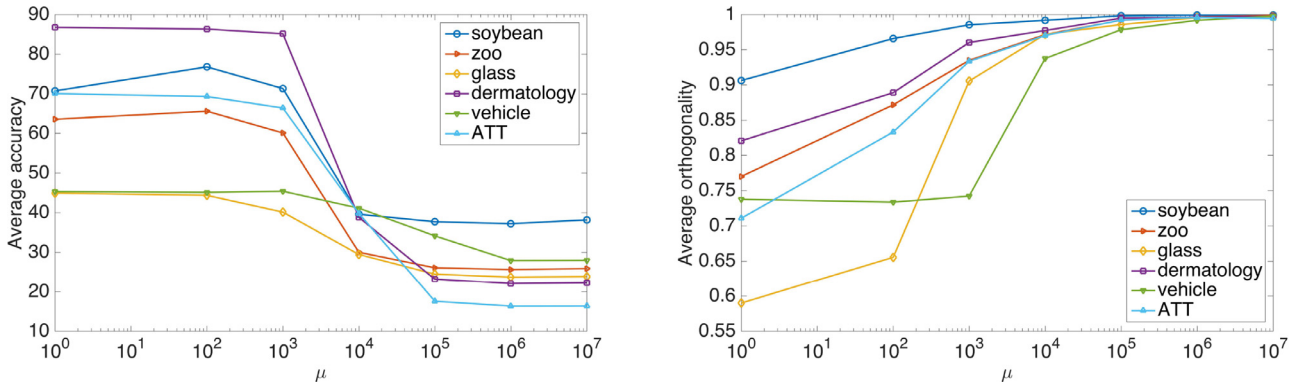


Fig. 6. Left: The average orthogonality of the clustering matrix $\mathbf{H}$ (KNSC-Rcut) over the 256 runs, plotted for fixed reconstruction error parameter $\alpha = 10$ and for a wide range of values of the orthogonality parameter μ on all six data sets. Right: The average clustering accuracy of KNSC-Rcut for fixed $\alpha = 10$ plotted for different values of the parameter μ . The average orthogonality of the clustering matrix $\mathbf{H}$ increases up to 1 if the parameter μ is increased. The average clustering accuracy is robust for all six data sets for a wide range of the trade-off parameter μ .

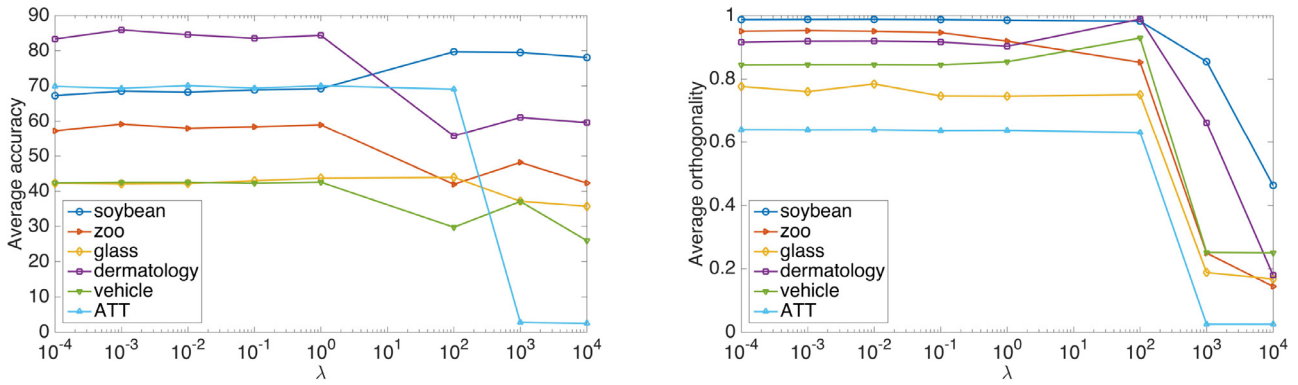


Fig. 7. Left: The average orthogonality of the clustering matrix $\mathbf{H}$ (KNSC-Rcut) over the 256 runs, plotted for fixed reconstruction error parameter $\alpha = 10$ and orthogonality regularization parameter $\mu = 100$ for different values of the graph regularization parameter λ on all six data sets. Right: The average clustering accuracy of KNSC-Rcut for fixed parameters $\alpha = 10$ and $\mu = 100$ plotted for different values of the parameter λ . The average clustering accuracy is robust for all six data sets for a wide range of the trade-off parameter λ .

Table 3

The average clustering accuracy on 5 UCI and AT&T data sets. The average clustering accuracy of KNSC-Ncut, KNSC-Rcut and KOGNMF compared with 9 recently proposed NMF-based NSC methods on the 5 UCI [52] data sets and the AT&T face database [53]. KNSC-Rcut performs considerably better on 4 data sets, and has a comparable accuracy on two data sets. KNSC-Ncut algorithm outperforms on 5 data sets, and has a comparable clustering accuracy on one data set. KOGNMF algorithm has considerably better accuracy on 4 data sets, including the difficult AT&T face images database, and is comparable on two data sets. All three algorithms have considerably higher clustering accuracy on the difficult AT&T face database.

	Dermatology	Glass	Soybean	Zoo	Vehicle	AT&T
Ncut	0.75	0.46	0.70	0.63	0.37	0.62
NSC-Ncut	0.71	0.25	0.71	0.61	0.39	0.35
NSSC-Ncut	0.71	0.34	0.71	0.66	0.41	0.02
GDBNSC-Ncut	0.82	0.41	0.79	0.65	0.46	0.38
KNSC-Ncut	0.87	0.50	0.78	0.80	0.45	0.70
Rcut	0.47	0.41	0.63	0.60	0.33	0.31
NSC-Rcut	0.66	0.25	0.69	0.61	0.38	0.35
NLE	0.34	0.25	0.47	0.49	0.28	0.20
NSSC-Rcut	0.67	0.26	0.69	0.61	0.38	0.35
GDBNSC-Rcut	0.73	0.36	0.80	0.64	0.388	0.36
KNSC-Rcut	0.87	0.45	0.75	0.65	0.45	0.69
KOGNMF	0.91	0.48	0.80	0.78	0.45	0.70

matrix has an interpretation that each row of $\mathbf{H}$ ($\mathbf{Z}$) can have only one nonzero element, which implies that each data object belongs only to 1 cluster. We plot the average orthogonality over 256 runs of the clustering matrix $\mathbf{H}$ (KNSC-Rcut) for a wide range of values of the parameter μ and fixed α . The average orthogonality per run is defined as $\sum_{i=1}^k (\mathbf{H}\mathbf{H}^T)_{i,i} / \sum_{i \neq j} (\mathbf{H}\mathbf{H}^T)_{i,j}$. For a wide range

Table 4

The average clustering accuracy on the hold-out validation set. The hold-out validation consists of randomly splitting each data set into two equally sized parts with the equally distributed cluster membership. The grid search optimization is performed on the first half of the data set, while the second half is used as a hold-out validation where optimized parameters are used. For each data set, we measure the average score over 256 independent runs on the hold-out data. We denote with bold our results that outperform the optimized clustering accuracy scores of the state-of-the-art NSC methods without the hold-out validation. The KNSC-Ncut and KNSC-Rcut algorithms have higher average clustering accuracy on the majority of data sets, while KOGNMF algorithm outperforms on all six data sets.

Datasets	Dermatology	Glass	Soybean	Zoo	Vehicle	AT&T
NLE	0.37	0.38	0.55	0.45	0.33	0.26
KNSC-Ncut	0.87	0.47	0.73	0.77	0.47	0.70
KNSC-Rcut	0.85	0.47	0.76	0.67	0.48	0.73
KOGNMF	0.89	0.49	0.76	0.78	0.48	0.73

of values of the ratio μ/α the orthogonality is preserved during the updates. In Fig. 6 we plot the corresponding average clustering accuracy for KNSC-Rcut. When μ becomes a few order of magnitude larger compared to the reconstruction error term, the objective function effectively becomes the optimization of the orthogonality term. At that point the reconstruction error term loses its significance and the average clustering accuracy starts to drop. In Fig. 6 we plot the clustering accuracy in a wide range of values of the parameter μ . The graph regularization λ is fixed to a constant value $\lambda = \alpha$ for simplicity. The average orthogonality is plotted for different values of λ and μ parameters in Figs. 6 and 7. The clus-

tering accuracy is robust for a wide range of λ , $\lambda = [10^{-4} - 10^2]$, and μ , $\mu = [10^0 - 10^7]$ throughout the experiments on all six data sets.

5. Conclusion

In this paper we study subspace clustering from nonlinear orthogonal non-negative matrix factorization perspective. We have constructed a nonlinear orthogonal NMF algorithm and derived three novel clustering algorithms. We have formally shown that the Rcut spectral clustering is equivalent to the nonlinear orthonormal NMF. The equivalence with the Ncut spectral clustering is obtained by introducing an additional scaling matrix into the nonlinear factorization. Based on this equivalence, we have proposed two kernel-based non-negative spectral clustering methods, KNSC-Ncut and KNSC-Rcut. By incorporating the graph regularization term into the nonlinear NMF framework we have formulated a kernel-based graph-regularized orthogonal non-negative matrix factorization (KOGNMF). To solve the subspace clustering we have derived general kernel-based orthogonal multiplicative updates with complexity $\mathcal{O}(kn^2)$. The monotonic convergence of all three algorithms is proven using an auxiliary function analogous to that used for proving convergence of the Expectation-Maximization algorithm. Experimental results show the effectiveness of our methods compared to state-of-the-art recently proposed NMF-based clustering methods.

Acknowledgments

The authors would like to thank to Mario Lucic and Maria Brbic for proofreading the article. The work of DT is funded by the by the [Croatian Science Foundation](#) with the project No. [I-1701-2014](#) "Machine learning algorithms for insightful analysis of complex data structures". The work of IK is funded by Croatian science foundation with the project [IP-2016-06-5235](#) "Structured Decompositions of Empirical Data for Computationally-Assisted Diagnoses of Disease" and was in part supported by [European Regional Development Fund](#) under the grant [KK.01.1.1.01.0009](#) (DATACROSS). The work of NAF is funded by the [EU Horizon 2020](#) SoBigData project under grant agreement No. [654024](#).

Appendix A

Proof of Theorem 1. The factorization $\Phi(\mathbf{X})\mathbf{D}^{-1/2} = \mathbf{W}\mathbf{Z}$ can be solved by the following optimization problem

$$\min_{\mathbf{Z}, \mathbf{W}} \|\Phi(\mathbf{X})\mathbf{D}^{-1/2} - \mathbf{W}\mathbf{Z}\|_F^2 \text{ s.t. } \mathbf{Z}\mathbf{Z}^T = \mathbf{I}, \quad (46)$$

where $\mathbf{Z}\mathbf{Z}^T = \mathbf{I}$ is the orthogonality constraint which can be included in the optimization implicitly or explicitly via Lagrange multipliers. Then objective function can be reformulated as $J(\mathbf{Z}, \mathbf{W})$

$$\frac{1}{2} \text{Tr}((\Phi(\mathbf{X})\mathbf{D}^{-1/2} - \mathbf{W}\mathbf{Z})^T (\Phi(\mathbf{X})\mathbf{D}^{-1/2} - \mathbf{W}\mathbf{Z})) \quad (47)$$

$$= \frac{1}{2} \text{Tr}((\mathbf{D}^{-1/2}\Phi^T(\mathbf{X}) - \mathbf{Z}^T\mathbf{W}^T)(\Phi(\mathbf{X})\mathbf{D}^{-1/2} - \mathbf{W}\mathbf{Z})) \quad (48)$$

$$= \frac{1}{2} \text{Tr}(\Phi(\mathbf{X})\mathbf{D}^{-1}\Phi(\mathbf{X})^T - 2\mathbf{W}\mathbf{Z}\mathbf{D}^{-1/2}\Phi(\mathbf{X})^T + \mathbf{W}\mathbf{W}^T). \quad (49)$$

The constraint $\mathbf{Z}\mathbf{Z}^T = \mathbf{I}$ is used in the last equality. Calculating the partial derivative of $J(\mathbf{Z}, \mathbf{W})$ with respect to $\mathbf{W}$ and letting it be equal to 0, it follows

$$\frac{\partial J(\mathbf{Z}, \mathbf{W})}{\partial \mathbf{W}} = -\Phi(\mathbf{X})\mathbf{D}^{-1/2}\mathbf{Z}^T + \mathbf{W} = 0. \quad (50)$$

From here, we have

$$\mathbf{W} = \Phi(\mathbf{X})\mathbf{D}^{-1/2}\mathbf{Z}^T \quad (51)$$

Substituting (51) back into (49), we obtain $J(\mathbf{Z}, \mathbf{W}) =$

$$\frac{1}{2} \text{Tr}(\Phi(\mathbf{X})\mathbf{D}^{-1}\Phi(\mathbf{X})^T - 2\Phi(\mathbf{X})\mathbf{D}^{-1/2}\mathbf{Z}^T\mathbf{Z}\mathbf{D}^{-1/2}\Phi(\mathbf{X})^T). \quad (52)$$

Since $\Phi(\mathbf{X})\mathbf{D}^{-1}\Phi(\mathbf{X})^T$ is not dependent on $\mathbf{Z}$ and $\mathbf{W}$, the minimization problem is equivalent to

$$\max_{\mathbf{Z}, \mathbf{W}} \text{Tr}(\mathbf{Z}\mathbf{D}^{-1/2}\Phi(\mathbf{X})^T\Phi(\mathbf{X})\mathbf{D}^{-1/2}\mathbf{Z}^T) \text{ s.t. } \mathbf{Z}\mathbf{Z}^T = \mathbf{I}. \quad (53)$$

For $\mathbf{A} = \Phi^T(\mathbf{X})\Phi(\mathbf{X})$ the objective function (53) is

$$\max_{\mathbf{Z}} \text{Tr}(\mathbf{Z}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Z}^T) \text{ s.t. } \mathbf{Z}\mathbf{Z}^T = \mathbf{I}. \quad (54)$$

Note, that the objective function for Ncut spectral clustering

$$\min_{\mathbf{Z}} \text{Tr}(\mathbf{Z}\mathbf{L}_{\text{sym}}\mathbf{Z}^T) \text{ s.t. } \mathbf{Z}\mathbf{Z}^T = \mathbf{I}. \quad (55)$$

can easily be transformed to (53).

$$\min_{\mathbf{Z}, \mathbf{Z}\mathbf{Z}^T = \mathbf{I}} \text{Tr}(\mathbf{Z}\mathbf{D}^{-1/2}(\mathbf{D} - \mathbf{A})\mathbf{D}^{-1/2}\mathbf{Z}^T) = \quad (56)$$

$$\min_{\mathbf{Z}, \mathbf{Z}\mathbf{Z}^T = \mathbf{I}} \text{Tr}(\mathbf{Z}\mathbf{D}^{-1/2}\mathbf{D}\mathbf{D}^{-1/2}\mathbf{Z}^T - \mathbf{Z}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Z}^T) = \quad (57)$$

and since the term $\mathbf{Z}\mathbf{D}^{-1/2}\mathbf{D}\mathbf{D}^{-1/2}\mathbf{Z}^T = \mathbf{I}$ due to the orthogonality $\mathbf{Z}\mathbf{Z}^T = \mathbf{I}$ this is equal to maximization of the second term.

$$\max_{\mathbf{Z}, \mathbf{Z}\mathbf{Z}^T = \mathbf{I}} \text{Tr}(\mathbf{Z}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Z}^T). \quad (58)$$

which concludes the proof.

Proof of Theorem 2. For the Rcut spectral clustering we solve the factorization $\Phi(\mathbf{X}) = \mathbf{W}\mathbf{H}$, with constraint $\mathbf{H}\mathbf{H}^T = \mathbf{I}$. The factorization $\Phi(\mathbf{X}) = \mathbf{W}\mathbf{H}$ can be solved by the optimization problem

$$\min_{\mathbf{H}, \mathbf{W}, \mathbf{H}\mathbf{H}^T = \mathbf{I}} \|\Phi(\mathbf{X}) - \mathbf{W}\mathbf{H}\|_F^2, \quad (59)$$

where $\mathbf{H}\mathbf{H}^T = \mathbf{I}$ is the orthogonality constraint which can be included in the optimization implicitly or explicitly. The objective function (59) can be reformulated as

$$J(\mathbf{H}, \mathbf{W}) = \frac{1}{2} \text{Tr}((\Phi(\mathbf{X}) - \mathbf{W}\mathbf{H})^T (\Phi(\mathbf{X}) - \mathbf{W}\mathbf{H})) = \quad (60)$$

$$= \frac{1}{2} \text{Tr}((\Phi^T(\mathbf{X}) - \mathbf{H}^T\mathbf{W}^T)(\Phi(\mathbf{X}) - \mathbf{W}\mathbf{H})) = \quad (61)$$

$$= \frac{1}{2} \text{Tr}(\Phi(\mathbf{X})^T\Phi(\mathbf{X}) - 2\Phi(\mathbf{X})^T\mathbf{W}\mathbf{H} + \mathbf{W}^T\mathbf{W}). \quad (62)$$

The constraint $\mathbf{H}\mathbf{H}^T = \mathbf{I}$ is used in the last equality. Calculating the partial derivative of $J(\mathbf{H}, \mathbf{W})$ with respect to $\mathbf{W}$ and letting it be equal to 0, it follows

$$\frac{\partial J(\mathbf{H}, \mathbf{W})}{\partial \mathbf{W}} = -\Phi(\mathbf{X})\mathbf{H}^T + \mathbf{W} = 0. \quad (63)$$

From here, we have

$$\mathbf{W} = \Phi(\mathbf{X})\mathbf{H}^T. \quad (64)$$

Substituting (64) back into (62), we obtain

$$J(\mathbf{H}) = \frac{1}{2} \text{Tr}(\Phi(\mathbf{X})^T\Phi(\mathbf{X}) - \Phi(\mathbf{X})^T\Phi(\mathbf{X})\mathbf{H}^T\mathbf{H}). \quad (65)$$

Since the first term is constant, not dependent on $\mathbf{H}$ and $\mathbf{W}$, the minimization problem is equivalent to

$$\max_{\mathbf{H}, \mathbf{W}, \mathbf{H}\mathbf{H}^T = \mathbf{I}} \text{Tr}(\mathbf{H}\Phi(\mathbf{X})^T\Phi(\mathbf{X})\mathbf{H}^T). \quad (66)$$

For $\mathbf{A} = \Phi^T(\mathbf{X})\Phi(\mathbf{X})$ the objective function (66) is the same as objective function (58) for the relaxed Rcut spectral clustering. To see

why, we start from the objective function of Rcut and come to the relaxed Rcut optimization function [29]:

$$\min_{\mathbf{H}, \mathbf{H}\mathbf{H}^T=\mathbf{I}} \text{Tr}(\mathbf{H}\mathbf{L}\mathbf{H}^T) = \min_{\mathbf{H}, \mathbf{H}\mathbf{H}^T=\mathbf{I}} \text{Tr}(\mathbf{H}\mathbf{D}\mathbf{H}^T - \mathbf{H}\mathbf{A}\mathbf{H}^T). \quad (67)$$

Now, the substitution is made $\mathbf{Q} = \mathbf{H}\mathbf{D}^{1/2}$ which implies $\mathbf{H} = \mathbf{Q}\mathbf{D}^{-1/2}$, $\mathbf{H}\mathbf{H}^T = \mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T$ and the objective function can be written as:

$$\begin{aligned} \min_{\mathbf{Q}, \mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T=\mathbf{I}} \text{Tr}(\mathbf{Q}\mathbf{D}^{-1/2}\mathbf{D}\mathbf{D}^{-1/2}\mathbf{Q}^T - \mathbf{Q}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Q}^T) \\ = \min_{\mathbf{Q}, \mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T=\mathbf{I}} \text{Tr}(\mathbf{Q}\mathbf{Q}^T - \mathbf{Q}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Q}^T). \end{aligned} \quad (68)$$

The expression (68) is equivalent to

$$\max_{\mathbf{Q}, \mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T=\mathbf{I}} \text{Tr}(\mathbf{Q}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Q}^T) \text{ s.t. } \mathbf{Q}\mathbf{Q}^T = \mathbf{I} \quad (69)$$

Next, we release the orthonormality constraint $\mathbf{Q}\mathbf{Q}^T = \mathbf{I}$. The relaxation is justified by the fact that the rows of $\mathbf{Q}$ are orthogonal to each other since $\mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T = \mathbf{I}$.

$$\max_{\mathbf{Q}, \mathbf{Q}\mathbf{D}^{-1}\mathbf{Q}^T=\mathbf{I}} \text{Tr}(\mathbf{Q}\mathbf{D}^{-1/2}\mathbf{A}\mathbf{D}^{-1/2}\mathbf{Q}^T) \quad (70)$$

and by substitution $\mathbf{Q} = \mathbf{H}\mathbf{D}^{1/2}$ this becomes:

$$\max_{\mathbf{H}, \mathbf{H}\mathbf{H}^T=\mathbf{I}} \text{Tr}(\mathbf{H}\mathbf{A}\mathbf{H}) \quad (71)$$

which is equal to objective function of (66), which concludes the proof.

Appendix B

Proof 3. The convergence analysis of the proposed algorithms.

We now show the algorithm KOGNMF converges to a feasible solution. We use the auxiliary function approach, following [7,32]. The convergence of KNSC-Ncut and KNSC-Rcut can be proven in a similar way.

The objective function of KOGNMF (36) is non-increasing under the alternative iterative updating rules in (37) and (38).

Definition 1. $A(h, h')$ is an auxiliary function for $B(h)$ when the following conditions are satisfied:

$$A(h, h') \geq B(h), \quad A(h, h) = B(h). \quad (72)$$

The auxiliary function is useful because of the following lemma:

Lemma 1. If A is an auxiliary function of B , then B is non-increasing under the updating formula

$$h^{(t+1)} = \arg \min_h A(h, h^{(t)}) \quad (73)$$

the function B is non-increasing.

Proof. $B(h^{(t+1)}) \leq A(h^{(t+1)}, h^{(t)}) \leq A(h^{(t)}, h^{(t)}) = B(h^{(t)})$. $\square$

We now rewrite the objective function $\mathcal{L}$ of KOGNMF in Eq. (36) as follows

$$\begin{aligned} \mathcal{L} = \alpha \|\Phi(\mathbf{X}) - \Phi(\mathbf{X})\mathbf{F}\mathbf{H}\|_F^2 + \lambda \text{Tr}(\mathbf{H}\mathbf{L}\mathbf{H}^T) + \mu \|\mathbf{H}\mathbf{H}^T - \mathbf{I}_k\|_F^2 \\ = \alpha \sum_{i=1}^D \sum_{j=1}^n \left(\Phi(x)_{ij} - \sum_{l=1}^k w_{il} h_{lj} \right)^2 + \lambda \sum_{m=1}^k \sum_{j=1}^n \sum_{l=1}^n h_{mj} L_{jl} h_{lm} \\ + \mu \sum_{i=1}^D \sum_{j=1}^n \left(\sum_{l=1}^n h_{il} h_{lj} - \delta_{ij} \right)^2 \end{aligned} \quad (74)$$

Considering any element h_{ab} in $\mathbf{H}$, we use B_{ab} to denote the part of $\mathcal{L}$ relevant to h_{ab} . Then it follows

$$\dot{B}_{ab} \equiv \left(\frac{\partial \mathcal{L}}{\partial h_{ab}} \right)_{ab} = (2\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} - 2\alpha \mathbf{F}^T \mathbf{K} + 2\lambda \mathbf{H} \mathbf{L} + 4\mu \mathbf{H}(\mathbf{H}^T \mathbf{H} - \mathbf{I}))_{ab} \quad (75)$$

Since multiplicative update rules are element-wise, we have to show that each B_{ab} is non-increasing under the update step given in Eq. (37).

Lemma 2. Function

$$\begin{aligned} A(h, h_{ab}^{(t)}) = B(h_{ab}^{(t)}) + \dot{B}_{ab}(h_{ab}^{(t)})(h - h_{ab}^{(t)}) \\ + \frac{(2\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + 2\lambda \mathbf{H} \mathbf{D})_{ab}}{h_{ab}^{(t)}} (h - h_{ab}^{(t)})^2. \end{aligned} \quad (76)$$

is an auxiliary function for B_{ab} , when $\mu = 0$.

Proof. By the above equation, we have $A(h, h) = B_{ab}(h)$, so we only need to show that $A(h, h_{ab}^{(t)}) \geq B_{ab}(h)$. To this end, we compare the auxiliary function given in Eq. (76) with the Taylor expansion of $B_{ab}(h)$. $\square$

$$\ddot{B}_{ab} \equiv \left(\frac{\partial^2 \mathcal{L}}{\partial h_{ab}^2} \right)_{ab} = (2\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} + 2\lambda \mathbf{L})_{ab} \quad (77)$$

$$B_{ab}(h) = B_{ab}(h_{ab}^{(t)}) + \dot{B}_{ab}(h - h_{ab}^{(t)}) + [\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} + \lambda \mathbf{L}]_{ab} (h - h_{ab}^{(t)})^2 \quad (78)$$

to find that $A(h, h_{ab}^{(t)}) \geq B_{ab}(h)$ is equivalent to

$$\frac{\alpha (\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H})_{ab} + \lambda (\mathbf{H} \mathbf{D})_{ab}}{h_{ab}^{(t)}} \geq (\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} + \lambda \mathbf{L})_{ab} \quad (79)$$

$$(\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H})_{ab} = \sum_{l=1}^k (\mathbf{F}^T \mathbf{K} \mathbf{F})_{al} h_{lb}^{(t)} \geq (\mathbf{F}^T \mathbf{K} \mathbf{F})_{aa} h_{ab}^{(t)} \quad (80)$$

$$(\mathbf{H} \mathbf{D})_{ab} = \sum_{l=1}^n h_{ab}^{(t)} \mathbf{D}_{lb} \geq h_{ab}^{(t)} \mathbf{D}_{bb} \geq h_{ab}^{(t)} (\mathbf{D} - \mathbf{A})_{bb} \quad (81)$$

In summary, we have the following inequality

$$\frac{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ab}}{h_{ab}^{(t)}} \geq \frac{1}{2} \ddot{B}_{ab} \quad (82)$$

Then the inequality $A(h, h_{ab}^{(t)}) \geq B_{ab}(h)$ is satisfied, and the Lemma is proven.

From Lemma 2, we know that $A(h, h_{ab}^{(t)})$ is an auxiliary function of $B_{ab}(h_{ab})$. We can now demonstrate the convergence of the update rules given in Eq. (37).

$$h^{t+1} = \arg \min_h A(h, h^{(t)}) \quad (83)$$

$$h_{ab}^{t+1} = h_{ab}^{(t)} \frac{(\alpha \mathbf{F}^T \mathbf{K} + \lambda \mathbf{H} \mathbf{A})_{ab}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ab}} \quad (84)$$

So the updating rule for H is as follows:

$$H_{ab} \leftarrow H_{ab} \frac{(\alpha \mathbf{F}^T \mathbf{K} + \lambda \mathbf{H} \mathbf{A})_{ab}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ab}} \quad (85)$$

Similarly, for $\mu > 0$, we use the following auxiliary function $A(h, h_{ab}^t) =$

$$A(h, h_{ab}^t) = B(h_{ab}^{(t)}) + \dot{B}_{ab}(h_{ab}^{(t)})(h - h_{ab}^{(t)}) + \frac{\alpha(\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H})_{ab} + \lambda(\mathbf{H} \mathbf{D})_{ab} + \mu(\mathbf{H} \mathbf{H}^T \mathbf{H})_{ab}}{h_{ab}^t} (h - h_{ab}^t)^2. \quad (86)$$

and by using this:

$$(\mathbf{H} \mathbf{H}^T \mathbf{H})_{ab} = \sum_{l=1}^n h_{al}^t (\mathbf{H}^T \mathbf{H})_{lb} \geq h_{ab}^t (\mathbf{H}^T \mathbf{H})_{bb} \quad (87)$$

we obtain the following inequality

$$\frac{\alpha(\mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H})_{ab} + \lambda(\mathbf{H} \mathbf{D})_{ab} + \mu(\mathbf{H} \mathbf{H}^T \mathbf{H})_{ab}}{h_{ab}^t} \geq (\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} + \mu \mathbf{H}^T \mathbf{H} + \lambda \mathbf{L})_{ab} \quad (88)$$

which is used to prove that (86) is an auxiliary function of (74). Finally, we get the update rule

$$H_{ab} \leftarrow H_{ab} \frac{(\alpha \mathbf{F}^T \mathbf{K} + 2\mu \mathbf{H} + \lambda \mathbf{H} \mathbf{A})_{ab}}{(\alpha \mathbf{F}^T \mathbf{K} \mathbf{F} \mathbf{H} + 2\mu \mathbf{H} \mathbf{H}^T \mathbf{H} + \lambda \mathbf{H} \mathbf{D})_{ab}}. \quad (89)$$

The proof of the convergence for the $\mathbf{F}$ update rule (38) can be derived by following Proposition 8 from [7]. The auxiliary function for our objective function $\mathcal{L}(\mathbf{F})$ (39) as a function of $\mathbf{F}$ is:

$$A(\mathbf{F}, \mathbf{F}') = - \sum_{i,k} 2(\mathbf{K} \mathbf{H}^T)_{i,k} \mathbf{F}'_{i,k} \left(1 + \log \frac{\mathbf{F}_{i,k}}{\mathbf{F}'_{i,k}} \right) + \sum_{i,k} \frac{(\mathbf{K} \mathbf{F}' \mathbf{H}^T)_{i,k} (\mathbf{F}_{i,k})^2}{\mathbf{F}_{i,k}}, \quad (90)$$

The proof that this is an auxiliary function of $\mathcal{L}(\mathbf{F})$ (39) is given in [7], with the change in notation $\mathbf{F} = \mathbf{W}$, $\mathbf{H} = \mathbf{G}^T$ and $\Phi(\mathbf{X}) = \mathbf{X}$.

This auxiliary function is a convex function of F and it's global minimum can be derived with the following update rule:

$$F_{ab} \leftarrow F_{ab} \frac{(\mathbf{K} \mathbf{H}^T)_{ab}}{(\mathbf{K} \mathbf{F} \mathbf{H} \mathbf{H}^T)_{ab}}. \quad (91)$$

Supplementary material

Supplementary material associated with this article can be found, in the online version, at doi:10.1016/j.patcog.2018.04.029.

References

- [1] H.S. Seung, D.D. Lee, Learning the parts of objects by non-negative matrix factorization, *Nature* 401 (6755) (1999) 788–791, doi:10.1038/44565.
- [2] C. Ding, T. Li, M.I. Jordan, Nonnegative matrix factorization for combinatorial optimization: Spectral clustering, graph matching, and clique finding, in: 2008 Eighth IEEE International Conference on Data Mining, Institute of Electrical and Electronics Engineers (IEEE), 2008, doi:10.1109/icdm.2008.130.
- [3] S. Yang, Z. Yi, M. Ye, X. He, Convergence analysis of graph regularized non-negative matrix factorization, *IEEE Trans. Knowl. Data Eng.* 26 (9) (2014) 2151–2165, doi:10.1109/tkde.2013.98.
- [4] D. Cai, X. He, J. Han, T.S. Huang, Graph regularized nonnegative matrix factorization for data representation, *IEEE Trans. Pattern Anal. Mach. Intell.* 33 (8) (2011) 1548–1560, doi:10.1109/tpami.2010.231.
- [5] S. Choi, Algorithms for orthogonal nonnegative matrix factorization, in: 2008 IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence), Institute of Electrical and Electronics Engineers (IEEE), 2008, doi:10.1109/ijcnn.2008.4634046.
- [6] C. Ding, T. Li, W. Peng, H. Park, Orthogonal nonnegative matrix t-factorizations for clustering, in: Proceedings of the 12th ACM SIGKDD international conference on Knowledge discovery and data mining - KDD 06, Association for Computing Machinery (ACM), 2006, doi:10.1145/1150402.1150420.
- [7] C. Ding, T. Li, M. Jordan, Convex and semi-nonnegative matrix factorizations, *IEEE Trans. Pattern Anal. Mach. Intell.* 32 (1) (2010) 45–55, doi:10.1109/tpami.2008.277.
- [8] F. Pompili, N. Gillis, P.-A. Absil, F. Glineur, Two algorithms for orthogonal non-negative matrix factorization with application to clustering, *Neurocomputing* 141 (2014) 15–25, doi:10.1016/j.neucom.2014.02.018.
- [9] H. Lee, A. Cichocki, S. Choi, Kernel nonnegative matrix factorization for spectral EEG feature extraction, *Neurocomputing* 72 (13–15) (2009) 3182–3190, doi:10.1016/j.neucom.2009.03.005.
- [10] B. Pan, J. Lai, W.-S. Chen, Nonlinear nonnegative matrix factorization based on Mercer kernel construction, *Pattern Recognit.* 44 (10–11) (2011) 2800–2810. Semi-Supervised Learning for Visual Content Analysis and Understanding.
- [11] E. E., V. R., Sparse subspace clustering: algorithm, theory, and applications, *IEEE Trans. Patt. Anal. Mach. Intel.* 35 (1) (2013) 2765–2781.
- [12] L. G., Z. Lin, Y. S., S. J., Y. Y., M. Y., Robust recovery of subspace structures by low-rank representation, *IEEE Trans. Patt. Anal. Mach. Intel.* 35 (1) (2013) 171–184.
- [13] L. C.-G., V. R., A structured sparse plus structured low-rank framework for subspace clustering and completion, *IEEE Trans. Sig. Proc.* 64 (24) (2016) 6557–6570.
- [14] L. C.-G., V. R., Diversity-induced multi-view subspace clustering, *IEEE Trans. Sig. Proc.* 64 (24) (2016) 6557–6570.
- [15] M. Brbic, I. Kopriva, Multi-view low-rank sparse subspace clustering, *Pattern Recognit.* 73 (2018) 247–258.
- [16] L.W. Liu, G. C., H. J., Multi-view clustering via joint nonnegative matrix factorization, *Proc. SIAM Int. Conf. Data Mining (SDM'13)* 73 (2013) 252–260.
- [17] D.-S. Pham, O. Arandjelovic, S. Venkatesh, Achieving stable subspace clustering by post-processing generic clustering results, *IEEE Int. Joint Conf. Neural Netw.* (2016).
- [18] M. Filippone, F. Camastra, F. Masulli, S. Rovetta, A survey of kernel and spectral methods for clustering, *Pattern Recognit.* 41 (1) (2008) 176–190, doi:10.1016/j.patcog.2007.05.018.
- [19] A.Y. Ng, M.I. Jordan, Y. Weiss, et al., On spectral clustering: analysis and an algorithm, *Adv. Neural Inf. Process Syst.* 2 (2002) 849–856.
- [20] Y. Liu, X. Li, C. Liu, H. Liu, Structure-constrained low-rank and partial sparse representation with sample selection for image classification, *Pattern Recognit.* 59 (2016) 5–13, doi:10.1016/j.patcog.2016.01.026.
- [21] S. White, P. Smyth, A spectral clustering approach to finding communities in graphs, in: Proceedings of the 2005 SIAM International Conference on Data Mining, Society for Industrial & Applied Mathematics (SIAM), 2005, pp. 274–285, doi:10.1137/1.9781611972757.25.
- [22] F.R. Bach, M.I. Jordan, Spectral clustering for speech separation, in: Automatic Speech and Speaker Recognition, Wiley-Blackwell, pp. 221–250, doi:10.1002/9780470742044.ch13.
- [23] H. Jia, S. Ding, H. Ma, W. Xing, Spectral clustering with neighborhood attribute reduction based on information entropy, *JCP* 9 (6) (2014), doi:10.4304/jcp.9.6.1316-1324.
- [24] H. Jia, S. Ding, X. Xu, R. Nie, The latest research progress on spectral clustering, *Neural Comput. Appl.* 24 (7–8) (2013) 1477–1486, doi:10.1007/s00521-013-1439-2.
- [25] J. Lurie, Review of spectral graph theory, *ACM SIGACT News* 30 (2) (1999) 14, doi:10.1145/568547.568553.
- [26] U. von Luxburg, A tutorial on spectral clustering, *Stat. Comput.* 17 (4) (2007) 395–416, doi:10.1007/s11222-007-9033-z.
- [27] W. Ju, D. Xiang, B. Zhang, L. Wang, I. Kopriva, X. Chen, Random walk and graph cut for co-segmentation of lung tumor on pet-ct images, *IEEE Trans. Image Process.* 25 (3) (2016) 1192.
- [28] R. Langone, R. Mall, C. Alzate, J.A.K. Suykens, Kernel spectral clustering and applications, in: Unsupervised Learning Algorithms, Springer Science Business Media, 2016, pp. 135–161, doi:10.1007/978-3-319-24211-8_6.
- [29] H. Lu, Z. Fu, X. Shu, Non-negative and sparse spectral clustering, *Pattern Recognit.* 47 (1) (2014) 418–426, doi:10.1016/j.patcog.2013.07.003.
- [30] C. Ding, X. He, H.D. Simon, On the equivalence of nonnegative matrix factorization and spectral clustering, in: Proceedings of the 2005 SIAM International Conference on Data Mining, Society for Industrial & Applied Mathematics (SIAM), 2005, pp. 606–610, doi:10.1137/1.9781611972757.70.
- [31] D. Luo, C. Ding, H. Huang, T. Li, Non-negative Laplacian embedding, in: 2009 Ninth IEEE International Conference on Data Mining, Institute of Electrical and Electronics Engineers (IEEE), 2009, doi:10.1109/icdm.2009.74.
- [32] R. Shang, Z. Zhang, L. Jiao, W. Wang, S. Yang, Global discriminative-based non-negative spectral clustering, *Pattern Recognit.* 55 (2016) 172–182, doi:10.1016/j.patcog.2016.01.035.
- [33] Y. Bengio, O. Delalleau, N.L. Roux, J.-F. Paiement, P. Vincent, M. Ouimet, Learning eigenfunctions links spectral embedding and kernel PCA, *Neural Comput.* 16 (10) (2004) 2197–2219, doi:10.1162/0899766041732396.
- [34] C. Alzate, J. Suykens, A weighted kernel PCA formulation with out-of-sample extensions for spectral clustering methods, in: The 2006 IEEE International Joint Conference on Neural Network Proceedings, Institute of Electrical and Electronics Engineers (IEEE), 2006, doi:10.1109/ijcnn.2006.246671.
- [35] P. Li, J. Bu, Y. Yang, R. Ji, C. Chen, D. Cai, Discriminative orthogonal nonnegative matrix factorization with flexibility for data representation, *Expert Syst. Appl.* 41 (4) (2014) 1283–1293, doi:10.1016/j.eswa.2013.08.026.
- [36] D.D. Lee, H.S. Seung, Algorithms for non-negative matrix factorization, in: NIPS, MIT Press, 2000, pp. 556–562.
- [37] X. Peng, H. Tang, L. Zhang, Z. Yi, S. Xiao, A unified framework for representation-based subspace clustering of out-of-sample and large-scale data, *IEEE Trans. Neural Netw. Learn. Syst.* (2015) 1–14, doi:10.1109/tnnls.2015.2490080.

- [38] C. Hou, F. Nie, D. Yi, D. Tao, Discriminative embedded clustering: a framework for grouping high-dimensional data, *IEEE Trans. Neural Netw. Learn. Syst.* 26 (6) (2015) 1287–1299, doi:[10.1109/tnnls.2014.2337335](https://doi.org/10.1109/tnnls.2014.2337335).
- [39] Y. Ma, H. Derksen, W. Hong, J. Wright, Segmentation of multivariate mixed data via lossy data coding and compression, *IEEE Trans. Pattern Anal. Mach. Intell.* 29 (9) (2007) 1546–1562, doi:[10.1109/tpami.2007.1085](https://doi.org/10.1109/tpami.2007.1085).
- [40] S.R. Rao, R. Tron, R. Vidal, Y. Ma, Motion segmentation via robust subspace separation in the presence of outlying, incomplete, or corrupted trajectories, in: 2008 IEEE Conference on Computer Vision and Pattern Recognition, Institute of Electrical and Electronics Engineers (IEEE), 2008, doi:[10.1109/cvpr.2008.4587437](https://doi.org/10.1109/cvpr.2008.4587437).
- [41] M. Belkin, P. Niyogi, Laplacian eigenmaps for dimensionality reduction and data representation, *Neural Comput.* 15 (6) (2003) 1373–1396, doi:[10.1162/089976603321780317](https://doi.org/10.1162/089976603321780317).
- [42] J. Shi, J. Malik, Normalized cuts and image segmentation, in: Proceedings of IEEE Computer Society Conference on Computer Vision and Pattern Recognition, Institute of Electrical and Electronics Engineers (IEEE), doi:[10.1109/cvpr.1997.609407](https://doi.org/10.1109/cvpr.1997.609407).
- [43] C.-K. Cheng, Y.-C. Wei, An improved two-way partitioning algorithm with stable performance (VLSI), *IEEE Trans. Comput. Aided Des. Integr. Circuits Syst.* 10 (12) (1991) 1502–1511, doi:[10.1109/43.103500](https://doi.org/10.1109/43.103500).
- [44] P.K. Chan, M.D.F. Schlag, J.Y. Zien, Spectral k-way ratio-cut partitioning and clustering, in: Proceedings of the 30th international on Design automation conference - DAC 93, Association for Computing Machinery (ACM), 1993, doi:[10.1145/157485.165117](https://doi.org/10.1145/157485.165117).
- [45] N. Gillis, S.A. Vavasis, Fast and robust recursive algorithms for separable non-negative matrix factorization, *IEEE Trans. Pattern Anal. Mach. Intell.* 36 (4) (2014) 698–714, doi:[10.1109/tpami.2013.226](https://doi.org/10.1109/tpami.2013.226).
- [46] A.N. Langville, C.D. Meyer, R. Albright, Initializations for the nonnegative matrix factorization, *KDD* (2006).
- [47] B. Scholkopf, A.J. Smola, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*, MIT Press, Cambridge, MA, USA, 2001.
- [48] T.-M. Huang, V. Kecman, I. Kopriva, *Kernel Based Algorithms for Mining Huge Data Sets: Supervised, Semi-supervised, and Unsupervised Learning (Studies in Computational Intelligence)*, Springer-Verlag New York, Inc., Secaucus, NJ, USA, 2006.
- [49] D. Cai, X. Wang, X. He, Probabilistic dyadic data analysis with local and global consistency, in: Proceedings of the 26th Annual International Conference on Machine Learning - ICML 09, Association for Computing Machinery (ACM), 2009, doi:[10.1145/1553374.1553388](https://doi.org/10.1145/1553374.1553388).
- [50] X. Niyogi, Locality preserving projections, in: *Neural information processing systems*, 16, MIT, 2004, p. 153.
- [51] J.J.-Y. Wang, J.Z. Huang, Y. Sun, X. Gao, Feature selection and multi-kernel learning for adaptive graph regularized nonnegative matrix factorization, *Expert Syst. Appl.* 42 (3) (2015) 1278–1286, doi:[10.1016/j.eswa.2014.09.008](https://doi.org/10.1016/j.eswa.2014.09.008).
- [52] A. Frank, A. Asuncion, UCI machine learning repository, <http://archive.ics.uci.edu/ml/> (2010).
- [53] F. Samaria, A. Harter, Parameterisation of a stochastic model for human face identification, in: Proceedings of 1994 IEEE Workshop on Applications of Computer Vision, Institute of Electrical and Electronics Engineers (IEEE), doi:[10.1109/acv.1994.341300](https://doi.org/10.1109/acv.1994.341300).



Dijana Tolić In 2015, she received Ph.D. degree in Physics at the University of Zagreb, in theoretical particle physics and quantum theory of fields. From 2010 to 2015 she worked at the Theoretical Physics Division and from 2015 in the Laboratory for Machine Learning and Knowledge Representations at the Rudjer Boskovic Institute, Croatia.



Nino Antulov-Fantulin In 2015, he received Ph.D. degree in Computer Science at the University of Zagreb, with the topic of statistical algorithms and complex networks. From 2016, he is the postdoctoral researcher at the ETH Zurich, Swiss Federal Institute of Technology, COSS, Switzerland.



Ivica Kopriva Received Ph.D. degree in electrical engineering from the University of Zagreb, Croatia, in 1998, with the topic on blind source separation. He has been senior research scientist the George Washington University, 2001–2005. Since 2006, he is a senior scientist at the Rudjer Boskovic Institute, Croatia.