



Nonlinear subspace clustering for image clustering



Wencheng Zhu^{a,b,c}, Jiwen Lu^{a,b,c,*}, Jie Zhou^{a,b,c}

^a Department of Automation, Tsinghua University, Beijing, 100084, China

^b State Key Lab of Intelligent Technologies and Systems, Beijing, 100084, China

^c Tsinghua National Laboratory for Information Science and Technology (TNList), Beijing, 100084, China

ARTICLE INFO

Article history:

Available online 19 August 2017

Keywords:

Subspace clustering

Neural network

Nonlinear transformation

Local similarity

ABSTRACT

We present in this paper a nonlinear subspace clustering (NSC) method for image clustering. Unlike most existing subspace clustering methods which only exploit the linear relationship of samples to learn the affine matrix, our NSC reveals the multi-cluster nonlinear structure of samples via a nonlinear neural network. While kernel-based clustering methods can also address the nonlinear issue of samples, this type of methods suffers from the scalability issue. Specifically, our NSC employs a feed-forward neural network to map samples into a nonlinear space and performs subspace clustering at the top layer of the network, so that the mapping functions and the clustering issues are iteratively learned. Otherwise, our NSC applies a similarity measure based on the grouping effect to capture the local structure of data. Experimental results illustrate that our NSC outperforms the state-of-the-arts.

© 2017 Elsevier B.V. All rights reserved.

1. Introduction

Subspace clustering has been one important visual analysis task and has many potential applications such as image and motion segmentation [17,26], face clustering [30] and so on. The objective of subspace clustering is to partition samples into different subspaces and seeks the multi-cluster structure of data [28]. Over the past decade, a number of subspace clustering methods have been proposed in the literature.

Existing subspace clustering methods can be mainly divided into four categories including algebraic based, iterative based, statistical based and spectral clustering based methods [28]. Methods in the first category apply the linear algebra or polynomial algebra theories to split the data into different subspaces, where the typical methods are matrix factorization based methods [2,15] and generalized PCA [29]. Matrix factorization based methods segment data by factorizing the data matrix into a low-rank matrix and a bases matrix [2]. These methods easily fail when the subspaces are connected or influenced by noise. Generalized PCA assumes that a set of polynomials of degree n can fit a union of n subspaces [29]. GPCA can be considered as polynomial fitting, thus GPCA can handle subspaces with different dimensions and is sensitive to noise. Methods in the second category first assign the samples to subspaces and then iteratively update the clusters of the subspaces to

partition the data. For example, the K-subspace method appoints a new sample to the nearest subspaces and updates the subspaces in the following [1]. Statistical based methods suppose that the data obey Gaussian distribution and apply Expectation Maximization (EM) to estimate the probabilities of subspaces. Two representative methods in this category include random sample consensus [9] and agglomerative lossy compression [22]. Methods in the last category [7,13,18,20,21] first utilize the self-representation property which reflects the similarity structure of data to reconstruct the original data and then impose sparse, low-rank or the grouping effect constraints on the self-representation matrix [28]. Representative methods in this category include sparse subspace clustering (SSC) [7], low-rank representation based subspace clustering (LRR) [18], least squares regression based subspace clustering (LSR) [20] and smooth representation clustering (SMR) [13]. Sparse subspace clustering (SSC) holds the assumption that each subspaces can be linear and sparse combination of other points. SSC seeks for the points in the same subspace with the sparsest representation. SSC can deal with noise and outliers [7]. Unlike SSC obtaining the sparsest representations of candidate points, LRR finds the lowest-rank representation of each data and uncovers the correlation structure of data. Otherwise, an effective way is provided to obtain robust data segmentation [18]. The block diagonal property of data between-cluster is used in SSC and LRR, however least squares regression based subspace clustering (LSR) utilizes the within-cluster correlations and segments data by using the group effect. The group effect describes that the close data have the close representations and the correlations in the real data

* Corresponding author at: Department of Automation, Tsinghua University, Beijing, 100084, China.

E-mail address: lujiwen@tsinghua.edu.cn (J. Lu).

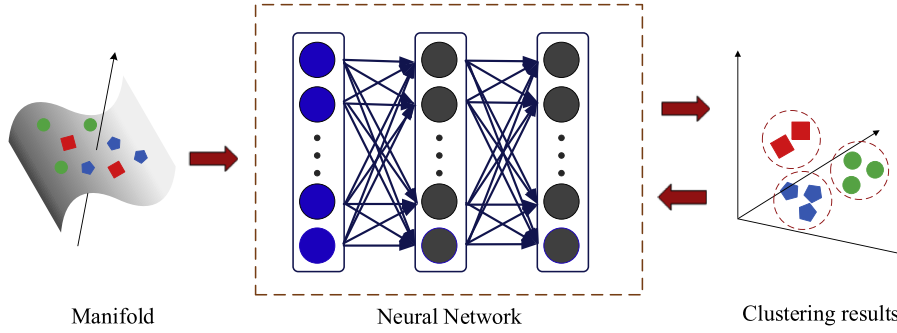


Fig. 1. The basic idea of our NSC method. We employ a feed-forward neural network to map samples into a nonlinear space and learn the self-representation matrix to perform subspace clustering at the top layer of the network. The parameters of our model are iteratively learned.

tend to be dense and highly correlated, the F -norm should be used to group data [20]. Smooth representation clustering (SMR) also applies the group effect of representation and proves the effectiveness for subspace clustering in the self-representation model. Further, SMR proposes a new form of the group effect which is more superior than the old one [13]. Spectral clustering based methods have achieved good performance, but these methods can not deal with the points between the intersection of subspaces and only utilize the linear structure of data.

Most existing subspace clustering methods tend to exploit the linear relationship of samples to learn the affine matrix, which are not powerful enough to model the nonlinear relationship of samples, especially when images are captured in wild conditions. Thus, many kernel-based clustering methods appear [24,31]. Kernel-based clustering methods embed the nonlinear data into a high-dimension space by a nonlinear mapping. However, this type of methods suffers from the scalability issue. To address this, we propose a nonlinear subspace clustering (NSC) method for image clustering. Specifically, we employ a feed-forward neural network to transform samples into a nonlinear space and learn the self-representation matrix to perform subspace clustering at the top layer of the network. The parameters of our model are iteratively learned. Fig. 1 shows the basic idea of the proposed NSC method.

We summary our contributions as follows:

1. We propose a nonlinear subspace clustering method (NSC), which utilizes a feed-forward neural network to embed samples into a nonlinear space.
2. NSC can capture the local structure of data via the grouping effect.
3. Experimental results illustrate that our NSC outperforms the state-of-the-arts.

2. Related work

Let $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] \in \mathbb{R}^{d \times n}$ denote the data matrix, $\mathbf{C} = [\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n] \in \mathbb{R}^{n \times n}$ is the self-representation matrix, where $\mathbf{x}_i$ is the i th sample of $\mathbf{X}$, the dimension of data is d and the number of data is n .

2.1. Subspace clustering

Recently, representation based subspace clustering methods [7,13,18–20] have attracted many attentions and achieved superior results. The self-representation approach is used to seek the block diagonal property of data between-cluster and then conduct spectral clustering on the self-representation matrix. Generally, these methods can be formulated as:

$$\min_{\mathbf{C}} \|\mathbf{X} - \mathbf{XC}\|_F + \lambda R(\mathbf{C}) \quad (1)$$

where $\|\cdot\|_F$ defines the suitable norm of loss function, $R(\mathbf{C})$ is the regularization term and λ is a trade-off parameter.

SSC aims to find the sparse representation of data and has the following form:

$$\min_{\mathbf{C}} \|\mathbf{C}\|_1 \quad s.t. \quad \mathbf{X} = \mathbf{XC}, \text{diag}(\mathbf{C}) = 0 \quad (2)$$

The sample can be sparse represented by other samples in the same subspace. SSC has difficulty dealing with the highly relevant samples as the only sample is selected. Otherwise, the l_1 -norm minimization problem is time-consuming to solve [7].

LRR assumes that the representation of data is low-rank as well as sparse which captures the global structure of data accurately. LRR solves the following problem:

$$\min_{\mathbf{C}} \|\mathbf{X} - \mathbf{XC}\|_{2,1} + \lambda \|\mathbf{C}\|_* \quad (3)$$

where $\|\cdot\|_*$ is the nuclear norm defined by the sum of singular values of $\mathbf{C}$. LRR is not sensitive to noise and outliers and can be effectively solved. However, whether the affine matrix forced to be low-rank leads to the good segmentation need to be analyzed further [18].

LSR has proved that if the objective function meets Enforced Block Diagonal (EBD) condition, the affine matrix is block diagonal. Thus, we can take simple norm like $\|\cdot\|_F$ to restrict the objective function.

SMR is a special case of EBD condition and applies the grouping effect to explore the local structure of data. SMR considers the following minimization problem:

$$\min_{\mathbf{C}} \|\mathbf{X} - \mathbf{XC}\|_F^2 + \lambda \text{tr}(\mathbf{C}\mathbf{L}\mathbf{C}^T) \quad (4)$$

Different from these subspace clustering methods, our proposed NSC embeds the samples into a nonlinear space by nonlinear transformation. Thus, NSC can utilize the nonlinearity of data. Otherwise, the grouping effect is applied to acquire the local structure of data.

2.2. Neural network

Due to the nonlinear transformation of neural network, neural network especially deep neural network has achieved great success in many applications, e.g. face recognition [3–6] and image classification [12]. Neural network is composed of many neural units. Each neural unit is linked to other units and has inputs, outputs, a summation function [11]. Multi-layers feed-forward neural network is a kind of neural network whose adjacent two layers are fully connected in the networks. In the feed-forward neural network, there is no connection in the same layer and cross-layer. The input of a neural unit is the outputs of the upper layer. The output of a neural unit is the result of the summation function with the weighted inputs. The properly adjusted multi-layers feed-forward

neural network can provide the suitable nonlinear transformation of input data. Since the complex neural network tends to over-fit, we use a shallow feed-forward neural network to embed the data into a nonlinear space and then conduct subspace clustering [27].

3. Nonlinear subspace clustering

In this section, we first describe the proposed model NSC and then details the optimization procedure.

3.1. Model

As described above, $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] \in \mathbb{R}^{d \times n}$ denotes the data matrix, $\mathbf{C} = [\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_n] \in \mathbb{R}^{n \times n}$ is the self-representation matrix, where $\mathbf{x}_i$ is the i th sample of $\mathbf{X}$, the dimension of data is d and the number of data is n . We utilize a multi-layer feed-forward neural network to map each sample $\mathbf{x}_i$ into a nonlinear feature space so that the nonlinear relationship of samples can be well discovered. Assume that there are $M + 1$ layers in our NSC model which conducts M times nonlinear transformations. To give a clear description of our model, we make some definitions. The input sample $\mathbf{x}_i$ is denoted as $\mathbf{h}_i^{(0)} = \mathbf{x}_i \in \mathbb{R}^d$ and the output of m th layer is denoted as

$$\mathbf{h}_i^{(m)} = g(\mathbf{W}^{(m)} \mathbf{h}_i^{(m-1)} + \mathbf{b}^{(m)}) \in \mathbb{R}^{d_m}, \quad (5)$$

where $m = 1, 2, \dots, M$ is the number of the layer in the network, $g(\cdot)$ denotes the activation function e.g. *tanh*, *linear*, *nssigmoid* [23], *sigmoid* and *ReLU* [16], d_m is the dimension of the outputs in the m th layer, $\mathbf{W}^{(m)} \in \mathbb{R}^{d_m \times d_{m-1}}$ and $\mathbf{b}^{(m)} \in \mathbb{R}^{d_m}$ are the weight and bias matrixes in the m th layer respectively [25].

Given the data matrix $\mathbf{X}$, the output $\mathbf{H}^{(M)}$ of the top layer in the neural network is defined as

$$\mathbf{H}^{(M)} = [\mathbf{h}_1^{(M)}, \mathbf{h}_2^{(M)}, \dots, \mathbf{h}_n^{(M)}]. \quad (6)$$

NSC first transforms the data matrix $\mathbf{X}$ into a nonlinear space by a multi-layers feed-forward neural network to obtain $\mathbf{H}^{(M)}$, then conducts the subspace clustering iteratively. The objective function J of NSC can be formulated as

$$\min_{\{\mathbf{W}^{(m)}, \mathbf{b}^{(m)}\}_{m=1}^M, \mathbf{C}} J = J_1 + \alpha J_2 + \beta J_3, \quad (7)$$

where J_1 is the loss function and guarantees the rebuilding ability of the self-representation matrix in the nonlinear space, which is defined as

$$J_1 = \frac{1}{2} \sum_{i=1}^n \|\mathbf{h}_i^{(M)} - \mathbf{H}^{(M)} \mathbf{c}_i\|_F^2. \quad (8)$$

J_2 utilizes the grouping effect to capture the local structure of data and the effectiveness of the grouping effect is proved in [13], which is formulated as

$$J_2 = \frac{1}{2} \text{tr}(\mathbf{C} \mathbf{L} \mathbf{C}^T), \quad (9)$$

where $\mathbf{L}$ is the Laplacian matrix and $\mathbf{L} = \mathbf{D} - \mathbf{S}$, $\mathbf{S}$ measures the similarity of data, $\mathbf{D}$ is the diagonal matrix with the element $D_{ii} = \sum_{j=1}^n S_{ij}$. J_3 is the regularization term and aims to avoid the model over-fitting, which is designed as

$$J_3 = \frac{1}{2} \sum_{m=1}^M (\|\mathbf{W}^{(m)}\|_F^2 + \|\mathbf{b}^{(m)}\|_2^2). \quad (10)$$

The corresponding α and β are the positive trade-off parameters.

Then, NSC can be expressed as

$$\min_{\{\mathbf{W}^{(m)}, \mathbf{b}^{(m)}\}_{m=1}^M, \mathbf{C}} J = \left\{ \underbrace{\frac{1}{2} \sum_{i=1}^n \|\mathbf{h}_i^{(M)} - \mathbf{H}^{(M)} \mathbf{c}_i\|_F^2}_{J_1} + \underbrace{\frac{\alpha}{2} \text{tr}(\mathbf{C} \mathbf{L} \mathbf{C}^T)}_{J_2} + \underbrace{\frac{\beta}{2} \sum_{m=1}^M (\|\mathbf{W}^{(m)}\|_F^2 + \|\mathbf{b}^{(m)}\|_2^2)}_{J_3} \right\}. \quad (11)$$

The proposed neural network model NSC facilitates the self-representation idea to learn the affine matrix at the top layer of network. The nonlinearity of data is exploited through neural network and the local structure is manipulated by the grouping effect.

3.2. Optimization

In this subsection, we present the detailed procedures of the optimization problem in (11). We update $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$ and $\mathbf{C}$ iteratively.

Update $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$: To update $\mathbf{W}^{(m)}$ and $\mathbf{b}^{(m)}$, we fix $\mathbf{C}$, $\mathbf{H}^{(M)}$ and remove the irrelevant term to obtain the following optimization problem:

$$\min_{\{\mathbf{W}^{(m)}, \mathbf{b}^{(m)}\}_{m=1}^M} \left\{ \frac{1}{2} \sum_{i=1}^n \|\mathbf{h}_i^{(M)} - \mathbf{H}^{(M)} \mathbf{c}_i\|_F^2 + \frac{\beta}{2} \sum_{m=1}^M (\|\mathbf{W}^{(m)}\|_F^2 + \|\mathbf{b}^{(m)}\|_2^2) \right\}. \quad (12)$$

The optimization problem in (12) can be solved by the sub-gradient descent algorithm.

We take the derivative of the objective in (12) with the parameters $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$ to zero and employ the chain rule [25] to obtain the following equations:

$$\frac{\partial J}{\partial \mathbf{W}^{(m)}} = \Delta^{(m)} (\mathbf{h}_i^{(m-1)})^T + \beta \mathbf{W}^{(m)}, \quad (13)$$

$$\frac{\partial J}{\partial \mathbf{b}^{(m)}} = \Delta^{(m)} + \beta \mathbf{b}^{(m)}, \quad (14)$$

where $\Delta^{(m)}$ has the following form:

$$\Delta^{(m)} = \begin{cases} (\mathbf{W}^{(m+1)})^T \Delta^{(m+1)} \odot g'(\mathbf{z}_i^{(m)}), & m = 1, \dots, M-1 \\ (\mathbf{h}_i^{(M)} - \mathbf{H}^{(M)} \mathbf{c}_i) \odot g'(\mathbf{z}_i^{(M)}), & m = M \end{cases} \quad (15)$$

where $\mathbf{z}_i^{(m)} = \mathbf{W}^{(m)} \mathbf{h}_i^{(m-1)} + \mathbf{b}^{(m)}$, $g(\cdot)$ is the activation function whose derivative is $g'(\cdot)$. The operator $\odot$ means the element-wise multiplication.

Thus, the neural network can be updated by the following paradigm:

$$\begin{cases} \mathbf{W}^{(m)} = \mathbf{W}^{(m)} - \tau \frac{\partial J}{\partial \mathbf{W}^{(m)}} \\ \mathbf{b}^{(m)} = \mathbf{b}^{(m)} - \tau \frac{\partial J}{\partial \mathbf{b}^{(m)}} \end{cases} \quad (16)$$

where $\tau > 0$ is the step size (we set $\tau = 10^{-4}$ in our experiment).

Update $\mathbf{C}$: To update $\mathbf{C}$, we fix $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$ and omit unrelated items, then we get the following optimization problem:

$$\min_{\mathbf{C}} \|\mathbf{H}^{(M)} - \mathbf{H}^{(M)} \mathbf{C}\|_F^2 + \alpha \text{tr}(\mathbf{C} \mathbf{L} \mathbf{C}^T). \quad (17)$$

We set the derivative of (17) with $\mathbf{C}$ to zero and have:

$$(\mathbf{H}^{(M)})^T (\mathbf{H}^{(M)}) \mathbf{C} + \alpha \mathbf{C} \mathbf{L} = (\mathbf{H}^{(M)})^T (\mathbf{H}^{(M)}), \quad (18)$$

the equation in (18) is the continuous Lyapunov equation which can be settled using the MATLAB “lyap” function.

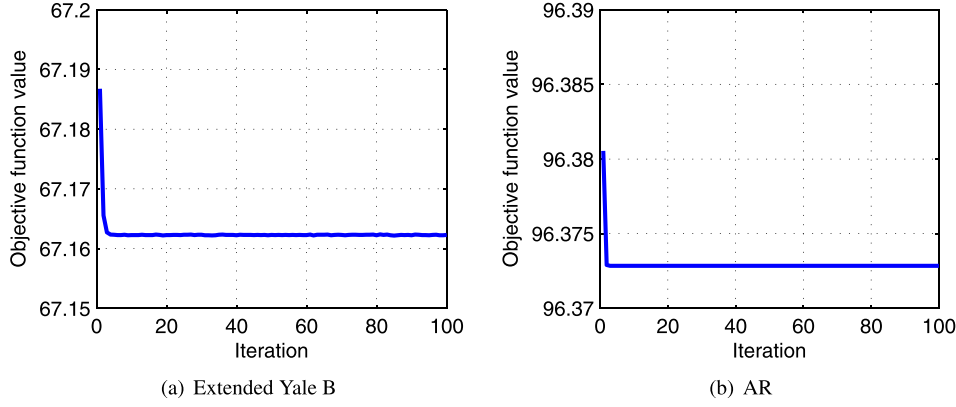


Fig. 2. The convergence curves on the Extended Yale Face B and AR datasets with parameters $\alpha = 1$ and $\beta = 1$. (a) shows the curve on Extended Yale Face B dataset and (b) shows curve on AR dataset.

We iteratively update $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$ and $\mathbf{C}$ until the objective function converges. Fig. 2 shows the convergence curve on the Extended Yale Face B and AR datasets with parameters $\alpha = 1$ and $\beta = 1$. Then, the self-representation matrix $\mathbf{C}$ is gotten and we build the graph $\mathbf{G} = |\mathbf{C}| + |\mathbf{C}^T|$. Finally, we perform spectral clustering on the graph $\mathbf{G}$. The detailed algorithm of our NSC method is summarized in Algorithm 1.

Algorithm 1: NSC.

Input:

The data matrix $\mathbf{X} = [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n] \in \mathbb{R}^{d \times n}$;
The parameters α and β ;

Output:

The neural network $\mathbf{W}^{(m)}$, $\mathbf{b}^{(m)}$ $m = 1, 2, \dots, M$;
The self-representation matrix $\mathbf{C}$;
The clustering results;

```

1 Initialize  $\mathbf{W}^{(m)}$ ,  $\mathbf{b}^{(m)}$ ,  $\mathbf{H}^{(M)}$  and  $\mathbf{C}$ ;
2 Compute the Laplacian matrix  $\mathbf{L}$ ;
3 while not converge do
4   for  $i = 1, 2, \dots, n$  do
5     Randomly select a sample  $\mathbf{x}_i$  and let  $\mathbf{h}_i^{(0)} = \mathbf{x}_i$ ;
6     Update  $\mathbf{W}^{(m)}$  and  $\mathbf{b}^{(m)}$  by (12);
7     Compute  $\mathbf{H}^{(M)}$  by (5);
8     Update  $\mathbf{C}$  by (18);
9   end
10 end
11 Build the graph  $\mathbf{G} = |\mathbf{C}| + |\mathbf{C}^T|$ ;
12 Acquire the clustering results via spectral clustering;
```

4. Experiments

In this section, we will present the implementation details and experimental results in our experiment.

4.1. Data sets and settings

We conducted experiments on two popular benchmark datasets: the Extended Yale Face B [13] and the AR [32].

The Extended Yale Face B dataset [10] is a face dataset containing 38 individuals with different pose and illumination conditions and the original size is 192×168 pixels. The first 10 persons are used, each person has 64 frontal face images and all images are resized to 48×42 pixels. We used PCA to reduce the dimension of

data into 168 dimensions. Fig. 3 shows the first 10 images for each individual with different illumination conditions.

The AR dataset [32] is a famous dataset in computer vision. A subset of AR dataset is used containing 50 subjects with variation in illumination and expression. Each subject has 7 images and the size of image is 128×128 pixels. We used VGG-16 (facenet) to extract the feature with 4096 dimensions and then applied PCA to reduce the dimension to 200 dimensions.

We employed two evaluation criteria [13,25] including the clustering error (CE) and normalized mutual information (NMI) to evaluate the performances of different subspace clustering methods. The clustering error is defined as

$$CE = 1 - \frac{1}{N} \sum_{i=1}^N \delta(p_i, q_i), \quad (19)$$

where p_i and q_i denote predicted label and the ground truth label for i th sample respectively, $\delta(\cdot)$ is the mapping function which matches the predicted labels to ground truth labels. Following, NMI is stated as

$$NMI(\mathbf{X}, \mathbf{Y}) = \frac{1}{2} \frac{I(\mathbf{X}, \mathbf{Y})}{H(\mathbf{X}) + H(\mathbf{Y})}, \quad (20)$$

where $I(\mathbf{X}; \mathbf{Y})$ and $H(\mathbf{X})$ are denoted as Eqs. (21) and (22) respectively.

$$I(\mathbf{X}; \mathbf{Y}) = \sum_{y \in \mathbf{Y}} \sum_{x \in \mathbf{X}} p(x, y) \log \left(\frac{p(x, y)}{p(x)p(y)} \right) \quad (21)$$

$$H(\mathbf{X}) = - \sum_{i=1}^n p(x_i) \log(p(x_i)) \quad (22)$$

4.2. Parameter settings

Following the work in [13], we built the similarity matrix $\mathbf{S}$ by using the k -nearest neighbor graph with 0–1 weights. Moreover, we set the neighbor size as 4. A small diagonal matrix $\theta \mathbf{I}$ is added to the Laplacian matrix $\mathbf{L}$ for the propose of numerical stability. $\mathbf{I}$ is the identity matrix and $0 < \theta \ll 1$ [13]. For a fair comparison, we used a ‘grid-search’ approach to tune parameters in the range of $\{10^{-3}, 10^{-2}, \dots, 10^3\}$. The $\tanh$ function is used as the nonlinear activation function in NSC and has the following form:

$$g(x) = \tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}, \quad (23)$$

and the derivative is denoted as

$$g'(x) = \tanh'(x) = 1 - \tanh^2(x), \quad (24)$$



Fig. 3. The Extended Yale Face B dataset. For each individual, the first 10 images are shown with different illumination conditions.

Table 1

Experimental results on the Extended Yale Face B dataset (%).

Method	SSC	LRR	LSR1	LSR2	SMR	NSC
CE	33.1	38.6	30.3	26.9	26.1	25.0
NMI	58.4	54.5	57.8	65.3	66.1	67.1

Table 2

Experimental results on the AR dataset (%).

Method	SSC	LRR	LSR1	LSR2	SMR	NSC
CE	18.3	24.9	16.0	17.7	13.7	9.1
NMI	91.4	80.1	90.7	90.7	92.1	93.6

where $\mathbf{W}^{(m)}$ and $\mathbf{b}^{(m)}$ are initialized as the identity matrix and zero matrix respectively. The neural network for the Extended Yale Face B dataset has three hidden layers with 168-100-70 neurons, α and β are set as 0.1 and 10^{-3} . We trained the neural network on AR dataset with two hidden layers and the number of each layer is 200 and 200. α and β are set as 10^3 and 0.1.

4.3. Results and analyses

We compared NSC with four state-of-the-art methods: SSC [7,8], LRR [18], LSR [20] and SMR [13]. The codes of comparison algorithms are acquired from the original authors. LSR has two different forms named LSR1 and LSR2 separately. In this subsection, we conduct the experiments on two datasets and then present and investigate the experimental results clearly.

Table 1 presents the clustering error and NMI of different subspace clustering methods on the Extended Yale Face B dataset. We see that NSC and SMR can hold the local similarity relationship and have good performances. Our NSC mines the nonlinear structure among data and outperforms SMR about 1.1% and 0.01 in clustering error and NMI.

Table 2 shows the clustering error and NMI of different subspace clustering methods on AR dataset. The best results are marked in bold. As can be seen, NSC achieves best performance in clustering error and NMI. As to clustering error, NSC outperforms LRR method by 15.8% and the best comparison method by 4.6%. For NMI, NSC improves LRR method by 0.135 and the best comparison method by 0.015.

We also investigated the sensitiveness of parameters α and β . Fig. 4 presents the clustering error and NMI of NSC on the AR and Extended Yale B datasets with different α and β . The X-axis is the parameter β , the Y-axis is the parameter α and the Z-axis represents the clustering error or NMI. For the ease of representation, we take the logarithms (base 10) of parameters. We see that NSC is not sensitive to the parameter α in clustering error and NMI. As

Table 3

Experimental results on the USPS dataset (%).

Method	SSC	LRR	LSR1	LSR2	SMR	NSC
CE	41.5	29.3	29.4	31.1	29.5	12.4
NMI	59.6	67.6	68.0	66.6	69.6	79.0

we know, the regularization term is crucial to avoid the overfitting of models. Thus, an appropriate β can lead to a good performance and the accuracies of clustering error and NMI change sharply with the variation of β . When the β is so large or small, the model has no effect. The grouping effect is applied to guide the learning of the affine matrix and only provides the tendency. For a given β , the clustering error and NMI change slightly with different α .

We also evaluated the performances of different nonlinear activation functions used in our NSC model such as the *tanh*, *linear*, *nssigmoid* [23], *sigmoid* and *ReLU* [16] functions on the AR dataset, where the clustering performances of different methods are shown in Fig. 5 (the best results are recorded). We see that the *ReLU* activation function has the worst performance in clustering error and *linear* activation function obtains the worst performance in NMI. The activation function *tanh* achieves the lowest clustering error and highest value in NMI. The difference between the activation functions *linear* and *nssigmoid* is limited.

4.4. Evaluation on the digital dataset

The USPS dataset [14] is used to evaluate the performance of NSC on digital dataset. The USPS dataset has 9298 handwritten digit images and the size of each image is 16×16 pixels. For each digit, we selected the first 100 images. The parameter settings of USPS dataset follows the parameter settings on Extended Yale Face B and AR datasets. We trained the neural network on the USPS dataset with three layers and the number of each layer is 256, 256 and 256, and α and β are set as 1000 and 10.

The clustering error and NMI results are shown in Table 3. As can be seen, NSC achieves best performance in clustering error and NMI. For clustering error, NSC outperforms SSC method by 29.1% and the best comparison method by 16.9%. As to NMI, NSC improves SSC method by 19.4% and the best comparison method by 9.4%. The digit dataset is utilized to validate that our proposed clustering method can handle the digit dataset as well as face datasets.

5. Conclusion

In this paper, we have proposed a nonlinear subspace clustering method (NSC) for image clustering. NSC simultaneously transforms the original feature space into a nonlinear space. Experimen-

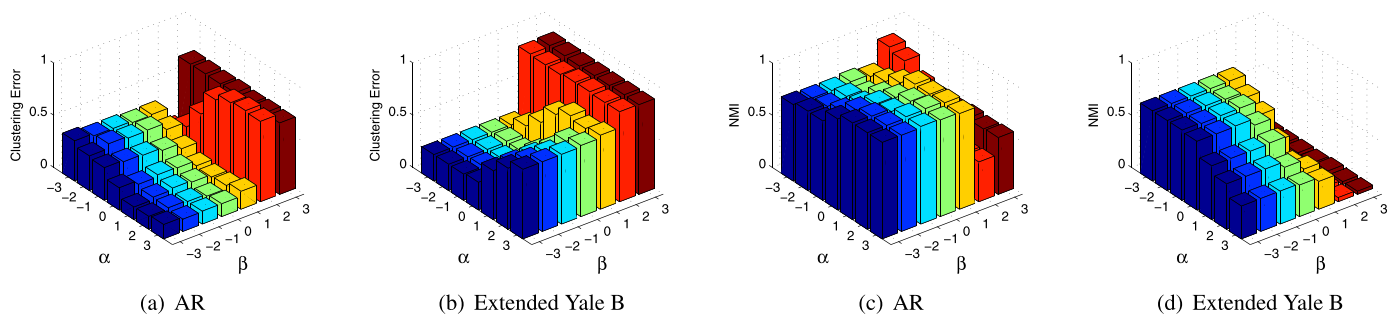


Fig. 4. The clustering error and NMI of NSC on AR and Extended Yale B datasets with different values of α and β . (a) the clustering error on AR, (b) the clustering error on Extended Yale B, (c) the NMI on AR, (d) the NMI on Extended Yale B.

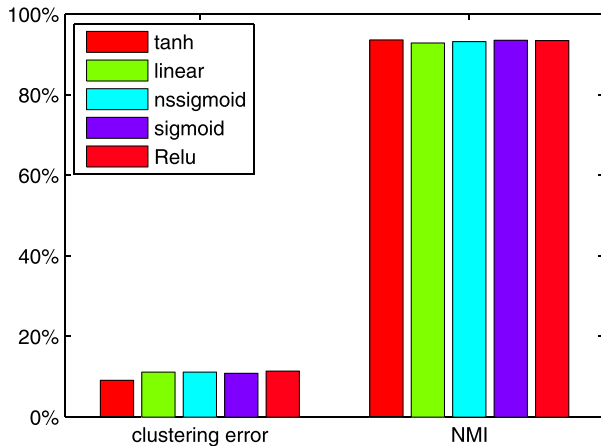


Fig. 5. The clustering error and NMI of NSC on the AR dataset with different activation functions.

tal results have clearly shown that our NSC achieve superior results than four state-of-the-art subspace clustering methods. In the future, we are going to enforce more constraints to discover the geometrical information of samples to further improve the clustering performance.

Acknowledgements

This work was supported in part by the National Key Research and Development Program of China under Grant 2016YFB1001001, in part by the National Natural Science Foundation of China under Grant 61672306, Grant 61225008, Grant 61572271, Grant 61527808, Grant 61373074, and Grant 61373090, in part by the National 1000 Young Talents Plan Program, the National Basic Research Program of China under Grant 2014CB349304, in part by the Ministry of Education of China under Grant 20120002110033, and in part by the Tsinghua University Initiative Scientific Research Program.

References

- [1] P.K. Agarwal, N.H. Mustafa, K-means projective clustering, in: ACM SIGMOD/PODS, ACM, 2004, pp. 155–165.
- [2] J.P. Costeira, T. Kanade, A multibody factorization method for independently moving objects, *IJCV* 29 (3) (1998) 159–179.
- [3] C. Ding, D. Tao, Robust face recognition via multimodal deep face representation, *TMM* 17 (11) (2015) 2049–2058.
- [4] C. Ding, D. Tao, Trunk-branch ensemble convolutional neural networks for video-based face recognition, *TPAMI* (2017).
- [5] C. Ding, D. Tao, Pose-invariant face recognition with homography-based normalization, *PR* 66 (2017) 144–152.
- [6] C. Ding, C. Xu, D. Tao, Multi-task pose-invariant face recognition, *TIP* 24 (3) (2015) 980–993.
- [7] E. Elhamifar, R. Vidal, Sparse subspace clustering, in: *CVPR*, IEEE, 2009, pp. 2790–2797.
- [8] E. Elhamifar, R. Vidal, Sparse subspace clustering: algorithm, theory, and applications, *TPAMI* 35 (11) (2013) 2765–2781.
- [9] M.A. Fischler, R.C. Bolles, Random sample consensus: a paradigm for model fitting with applications to image analysis and automated cartography, *Commun. ACM* 24 (6) (1981) 381–395.
- [10] A.S. Georgiades, P.N. Belhumeur, D.J. Kriegman, From few to many: illumination cone models for face recognition under variable lighting and pose, *TPAMI* 23 (6) (2001) 643–660.
- [11] D. Graupe, Principles of artificial neural networks, 7, World Scientific, 2013.
- [12] K. He, X. Zhang, S. Ren, J. Sun, Deep residual learning for image recognition, in: *CVPR*, 2016, pp. 770–778.
- [13] H. Hu, Z. Lin, J. Feng, J. Zhou, Smooth representation clustering, in: *CVPR*, 2014, pp. 3834–3841.
- [14] J.J. Hull, A database for handwritten text recognition research, *TPAMI* 16 (5) (1994) 550–554.
- [15] K.-i. Kanatani, Motion segmentation by subspace separation and model selection, in: *ICCV*, 2, IEEE, 2001, pp. 586–591.
- [16] A. Krizhevsky, I. Sutskever, G.E. Hinton, Imagenet classification with deep convolutional neural networks, in: *NIPS*, 2012, pp. 1097–1105.
- [17] F. Lauer, C. Schnörr, Spectral clustering of linear subspaces for motion segmentation, in: *ICCV*, IEEE, 2009, pp. 678–685.
- [18] G. Liu, Z. Lin, S. Yan, J. Sun, Y. Yu, Y. Ma, Robust recovery of subspace structures by low-rank representation, *TPAMI* 35 (1) (2013) 171–184.
- [19] G. Liu, S. Yan, Latent low-rank representation for subspace segmentation and feature extraction, in: *ICCV*, IEEE, 2011, pp. 1615–1622.
- [20] C.-Y. Lu, H. Min, Z.-Q. Zhao, L. Zhu, D.-S. Huang, S. Yan, Robust and efficient subspace segmentation via least squares regression, in: *ECCV*, Springer, 2012, pp. 347–360.
- [21] D. Luo, F. Nie, C. Ding, H. Huang, Multi-subspace representation and discovery, in: *ECML-PKDD*, Springer, 2011, pp. 405–420.
- [22] Y. Ma, H. Derksen, W. Hong, J. Wright, Segmentation of multivariate mixed data via lossy data coding and compression, *TPAMI* 29 (9) (2007).
- [23] V. Nair, G.E. Hinton, Rectified linear units improve restricted boltzmann machines, in: *ICML*, 2010, pp. 807–814.
- [24] V.M. Patel, R. Vidal, Kernel sparse subspace clustering, in: *ICIP*, IEEE, 2014, pp. 2849–2853.
- [25] X. Peng, S. Xiao, J. Feng, W.-Y. Yau, Z. Yi, Deep subspace clustering with sparsity prior, in: *IJCAI*, 2016, pp. 1925–1931.
- [26] S. Rao, R. Tron, R. Vidal, Y. Ma, Motion segmentation in the presence of outlying, incomplete, or corrupted trajectories, *TPAMI* 32 (10) (2010) 1832–1845.
- [27] J. Schmidhuber, Deep learning in neural networks: an overview, *Neural Netw.* 61 (2015) 85–117.
- [28] R. Vidal, Subspace clustering, *IEEE SPM* 28 (2) (2011) 52–68.
- [29] R. Vidal, Y. Ma, S. Sastry, Generalized principal component analysis (gpca), *TPAMI* 27 (12) (2005) 1945–1959.
- [30] S. Xiao, M. Tan, D. Xu, Weighted block-sparse low rank representation for face clustering in videos, in: *ECCV*, 2014, pp. 123–138.
- [31] S. Xiao, M. Tan, D. Xu, Z.Y. Dong, Robust kernel low-rank representation, *TNNLS* 27 (11) (2016) 2268–2281.
- [32] L. Zhang, M. Yang, X. Feng, Sparse representation or collaborative representation: which helps face recognition? in: *ICCV*, IEEE, 2011, pp. 471–478.