



Robust Subspace Segmentation by Low-Rank Representation

ICML 2010

Latent Low-Rank Representation for Subspace Segmentation and Feature Extraction

ICCV 2011

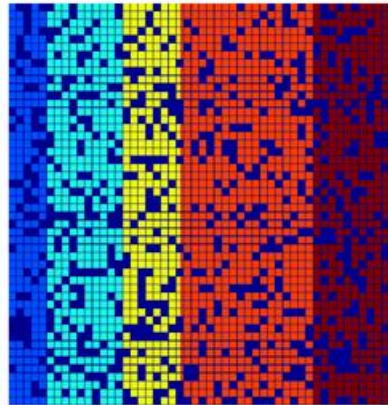
Integrated Low-Rank-Based Discriminative Feature Learning for Recognition

TNNLS 2016

2018.5.17

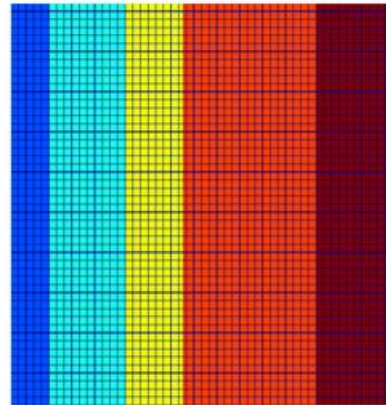
Problem Description

Rank- r matrix A of size $m \times n$, where $r \ll \min(m, n)$. Model the observed matrix D to **be a set of linear measurements** on the matrix A , subject to **noise and gross corruptions** i.e., $D = L(A) + \eta$, where L is a **linear operator**, and η represents **the matrix of corruptions**. We seek to recover the true matrix A from D .



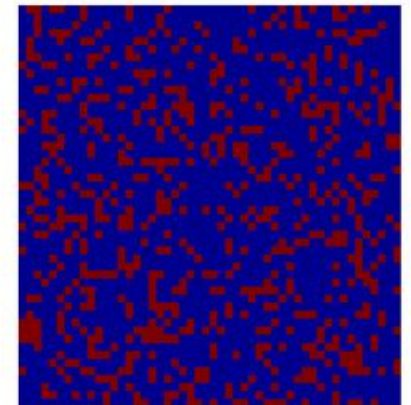
Matrix of corrupted observations

=



Underlying low-rank matrix

+

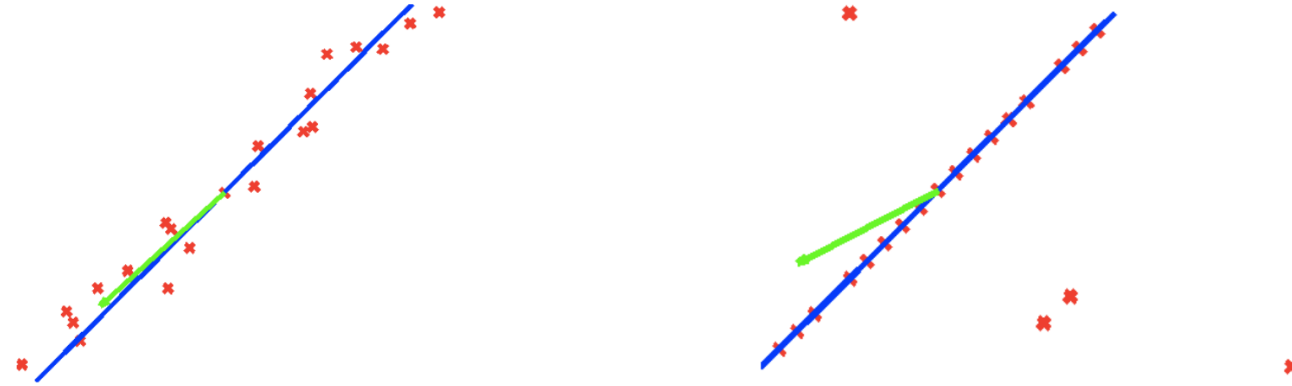


Sparse error matrix



PCA

L is the **identity operator** and the entries of η are **independent and identically distributed according to a isotropic Gaussian distribution**, then classical PCA provides the optimal estimate to A .



Matrix Completion

η is **zero**, L is the **matrix subsampling operator**, the problem is to use information from some entries of A to infer its missing entries.

$$\min_X \text{rank}(X) \quad s.t. \quad L(X) = D$$

$$\min_X \|X\|_* \quad s.t. \quad L(X) = D$$

nuclear norm: the sum of the singular values of the matrix

$$\text{rank}(A^H A) = \text{rank}(A A^H) = \text{rank}(A) \quad A^H A = \sigma^2 u$$

L is the identity operator and η is a sparse matrix, the problem is to **find the matrix of lowest rank that could have generated D when added to an unknown sparse matrix η .**

$$\min_{X,E} \text{rank}(X) + \gamma \|E\|_0 \quad s.t. \quad D = X + E \quad \xrightarrow{\text{convex relaxation}} \quad \min_{X,E} \|X\|_* + \gamma \|E\|_1 \quad s.t. \quad D = X + E$$

The **Augmented Lagrange Multiplier** method for exact recovery of corrupted low-rank matrices

The general method of augmented Lagrange multipliers is introduced for solving constrained optimization problems of the kind:

$$\min f(X) \quad s.t. \quad h(X) = 0$$

$$L(X, Y, \mu) = f(X) + \langle Y, h(X) \rangle + \frac{\mu}{2} \|h(X)\|_F^2$$

$$\langle Y, h(X) \rangle = \text{tr}(Y^T h(X))$$



Algorithm 3 (General Method of Augmented Lagrange Multiplier)

```
1:  $\rho \geq 1$ .  
2: while not converged do  
3:   Solve  $X_{k+1} = \arg \min_X L(X, Y_k, \mu_k)$ .  
4:    $Y_{k+1} = Y_k + \mu_k h(X_{k+1})$ ;  
5:   Update  $\mu_k$  to  $\mu_{k+1}$ .  
6: end while
```

Output: X_k .

$$\min f(X) \quad s.t. \quad h(X) = 0 \qquad L(X, Y, \mu) = f(X) + \langle Y, h(X) \rangle + \frac{\mu}{2} \|h(X)\|_F^2$$

$$\min_{X, E} \|X\|_* + \gamma \|E\|_1 \quad s.t. \quad D = X + E$$

$$L(X, E, Y, \mu) = \|X\|_* + \gamma \|E\|_1 + \langle Y, D - X - E \rangle - \frac{\mu}{2} \|D - X - E\|_F^2$$



Problem

Given a set of sufficiently dense data vectors $X = [x_1, x_2, \dots, x_n]$ (each column is a sample) drawn from a union of k subspaces $\{S_i\}_{i=1}^k$ of unknown dimensions, in a D -dimensional Euclidean space, segment all data vectors into their respective subspaces.

Assumption

The subspaces are low-rank and independent, and the data is noiseless. $\sum_{i=1}^k S_i = \bigoplus_{i=1}^k S_i$

A fraction of the data vectors are corrupted by noise or contaminated by outliers, or to be more precise, the data **contains sparse and properly bounded errors**.



Consider data vectors $X = [x_1, x_2, \dots, x_n]$, $x_i \in R^D$, each of which can be represented by the linear combination of the basis in a dictionary $A = [a_1, a_2, \dots, a_m]$

$$X = AZ \quad Z = [z_1, z_2, \dots, z_n]$$

[2009] Sparse representations using an appropriate dictionaries A may reveal the clustering of the points x_i . However, sparse representation **may not capture the global structures** of the data X . **Low rankness** may be a more appropriate criterion.

$$\min_Z \text{rank}(Z) \quad \text{s.t. } X = AZ$$

a good surrogate

$$\min_Z ||Z||_* \quad \text{s.t. } X = AZ$$

$$X = [x_1, x_2, \dots, x_n] \quad \{S_i\}_{i=1}^k \quad \{d_i\}_{i=1}^k \quad \{n_i\}_{i=1}^k \quad X = [X_1, X_2, \dots, X_k]$$

In order to segment the data into their respective subspaces, we **need to compute an affinity matrix that encodes the pairwise affinities between data vectors**. So we use the data X itself as the dictionary.

$$\min_Z ||Z||_* \quad \text{s. t. } X = XZ$$

There always exist feasible solutions even when the data sampling is insufficient.

Theorem 3.1

Assume that the data sampling is sufficient such that $n_i > \text{rank}(X_i) = d_i$. If the subspaces are independent then there exists an optimal solution Z^* .

$$Z^* = \begin{bmatrix} Z_1^* & 0 & 0 & 0 \\ 0 & Z_2^* & 0 & 0 \\ 0 & 0 & \ddots & 0 \\ 0 & 0 & 0 & Z_k^* \end{bmatrix}_{n \times n}$$



Theorem 3.1 **does not guarantee that an arbitrary optimal solution to the problem is block-diagonal**. The difficulty is essentially that the minimizer is **non-unique**. However, in our simulations we have observed that the solution obtained is always block-diagonal, and so we do not pursue this here.

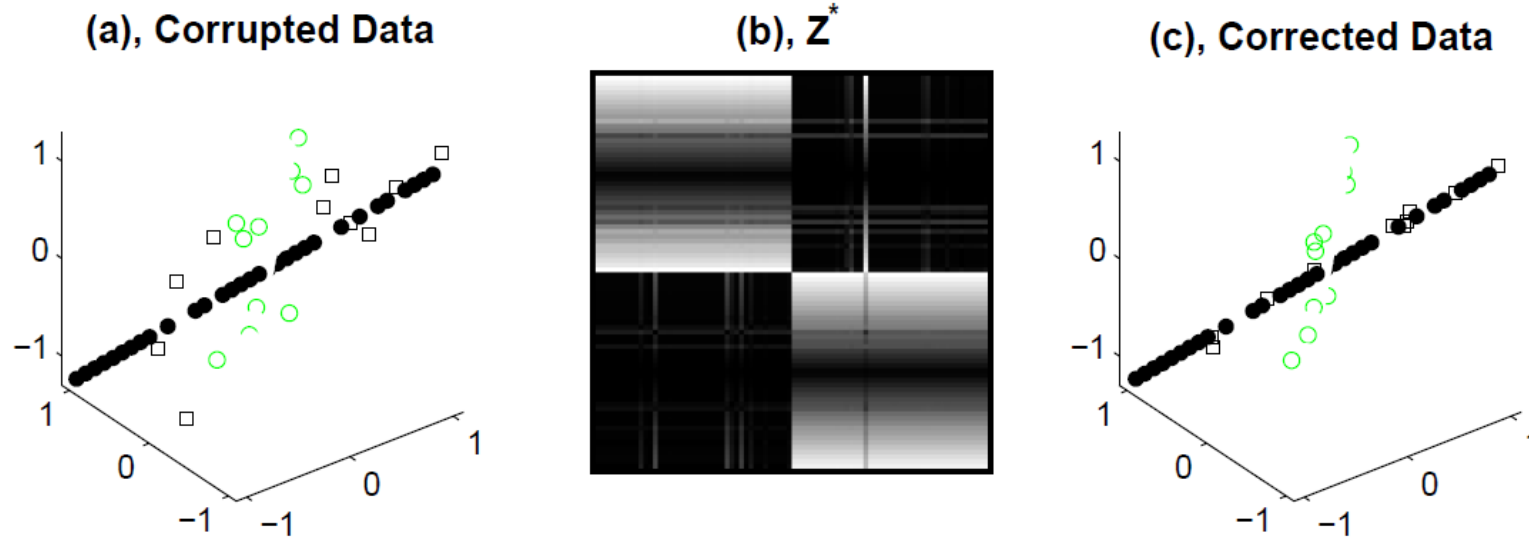
Robustness to Noise and Outliers

$$\min_{Z,E} \|Z\|_* + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E \quad \|E\|_{2,1} = \sum_{j=1}^n \sqrt{\sum_{i=1}^n (E_{ij})^2}$$

$$(Z^*, E^*) \xrightarrow{\text{recover data}} X - E^*/XZ^*$$

$$X = X_l + X_c$$

In the case that the remainder clean data is still sufficient to represent the subspaces, and the corruptions are properly bounded, it shall **automatically correct the corruptions** so as to obtain the lowest-rank representation.



There are about 80 data vectors sampled from two one-dimensional subspaces embedded in R^3 , and about 25% data vectors are corrupted by large Gaussian errors.

sufficient
bounded



it shall automatically correct the corruptions so as to obtain the lowest-rank representation.



$$\min_{Z,E} \|Z\|_* + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E \quad \longleftrightarrow \quad \min_{Z,E,J} \|J\|_* + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E, Z = J$$

$$\xrightarrow{\text{ALM}} \quad \min f(X) \quad \text{s.t. } h(X) = 0 \quad L(X, Y, \mu) = f(X) + \langle Y, h(X) \rangle + \frac{\mu}{2} \|h(X)\|_F^2$$

$$\min_{Z,E,J,Y_1,Y_2} \|J\|_* + \lambda \|E\|_{2,1} + \text{tr}[Y_1^T (X - XZ + E)] + \text{tr}[Y_2^T (Z - J)] + \frac{\mu}{2} (\|X - XZ + E\|_F^2 + \|Z - J\|_F^2)$$

1. fix the others and update J
2. fix the others and update Z
3. fix the others and update E
4. update the multipliers Y_1, Y_2
5. update the parameter μ
6. check the convergence conditions

$$\|X - XZ - E\|_\infty < \varepsilon \quad \text{and} \quad \|Z - J\|_\infty < \varepsilon$$

$$\min_{Z,E,Y} \|Z\|_* + \lambda \|E\|_{2,1} + \text{tr}[Y^T (X - XZ + E)] + \frac{\mu}{2} (\|X - XZ + E\|_F^2)$$

$$\|X - XZ - E\|_\infty < \varepsilon \quad \text{and} \quad \|Z^{t+1} - Z^t\|_\infty < \varepsilon$$



稀疏表示

$$SR_1: \min_{Z,E} \|Z\|_1 + \lambda \|E\|_1 \quad \text{s.t. } X = XZ + E, Z_{ii} = 0$$

$$SR_{2,1}: \min_{Z,E} \|Z\|_1 + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E, Z_{ii} = 0$$

$$\|x\|_1 = \sum_{i=1}^N |x_i|$$

$$\|Z\|_1 = \max_j \sum_{i=1}^m |Z_{i,j}|$$

鲁棒PCA

$$\min_{X,E} \|X\|_* + \gamma \|E\|_1 \quad \text{s.t. } D = X + E$$

低秩稀疏表示LRR

$$\min_{Z,E} \|Z\|_* + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E$$

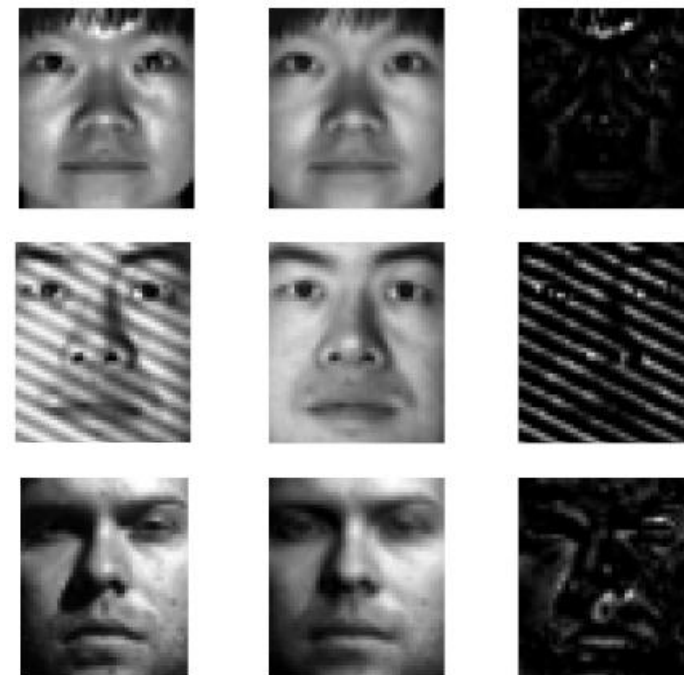
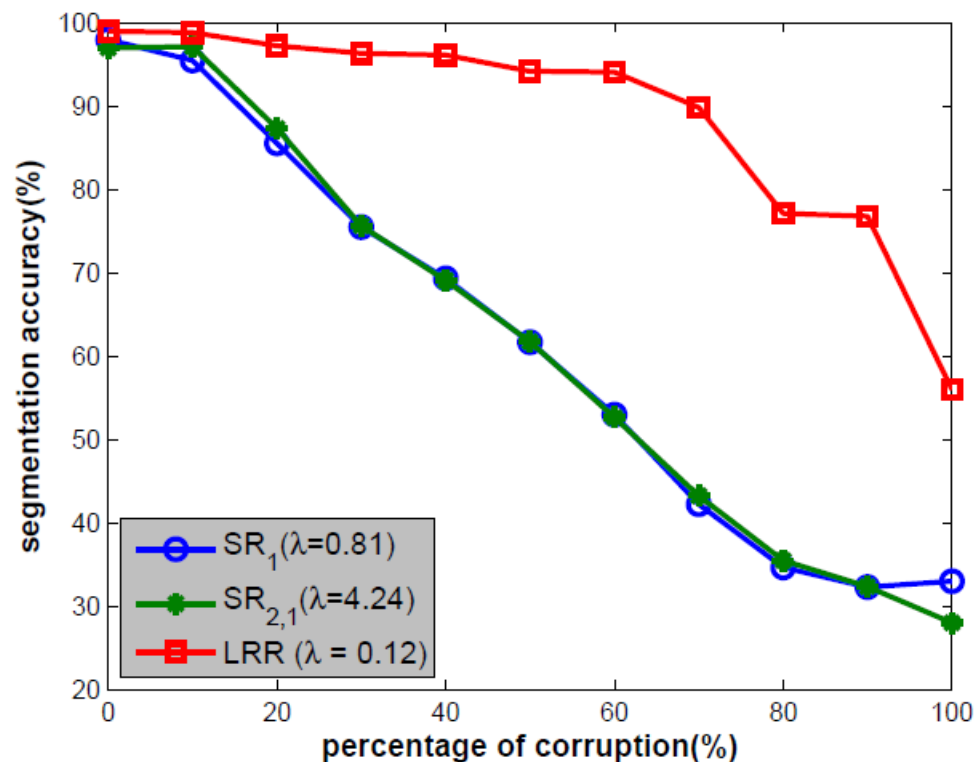
Algorithm 2 Subspace Segmentation by LRR

Input: data matrix X , number of subspaces k

1. obtain the lowest-rank representation by solving problem (5)
2. construct an undirected graph by using the lowest-rank representation to define the affinity matrix of the graph
3. use NCut to segment the vertices of the graph into k clusters

Some examples of using LRR to correct the corruptions in faces.

$$X = XZ^* + E^*$$





$$\min_Z \|Z\|_* \quad s.t. \quad X_O = X_O Z$$

$$Z_O^* = I$$



cannot use X_O as the dictionary to represent the subspaces if the data sampling is insufficient.



LRR requires that sufficient noiseless data is available in the dictionary A , i.e., only a part of A is corrupted. this assumption may be **invalid** and the robustness of LRR may be depressed in reality.

$$\min_Z \|Z\|_* \quad s.t. \quad X_O = [X_O, X_H]Z$$

X_O is the observed data matrix and X_H represents the unobserved, hidden data

$$Z_{O,H}^* = [Z_{O|H}^*; Z_{H|O}^*]$$



Problem 1 (Noiseless Data)

$$\min_Z \|Z\|_* \quad s.t. \quad X_O = [X_O, X_H]Z$$

Problem 2 (Corrupted Data)

$$\min_{Z,E} \|Z\|_* + \lambda \|E\|_1 \quad s.t. \quad X_O = [X_O, X_H]Z + E$$

Suppose $Z_{O,H}^* = [Z_{O|H}^*; Z_{H|O}^*]$ is the optimal solution (with respect to the variable Z) and $Z_{O|H}^*$ is the submatrix corresponding to X_O , then our goal is to **recover by using only the observed data X_O** .



$$[X_O, X_H] = U\Sigma V^T = U\Sigma[V_O; V_H]^T \quad X_O = U\Sigma V_O^T \quad X_H = U\Sigma V_H^T$$

$$U\Sigma V_O^T = U\Sigma V^T Z \quad V_O^T = V^T Z$$

$$\min_Z \|Z\|_* \quad s.t. \quad X_O = [X_O, X_H]Z \quad \xrightarrow{\text{a unique minimizer [PAMI2013]}} \quad Z_{O,H}^* = VV_O^T = [V_O V_O^T; V_H V_O^T]$$



$$\begin{aligned} X_O &= [X_O, X_H]Z_{O,H}^* = X_O Z_{O|H}^* + X_H Z_{H|O}^* = X_O Z_{O|H}^* + X_H V_H V_O^T \\ &= X_O Z_{O|H}^* + U \Sigma V_H^T V_H V_O^T = X_O Z_{O|H}^* + U \Sigma V_H^T V_H \Sigma^{-1} U^T X_O \end{aligned}$$

$$L_{H|O}^* = U \Sigma V_H^T V_H \Sigma^{-1} U^T$$

$$X_O = X_O Z_{O|H}^* + L_{H|O}^* X_O$$

X_O and X_H are sampled from the same collection of low-rank subspaces

$$\text{rank}(Z_{O|H}^*) \leq r \quad \text{and} \quad \text{rank}(L_{H|O}^*) \leq r$$

$$\min_Z \|Z\|_* \quad \text{s.t.} \quad X_O = [X_O, X_H]Z$$

$$\min_{Z,L} \|Z\|_* + \|L\|_* \quad \text{s.t.} \quad X = XZ + LX$$



$$\min_{Z_{O|H}, L_{H|O}} \text{rank}(Z_{O|H}) + \text{rank}(L_{H|O}) \quad \text{s.t.} \quad X_O = X_O Z_{O|H} + L_{H|O} X_O$$



$$\min_Z \|Z\|_* \quad s.t. X = XZ \quad Z_Z^*$$

$$\min_L \|L\|_* \quad s.t. X = LX \quad L_L^*$$

[PAMI2013] $\|Z_Z^*\|_* = rank(X) = rank(X^T) = \|L_L^*\|_*$ So the strengths of L and Z are balanced naturally.

Setting

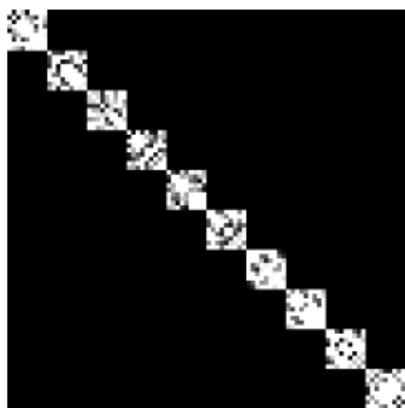
$$\{S_i\}_{i=1}^{10} \quad \{U_i\}_{i=1}^{10} (U_{i+1} = TU_i), \quad T \text{ is a random rotation and } U_1 \in R^{200 \times 10}$$

$$X_O = [X_1, X_2, \dots, X_{10}] \quad X_i = U_i C_i, 1 \leq i \leq 10$$

hidden matrix $X_H (200 \times 50)$

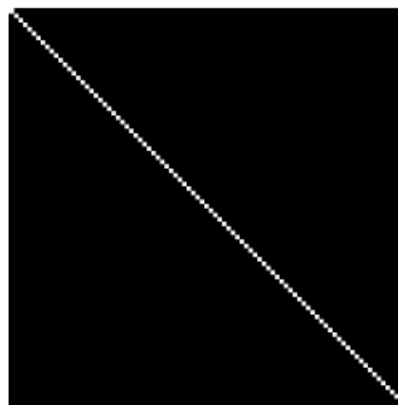
Example

$$\min_Z \|Z\|_* \quad s.t. X_O = [X_O, X_H]Z$$



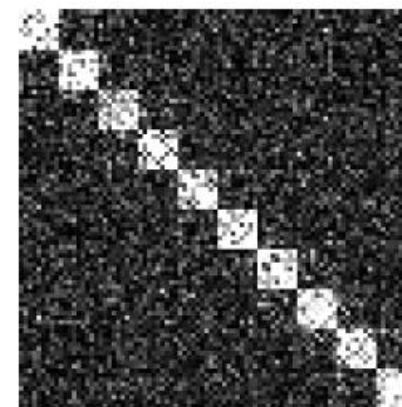
(a) $Z_{O|H}^*$

$$\min_Z \|Z\|_* \quad s.t. X_O = X_O Z$$



(b) LRR

$$\min_{Z,L} \|Z\|_* + \|L\|_* \quad s.t. X = XZ + LX$$

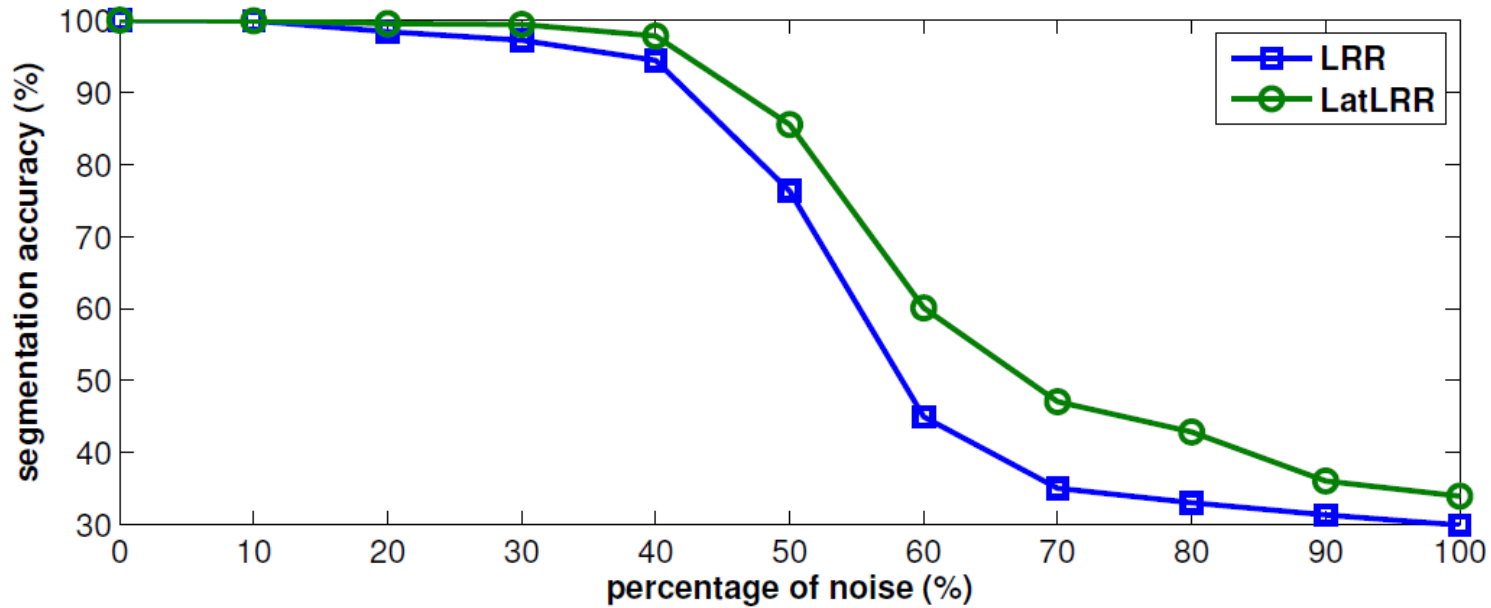


(c) LatLRR

Corrupted Data

$$\min_{Z,L,E} \|Z\|_* + \|L\|_* + \lambda \|E\|_1 \quad s.t. X = XZ + LX + E$$

$$\min_{Z,L,J,S,E} \|J\|_* + \|S\|_* + \lambda \|E\|_1 \quad s.t. \quad X = XZ + LX + E, Z = J, L = S$$





In summary, let (Z^*, L^*, E^*) be the minimizer, then Z^*, L^* are useful for **subspace segmentation** and **feature extraction**, respectively.

Subspace Segmentation

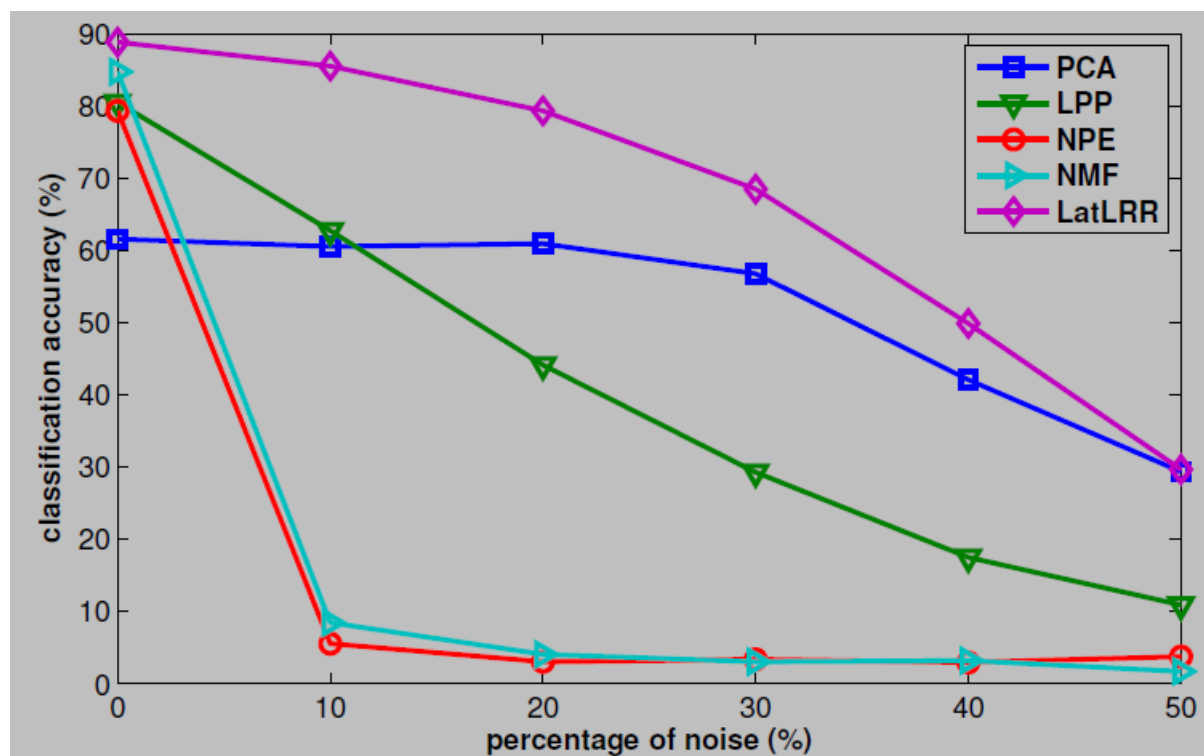
Utilize the **affinity matrix** identified by to define edge weights of an undirected graph, and then use Normalized Cuts (NCut) [23] to produce the final segmentation results.

Comparison under the same setting					
	LSA	RANSAC	SR	LRR	LatLRR
Mean	8.99	8.22	3.89	3.16	2.95
Std.	9.80	10.26	7.70	5.99	5.86
Max	37.74	47.83	32.57	37.43	37.97
Comparison to state-of-the-art methods					
	LSA	ALC	SSC	SC	LatLRR
Mean	4.94	3.37	1.24	1.20	0.85

Feature Extraction

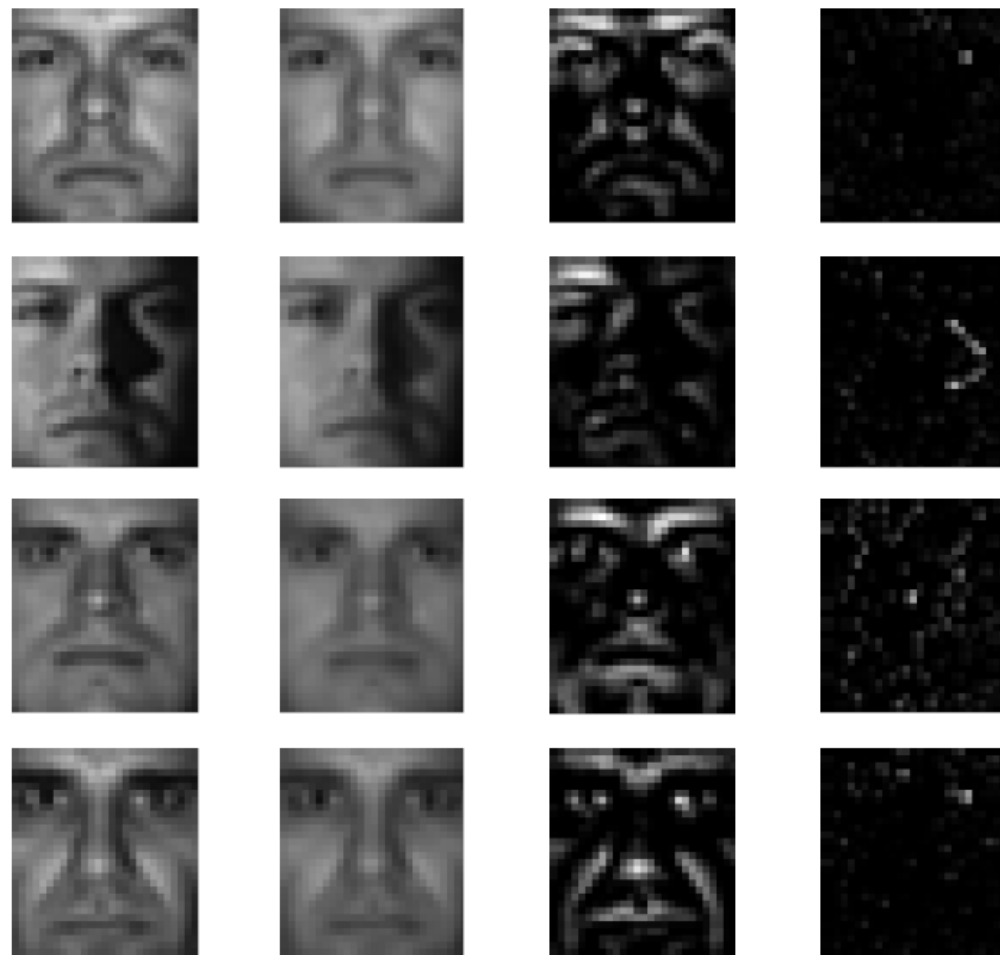
L^* may be useful for feature extraction, and experimentally find that L^* is able to extract “salient features” (i.e., notable features such as the eyes of faces) from data.

$$y = L^* x$$



The salient features correspond to the key object parts (e.g., the eyes), which are usually discriminative for recognition

$$\begin{array}{ccccc} X & = & XZ^* & + & L^* X & + & E^* \\ \text{data} & = & \text{principal features} & + & \text{salient features} & + & \text{sparse noise} \end{array}$$





Suppose P is a low-dimensional projection learnt by using L^*x as inputs for some dimension reduction methods, the reduced feature vector y of a testing data vector x can be computed by $y = P^T L^* x$.

	Raw Data	PCA (317D)	LPP (83D)	NPE (325D)	NMF (195D)	LatLRR	PCA(400D)	LatLRR + LPP(52D)	NPE(400D)
1-NN	61.07	61.54	80.46	79.28	84.69	88.76	87.28	87.60	82.18
3-NN	59.81	60.03	78.73	79.28	84.07	87.76	85.95	87.13	81.71
5-NN	58.16	58.54	76.69	76.69	82.58	86.03	85.87	85.56	80.85

LatLRR extends LRR to **handle the hidden effects**

As a subspace **segmentation algorithm**, LatLRR outperforms the state-of-the-art algorithms for motion segmentation.

Conclusions

Be empirically able to automatically **extract salient features from corrupted data**.

Compared to dimension reduction based methods, LatLRR is **more robust to noise**.



The basic idea is to **utilize the supervised information**, e.g., the labels of training samples, to learn discriminative features LX resulting from the LatLRR model.

Theorem

$$\min_{Z,L} \|Z\|_* + \|L\|_* \quad s.t. X = XZ + LX \xrightarrow{\text{complete solutions}} Z^* = V_X(I - S)V_X^T \quad \text{and} \quad L^* = U_X S U_X^T$$

$SVD(X) = U_X \Sigma_X U_X^T$ S is any block-diagonal matrix that satisfies two constraints:

- 1) its blocks are compatible with Σ_X , i.e., if $(\Sigma_X)_{ii} \neq (\Sigma_X)_{jj}$, then $S_{ij} = 0$
- 2) both S and $I - S$ are positive semidefinite.

Note that S can usually be chosen as diagonal with diagonal entries being any number between 0 and 1.



$$\min_{L,W} \sum_{i=1}^m \varphi(h_i, f(Lx_i, W)) + \alpha ||W||_F^2$$
$$s. t. \quad L = U_X S U_X^T$$

$$U_X \in R^{d \times r} \quad S \in R^{r \times r}$$

$$f(x, W) = Wx, \quad W \in R^{c \times d}$$

$$\min_{W,L} ||H - WLX||_F^2 + \alpha ||W||_F^2$$
$$s. t. \quad L = U_X S U_X^T$$

$$H = [h_1, h_2, \dots, h_n] \in R^{c \times m}$$
$$h_i = [0, 0, \dots, 1, \dots, 0, 0]^T \in R^c$$

solve

1) the singular values of the data matrix X are usually distinct from each other, i.e. $(\Sigma_X)_{ii} \neq (\Sigma_X)_{jj}$, when $S_{ij} = 0$.

2) since only focus on learning the discriminative features, the constraint that $I - S$ is positive semidefinite is not necessary, so, only need to bound $S_{ii} > 0$.



$$\text{diag}(\Lambda) = (S_{11}, S_{22}, \dots, S_{rr}, 0, 0, \dots, 0) \in R^d \quad \Lambda \in R^{d \times d}$$

$$L = U_X S U_X^T = U \Lambda U^T \quad U U^T = U^T U = V V^T = I$$

$$\begin{aligned} \|H - WLX\|_F^2 + \alpha \|W\|_F^2 &= \|H - W U \Lambda U^T U \Sigma V^T\|_F^2 + \alpha \|W\|_F^2 \\ &= \|H V - W U \Lambda \Sigma\|_F^2 + \alpha \|W\|_F^2 = \|H V - W U \Lambda \Sigma\|_F^2 + \alpha \|W U\|_F^2 \end{aligned}$$

$$\tilde{H} = H V \quad \tilde{W} = W U$$

$$\begin{aligned} \|H - WLX\|_F^2 + \alpha \|W\|_F^2 &= \|\tilde{H} - \tilde{W} \Lambda \Sigma\|_F^2 + \alpha \|\tilde{W}\|_F^2 \\ &= \sum_{i=1}^r \|\tilde{H}_i - S_{ii} \sigma_i \tilde{W}_i\|_2^2 + \sum_{i=r+1}^m \|\tilde{H}_i\|_2^2 + \alpha \sum_{i=1}^r \|\tilde{W}_i\|_2^2 + \alpha \sum_{i=r+1}^m \|\tilde{W}_i\|_2^2 \end{aligned}$$

$$\tilde{W}_i = 0, i = r + 1, \dots, m$$



$$\min_{\substack{S_{11}, \dots, S_{rr} \\ \tilde{W}_1, \dots, \tilde{W}_r}} \sum_{i=1}^r \left(\|\tilde{H}_i - S_{ii} \sigma_i \tilde{W}_i\|_2^2 + \alpha \|\tilde{W}_i\|_2^2 \right) \quad \text{s.t. } S_{ii} \geq 0, i = 1, 2, \dots, r$$

$$\sum_{i=1}^r S_{ii} = t \quad \sum_{i=1}^r S_{ii} \sigma_i = t \quad g = [S_{11} \sigma_1, \dots, S_{rr} \sigma_r]^T \quad Q = [S_{11} \sigma_1 \tilde{W}_1, \dots, S_{rr} \sigma_r \tilde{W}_r]$$

$$\min_{g, Q} \sum_{i=1}^r \left(\|\tilde{H}_i - Q_i\|_2^2 + \frac{\alpha}{g_i^2} \|Q_i\|_2^2 \right) \quad \text{s.t. } \sum_{i=1}^r g_i = t, \quad g_i \geq 0, i = 1, 2, \dots, r$$

fix g update of Q

$$Q_i = \min_{Q_i} \|\tilde{H}_i - Q_i\|_2^2 + \frac{\alpha}{g_i^2} \|Q_i\|_2^2 = \frac{g_i^2}{g_i^2 + \alpha} \tilde{H}_i \quad i = 1, 2, \dots, r$$

fix Q update of g

$$\operatorname{argmin}_{g_i} \sum_{i=1}^r \frac{\alpha}{g_i^2} \|Q_i\|_2^2 \quad \text{s.t. } \sum_{i=1}^r g_i = t, \quad g_i \geq 0, i = 1, 2, \dots, r$$



$$L(g, \tau) = \sum_{i=1}^r \frac{\alpha}{g_i^2} \|Q_i\|_2^2 + \tau \left(\sum_{i=1}^r g_i - t \right)$$

$$\frac{\partial L}{\partial g_i} = -\frac{2\alpha \|Q_i\|_2^2}{g_i^3} + \tau = 0 \quad \sum_{i=1}^r g_i = \tau \quad g_i = \frac{t \|Q_i\|_2^{2/3}}{\sum_{i=1}^r \|Q_i\|_2^{2/3}}$$

$$L = U\Lambda U^T \quad \Lambda_{ii} = \frac{g_{ii}}{\sigma_i} \quad (i = 1, \dots, r) \quad \tilde{W}_i = \frac{Q_i}{g_i} \quad Z = LX$$

[Neural Comput. 2015] proposed denoise X first, apply the noiseless LRR or the LatLRR to the denoised data.

$$\min_{Z, L, E} \|Z\|_* + \|L\|_* + \lambda \|E\|_1 \quad s.t. X - E = (X - E)Z + L(X - E)$$

[Neural Comput. 2015] proved solving the above problem is equivalent to denoising X with the robust PCA first to obtain (A, E) , and then solving noiseless LatLRR with X replaced by A .



Let the pair (A^*, E^*) be any optimal solution to the robust PCA problem. Then, the new noisy LatLRR model has minimizers (A^*, L^*, E^*) where $Z^* = V_A(I - S)V_A^T$ $L^* = U_A S U_A^T$



$$\min_{A,E} \|X\|_* + \gamma \|E\|_1 \quad s.t. \quad D = A + E$$

$$1) \quad X_{tr} \rightarrow A_{tr} \rightarrow U_{tr} \Sigma_{tr} V_{tr}^T$$

$$A_{ts} \rightarrow U_{tr} U_{tr}^T X_{ts}$$

$$2) \quad Z_{tr} \quad Z_{ts} \quad W$$

$$\min_{Z,L} \|Z\|_* + \|L\|_* \quad s.t. \quad X = XZ + LX$$

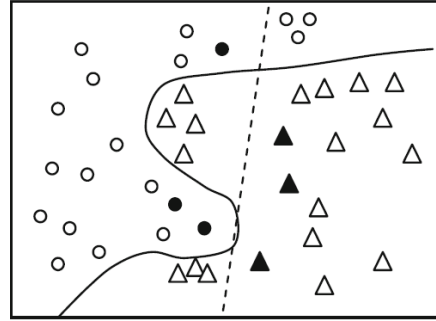


$$\min_{L,W} \sum_{i=1}^m \varphi(h_i, f(Lx_i, W)) + \alpha \|W\|_F^2$$

$$s.t. \quad L = U_X S U_X^T$$

Active Learning

Cold start ?



Limited budget



Active Learning via
Transductive Experimental Design

The goal of experimental design is to **find a set of experiments x_i that together are maximally informative.**

$$X: [x_1, \dots, x_m]^T \in \mathbb{R}^{m \times d}, |X| = m \quad V: [v_1, \dots, v_n]^T \in \mathbb{R}^{n \times d}, |V| = n$$

$$\min_{\mathbf{w}} \left\{ J(\mathbf{w}) = \sum_{i=1}^m (\mathbf{w}^\top \mathbf{x}_i - y_i)^2 + \mu \|\mathbf{w}\|^2 \right\} \quad \mathbf{C}_{\mathbf{w}} = \left(\frac{\partial J(\mathbf{w})}{\partial \mathbf{w} \partial \mathbf{w}^\top} \right)^{-1} = (\mathbf{X}^\top \mathbf{X} + \mu \mathbf{I})^{-1}$$

$\mathbf{f} = [f(v_1), \dots, f(v_n)]$ be the function values on all the available data $\mathbf{V}$, the predictive error $\mathbf{f} - \hat{\mathbf{f}}$ has the covariance matrix $\sigma^2 \mathbf{C}_{\mathbf{f}}$



$$\begin{aligned} \mathbf{C}_f &= \mathbf{V}\mathbf{C}_w\mathbf{V}^\top = \mathbf{V}(\mathbf{X}^\top\mathbf{X} + \mu\mathbf{I})^{-1}\mathbf{V}^\top \\ &= \frac{1}{\mu} \left[\mathbf{V}\mathbf{V}^\top - \mathbf{V}\mathbf{X}^\top(\mathbf{X}\mathbf{X}^\top + \mu\mathbf{I})^{-1}\mathbf{X}\mathbf{V}^\top \right] \end{aligned}$$

$$\begin{aligned} &\max_{\mathbf{X}} \quad \text{Tr} \left[\mathbf{V}\mathbf{X}^\top(\mathbf{X}\mathbf{X}^\top + \mu\mathbf{I})^{-1}\mathbf{X}\mathbf{V}^\top \right] \\ &\text{subject to} \quad \mathbf{X} \subset \mathbf{V}, |\mathbf{X}| = m \end{aligned}$$

$$\max_{X_o} \text{Tr}[V_o X_o^T (X_o X_o^T + \mu I)^{-1}] X_o V_o^T \longrightarrow X_o = (LX)^T \quad V_o = (LV)^T \quad L = USU^T$$

$$V \in R^{d \times n} \quad X \in R^{d \times m}$$

$$\max_{X,L} \text{Tr}[V^T L^T L X (X^T L^T L X + \mu I)^{-1}] X^T L^T L V \quad s.t. \quad L = USU^T$$

$$= \max_{X,S} \text{Tr}[V^T US^2 U^T X (X^T US^2 U^T X + \mu I)^{-1}] X^T US^2 U^T V$$

$$X = VP \quad s.t. \quad X \in R^{m \times n} \quad \sum_i P_{ij} = 1 \quad \sum_j P_{ij} \leq 1 \quad \sum_{i,j} P_{ij} = m$$

$$\max_{P,S} \text{Tr}[V^T US^2 U^T V P (P^T V^T US^2 U^T V P + \mu I)^{-1}] P^T V^T US^2 U^T V \quad s.t. \quad P, S$$

$$\min_{\mathbf{X}, \mathbf{A}} \quad \sum_{i=1}^n \|\mathbf{v}_i - \mathbf{X}^\top \mathbf{a}_i\|^2 + \mu \|\mathbf{a}_i\|^2 \quad \text{subject to} \quad \mathbf{X} \subset \mathbf{V}, |\mathbf{X}| = m \quad \mathbf{A} = [\mathbf{a}_1, \dots, \mathbf{a}_n]^\top \in \mathbb{R}^{n \times m}$$

$$L(X, A) = \|V - AX\|_F^2 + \mu \text{Tr}(AA^T) = \text{Tr}[(V - AX)(V - AX)^T] + \mu \text{Tr}(AA^T)$$

$$= \text{Tr}[VV^T - AXV^T - VX^T A^T + AXX^T A^T + \mu AA^T]$$

$$= \text{Tr}[VV^T] - \text{Tr}[AXV^T + VX^T A^T - A(XX^T + \mu I)A^T]$$

$$\frac{\partial L(X, A)}{\partial A} = 0 \quad \longrightarrow \quad A^* = VX^T(XX^T + \mu I)^{-1}$$

$$\min \|\mathbf{V} - \mathbf{A}\mathbf{X}\|_F^2 + \mu \text{Tr}(\mathbf{A}\mathbf{A}^\top) = \text{Tr}(\mathbf{V}\mathbf{V}^\top) - \text{Tr} \left[\mathbf{V}\mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + \mu \mathbf{I})^{-1} \mathbf{X}\mathbf{V}^\top \right]$$



$$\max_{\mathbf{X}} \quad \text{Tr} \left[\mathbf{V}\mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + \mu \mathbf{I})^{-1} \mathbf{X}\mathbf{V}^\top \right]$$

$$\text{subject to} \quad \mathbf{X} \subset \mathbf{V}, |\mathbf{X}| = m$$

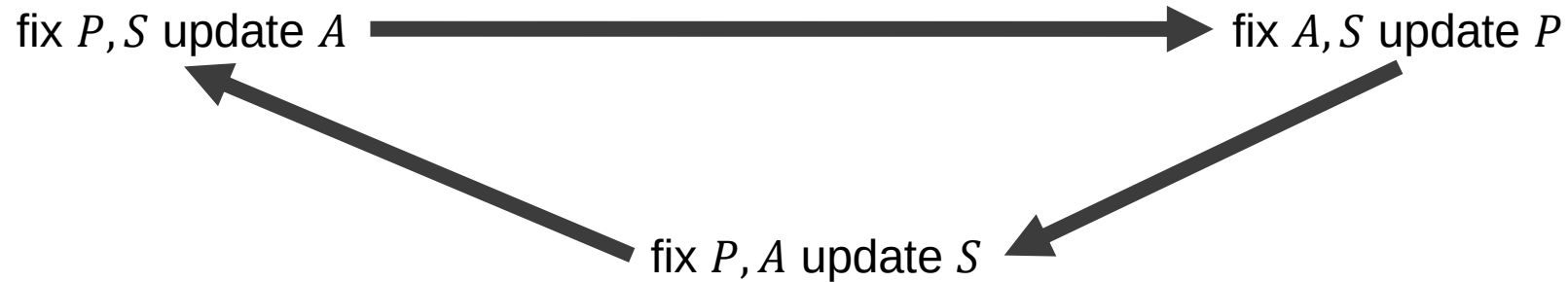


$$\min \|\mathbf{V} - \mathbf{A}\mathbf{X}\|_F^2 + \mu \text{Tr}(\mathbf{A}\mathbf{A}^\top) = \text{Tr}(\mathbf{V}\mathbf{V}^\top) - \text{Tr} \left[\mathbf{V}\mathbf{X}^\top (\mathbf{X}\mathbf{X}^\top + \mu \mathbf{I})^{-1} \mathbf{X}\mathbf{V}^\top \right]$$

$$\min_{A,P} \|V_o - AV_oP\|_F^2 + \mu \|A\|_F^2 = \min_{A,P} \|V^T L^T - AV^T L^T P\|_F^2 + \mu \|A\|_F^2, s.t. L = USU^T \quad V = U\Sigma V_v^T$$

$$= \min_{A,P,S} \|V_v \Sigma S U^T - AV \Sigma S U^T P\|_F^2 + \mu \|A\|_F^2$$

$$L = USU^T = U\Lambda U^T \quad \text{diag}(\Lambda) = (S_{11}, S_{22}, \dots, S_{rr}, 0, 0, \dots, 0) \in R^d \quad \Lambda \in R^{d \times d}$$





Early Active Learning via Robust Representation and Structured Sparsity

IJCAI 2013

The key idea of TED is to select the samples that can best represent the whole data using a linear representation.

$$\min_{A,B} \sum_{i=1}^n (\|x_i - Ba_i\|_2^2 + \gamma \|a_i\|_2^2) \quad s.t. A = [a_1, \dots, a_n] \in R^{m \times n}, B \subset X, |B| = m$$

The TED objective minimizes the **least square error**, hence it is **sensitive to the data outliers**. To solve these two deficiencies, we formulate the early active learning problem **using the structured sparsity-inducing norms** and propose a new robust formulation with efficient optimization algorithm.

$$\min_{A,B} \sum_{i=1}^n (\|x_i - Xa_i\|_2^2 + \gamma \|a_i\|_2^2) \quad s.t. A = [a_1, \dots, a_n] \in R^{n \times n} \quad \|A\|_{2,0}$$

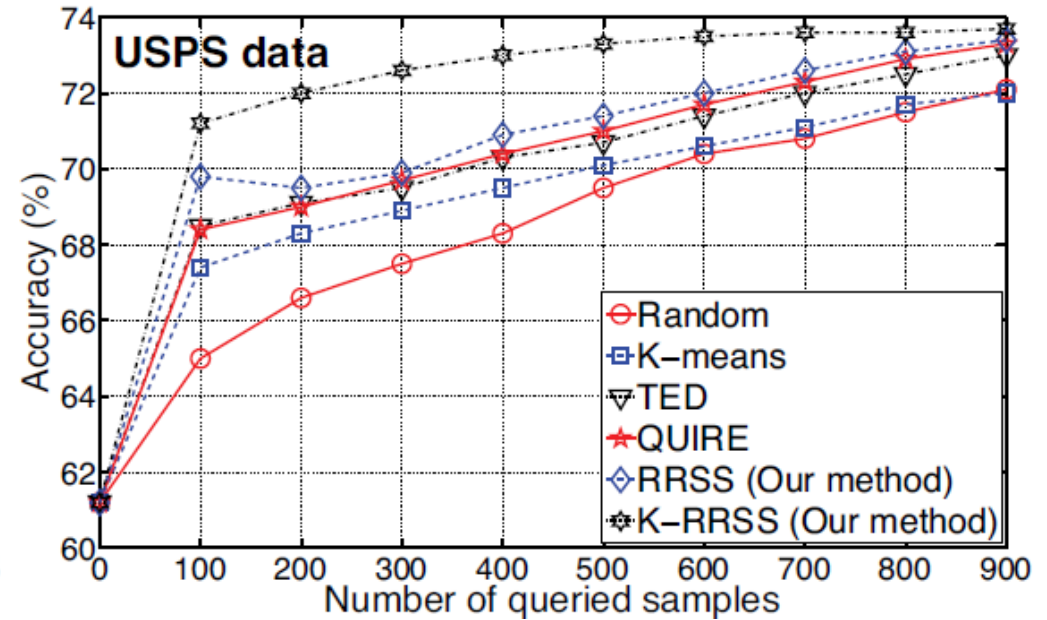
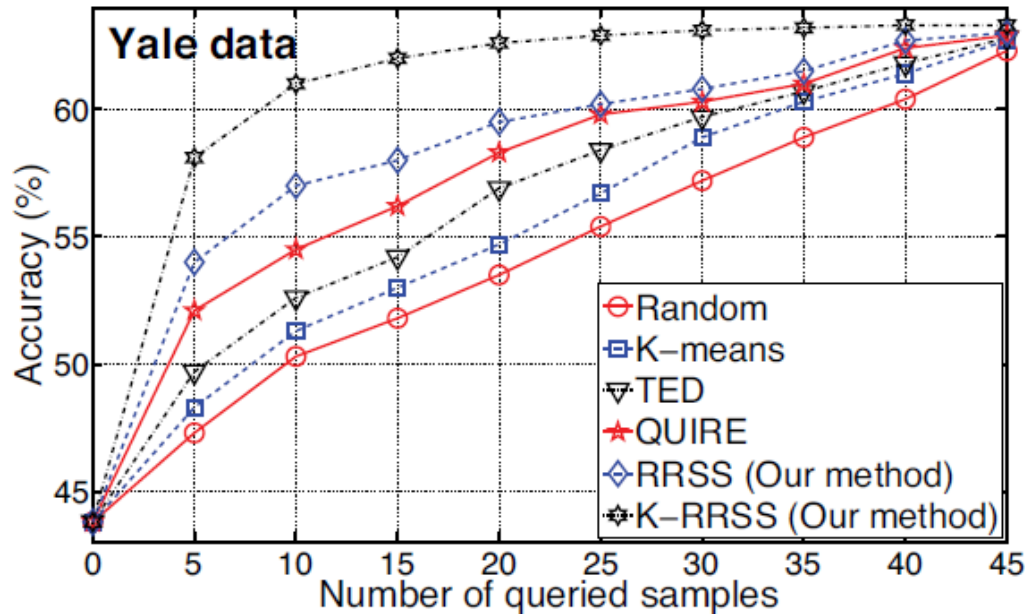
$$\min_A \sum_{i=1}^n (\|x_i - Xa_i\|_2^2 + \gamma \|A\|_{2,0})$$

$\|A\|_{2,1}$ is the minimum convex hull of $\|A\|_{2,0}$, and when A is row-sparse enough, one can always minimize $\|A\|_{2,1}$ to obtain the same result of minimizing $\|A\|_{2,0}$.

$$\min_A \sum_{i=1}^n (\|x_i - Xa_i\|_2^2 + \gamma \|A\|_{2,1}) = \min_A \|(X - XA)^T\|_2^2 + \gamma \|A\|_{2,1}$$

$$J = \min_A \|(X - XA)^T\|_{2,1} + \gamma \|A\|_{2,1}$$

L1-norm is imposed among data points and the L2-norm is used for features





IJCAI 2013

$$\min_A \sum_{i=1}^n (\|x_i - Xa_i\|_2^2 + \gamma \|A\|_{2,1}) = \min_A \|(X - XA)^T\|_2^2 + \gamma \|A\|_{2,1}$$

$$J = \min_A \|(X - XA)^T\|_{2,1} + \gamma \|A\|_{2,1}$$

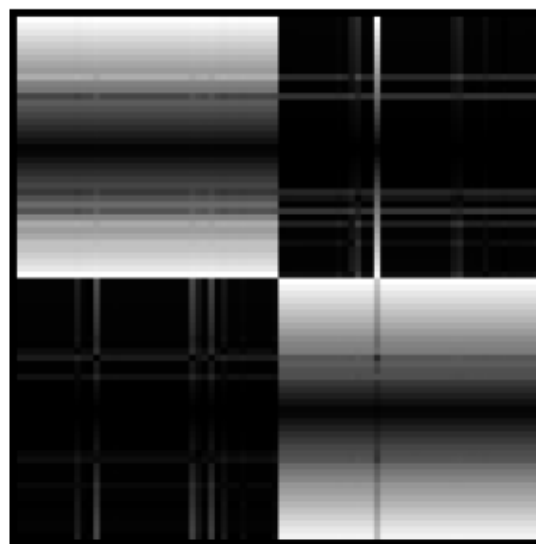
ICML 2010

$$\min_Z \|Z\|_* \quad \text{s.t. } X = XZ$$

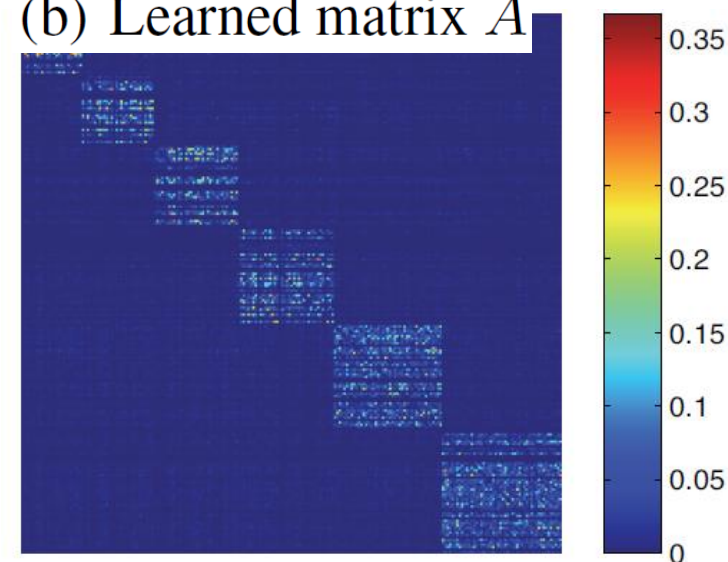
$$\min_{Z,E} \|Z\|_* + \lambda \|E\|_{2,1} \quad \text{s.t. } X = XZ + E$$

$$\min_Z \lambda \|X - XZ\|_{2,1} + \|Z\|_*$$

(b), Z^*



(b) Learned matrix A



SIGIR 08

$$\min_{\beta, \alpha_i \in \mathbb{R}^N} \sum_{i=1}^M \|\mathbf{x}_i - \mathbf{X}_C^\top \alpha_i\|^2 + \sum_{j=1}^N \frac{\alpha_{i,j}^2}{\beta_j} + \gamma \|\beta\|_1 \quad \text{subject to} \quad \mathbf{x}_i \in \mathbf{X}_{\mathcal{P}}, \quad \beta_j \geq 0, \quad j = 1, \dots, N$$



$$J = \min_A \|(X_o - X_o A)^T\|_{2,1} + \gamma \|A\|_{2,1} \quad X_o \in R^{d \times n} \quad A \in R^{n \times n}$$

$$X_o = LX \quad L = USU^T = \textcolor{red}{U}\Lambda\textcolor{red}{U}^T$$

$$\min_A \|(LX - LXA)^T\|_{2,1} + \gamma \|A\|_{2,1} = \min_{A,S} \|(US\Sigma V^T - US\Sigma V^T A)^T\|_{2,1} + \gamma \|A\|_{2,1}$$

$$\min_{A,P,S} \|V_v \Sigma S U^T - A V \Sigma S U^T P\|_F^2 + \mu \|A\|_F^2$$

1) $(\Sigma_X)_{ii} \neq (\Sigma_X)_{jj}$, when $S_{ij} = 0$. 2) S and $I - S$ are positive semidefinite

since only focus on learning the discriminative features, the constraint that $I - S$ is positive semidefinite is not necessary, so, only need to bound $S_{ii} > 0$.

Kernel Extension

$$\begin{aligned} J &= \min_A \|(\phi(X_o) - \phi(X_o)A)^T\|_{2,1} + \gamma \|A\|_{2,1} \\ &= \min_A \|(\phi(US\Sigma V^T) - \phi(US\Sigma V^T)A)^T\|_{2,1} + \gamma \|A\|_{2,1} \end{aligned}$$