

MIXTURE OF COMMON SUBSPACE VIA RIEMANNIAN GRADIENT: STATISTICAL AND COMPUTATIONAL LIMIT



Zhexu Jin, Joshua Agterberg

University of Illinois, Urbana Champaign

The Motivating Problem

Given L noisy symmetric low rank matrices (e.g. covariance matrices, networks), when can we cluster those matrices with the same subspace structure in a statistically consistent and computationally efficient manner?

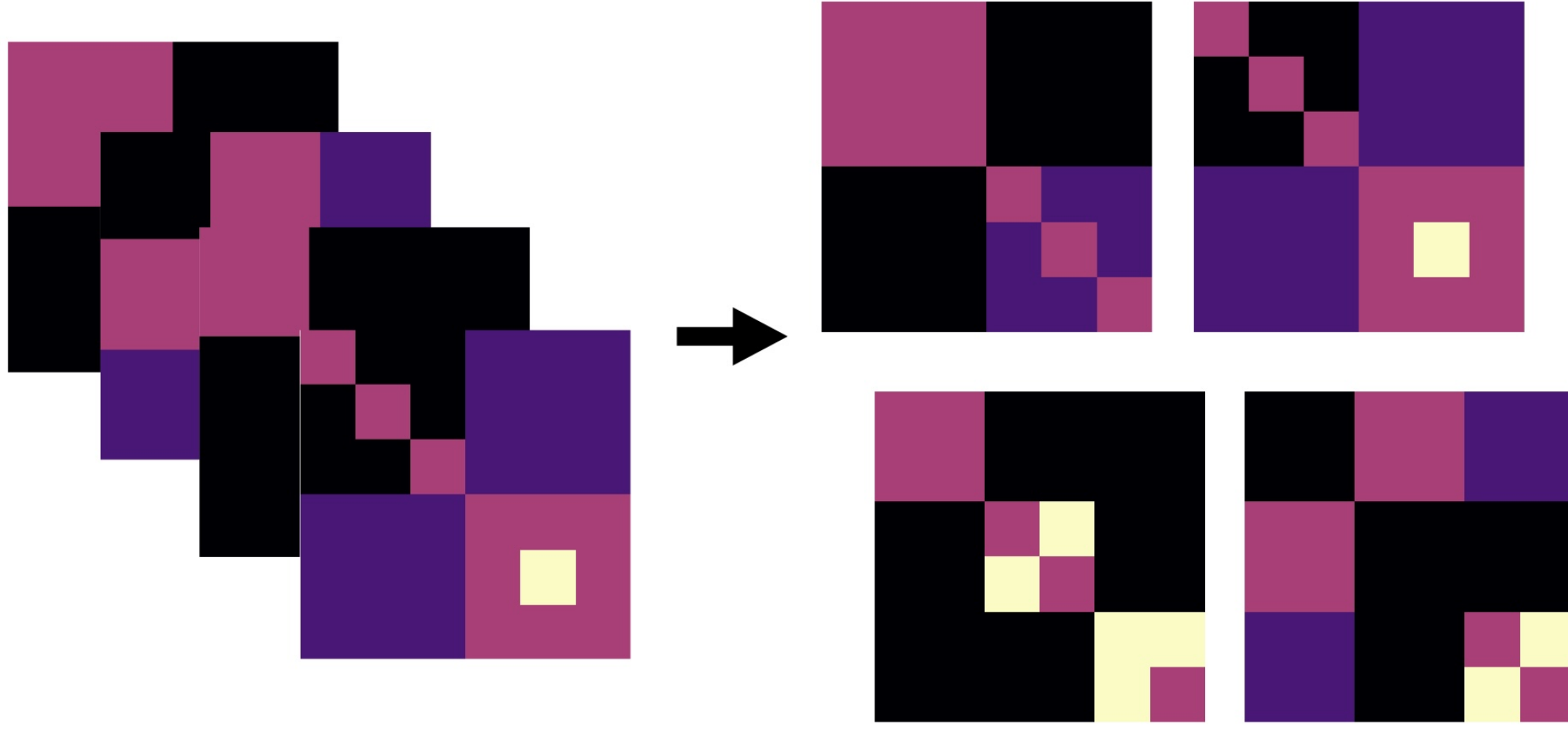


Fig. 1: Motivating Problem

Main Contributions

- Devise scalable initialization and riemannian gradient method for clustering mixture of common subspaces.
- Provide theoretical guarantee for clustering under a uniform signal to noise ratio condition across all layers.
- Derive computational and statistical lower bound for the uniform signal to noise ratio condition.

Mixture of Common Subspace Model

The model generalizes common subspace model [1]:

$$\mathbf{A}^{(l)} = \mathbf{S}^{(l)} + \mathbf{N}^{(l)} \quad l \in [L]$$

Here

$$\mathbf{S}^{(l)} = \mathbf{U}^{(m)} \mathbf{R}^{(l)} \mathbf{U}^{(m)\top} \quad m \in [M] \quad m := \tau(l)$$

- signal strength λ :

$$\lambda := \sigma_{\min}(\mathbf{S}^{(l)}) = \sigma_{\min}(\mathbf{R}^{(l)})$$

- noise level σ :

$$\mathbf{N}_{ij}^{(l)} \begin{cases} \sim \text{SG}(0, \sigma^2/2) & i > j \\ \sim \text{SG}(0, \sigma^2) & i = j \\ = \mathbf{N}_{ji}^{(l)} & i < j \end{cases}$$

- subspace separation:

$$\Delta := \min_{m \neq m'} \left\| \sin \Theta(\mathbf{U}^{(m)}, \mathbf{U}^{(m')}) \right\|$$

Initialization Algorithm

The unbiased estimator for $\mathbf{S}^{(l)2}$ is

$$\hat{\mathbf{S}}^{(l)2} := \mathbf{A}^{(l)2} - \frac{n+1}{2} \hat{\sigma}^2 \mathbf{I}_n$$

Algorithm 1 Clustering Initialization for Mixture of Common Subspace Model

Input: $\{\mathbf{A}^{(l)}\}_{l \in [L]}$, Matrix Cluster Number M

- 1: Compute $\hat{\mathbf{U}}$ as the top R eigenvectors of sum of squares $\mathbf{A} := \sum_{l \in [L]} \mathbf{A}^{(l)2}$, denoted as $\hat{\mathbf{U}}$.
- 2: Estimate $\hat{\sigma}$ as (18)
- 3: For each $l \in [L]$, compute $\hat{\mathbf{X}}^{(l)} = \hat{\mathbf{U}}^\top \hat{\mathbf{S}}^{(l)2} \hat{\mathbf{U}} / \left\| \hat{\mathbf{U}}^\top \hat{\mathbf{S}}^{(l)2} \hat{\mathbf{U}} \right\|_F$.
- 4: Run K-Means on $\{\hat{\mathbf{X}}^{(l)}\}_{l \in [L]}$, output the cluster membership $\hat{\tau}_0$.

Output: Estimated matrix cluster membership $\hat{\tau}_0$

Riemannian Gradient Algorithm

Given a fixed $\hat{\tau}$, we minimize $f(\mathbf{U}^{(m)})$ for each m via Riemannian Gradient Descent on Stiefel manifold:

$$f(\mathbf{U}^{(m)}) := \sum_{l \in \hat{\tau}^{-1}(m)} \left\| \mathbf{U}^{(m)} \mathbf{U}^{(m)\top} \mathbf{A}^{(l)} \mathbf{U}^{(m)} \mathbf{U}^{(m)\top} - \mathbf{A}^{(l)} \right\|_F^2$$

$$\text{grad } f(\mathbf{U}^{(m)}) := 4 \sum_{l \in \hat{\tau}^{-1}(m)} \left(\mathbf{U}^{(m)} \mathbf{U}^{(m)\top} - \mathbf{I} \right) \mathbf{A}^{(l)} \mathbf{U}^{(m)} \mathbf{U}^{(m)\top} \mathbf{A}^{(l)} \mathbf{U}^{(m)}$$

We optimize the nonconvex optimization above using an idea close to EM:

Algorithm 2 Riemannian Gradient Method for Subspace Refinement and Reclustering

Input: $\{\mathbf{A}^{(l)}\}_{l \in [L]}$, $\hat{\tau}_0, r_m$

- 1: **while** $\frac{1}{L} h_c(\hat{\tau}_k, \hat{\tau}_{k+1}) \geq \epsilon_1$ **and** $k \leq K$ **do**
- 2: **for** $m \in [M]$ **do**
- 3: compute $\hat{\mathbf{U}}_0^{(m)}$ as the top r_m eigenvectors of $\sum_{l \in \hat{\tau}^{-1}(m)} \mathbf{A}^{(l)2}$.
- 4: **while** $\left\| \frac{1}{L_m} \text{grad } f(\hat{\mathbf{U}}_t^{(m)}) \right\|_F \geq \epsilon_2$ **and** $t \leq T$ **do**
- 5: $\hat{\mathbf{U}}_{t+0.5}^{(m)} = \hat{\mathbf{U}}_t^{(m)} - \frac{\eta}{L} \text{grad } f(\hat{\mathbf{U}}_t^{(m)})$
- 6: $\hat{\mathbf{U}}_{t+1}^{(m)} = \mathbf{Q}_{0.5} \mathbf{R}_{0.5}^\top$ where $\mathbf{Q}_{0.5} \mathbf{\Sigma}_{0.5} \mathbf{R}_{0.5}^\top = \hat{\mathbf{U}}_{t+0.5}^{(m)}$ is the singular value decomposition of $\mathbf{U}_{0.5}^{(m)}$.
- 7: Update $\hat{\tau}_{k+1}$ such that $\hat{\tau}_{k+1}(l) = m$ iff $m = \arg \min_{m \in [M]} f_t(\hat{\mathbf{U}}_t^{(m)})$

Output: Estimated subspace $\{\hat{\mathbf{U}}_t^{(m)}\}_{m \in [M]}$, and cluster $\hat{\tau}_t$

Minimax Lower Bound

Define the loss function and parameter set

$$h_c(\hat{\tau}, \tau) := \inf_{\pi} \sum_{l \in [L]} \mathbb{I}\{\hat{\tau}(l) \neq \pi \circ \tau(l)\}. \quad \Theta(\lambda, n, L, r, M)$$

Theorem 1. If we have when $\frac{\lambda^2 \Delta^2}{\sigma^2} (\log(\frac{\alpha L \sigma}{M \lambda \Delta}))^{-1} \rightarrow \infty$ as $n \rightarrow \infty$, then it holds that

$$\inf_{\mathbf{z}} \sup_{\theta \in \Theta} \mathbb{E}_{\mathbb{P}_{\theta}} \frac{1}{L} h_c(\hat{\tau}, \tau^*) \geq \exp \left(- (1 + o(1)) \frac{\lambda^2 \Delta^2}{8 \sigma^2} \right).$$

Identifiability Condition

Quotient by rotation of $\mathbf{U}^{(m)}$, permutation of τ , and $\left\| \mathbf{U}^{(m)\top} \mathbf{U}^{(m')} \right\|_F^2 < r$.

Computational Lower Bound

Under the low degree likelihood ratio framework. Consider the following test

$$H_0^{(L)} : \mathbf{A}^{(l)} = \epsilon^{(l)} \lambda \mathbf{u} \mathbf{u}^\top + \mathbf{Z} \quad \text{vs} \quad H_1^{(L)} : \mathbf{A}^{(l)} = \epsilon^{(l)} \lambda \left(b^{(l)} \mathbf{u}_\delta \mathbf{u}_\delta^\top + (1 - b^{(l)}) \mathbf{u} \mathbf{u}^\top \right) \quad (1)$$

Theorem 2. For the hypothesis testing above, if $\lambda = o(\frac{n^{1/2}}{L^{1/4} \delta})$, then $\left\| \mathcal{L}_L^{\leq D} \right\|^2 = 1 + o(1)$. (Suggests exponential computational complexity under this regime)

Initialization Guarantee

Theorem 3. If dispersion condition is met and signal strength condition:

$$\frac{\lambda}{\sigma} \gtrsim \left(\frac{M^{1/4} \kappa r^{1/2}}{\alpha^{1/4}} \vee \frac{M^{3/8} \kappa^{1/2} (rR)^{1/4}}{\alpha^{3/8} s^{1/4}} \right) \frac{n^{1/2}}{L^{1/4} \Delta}$$

With probability at least $1 - \exp(-cn)$:

$$\frac{1}{L} h_c(\tau^*, \hat{\tau}) \lesssim \frac{\alpha \Delta^2}{M r \kappa^4} \lesssim 1$$

Simulation

For $M = 2$, we sample $\left\| \mathbf{U}_1^\top \mathbf{U}_2 \right\|_F^2 = 1.5$, $n = 200$, $r = 4$ and repeatedly cluster for 20 times

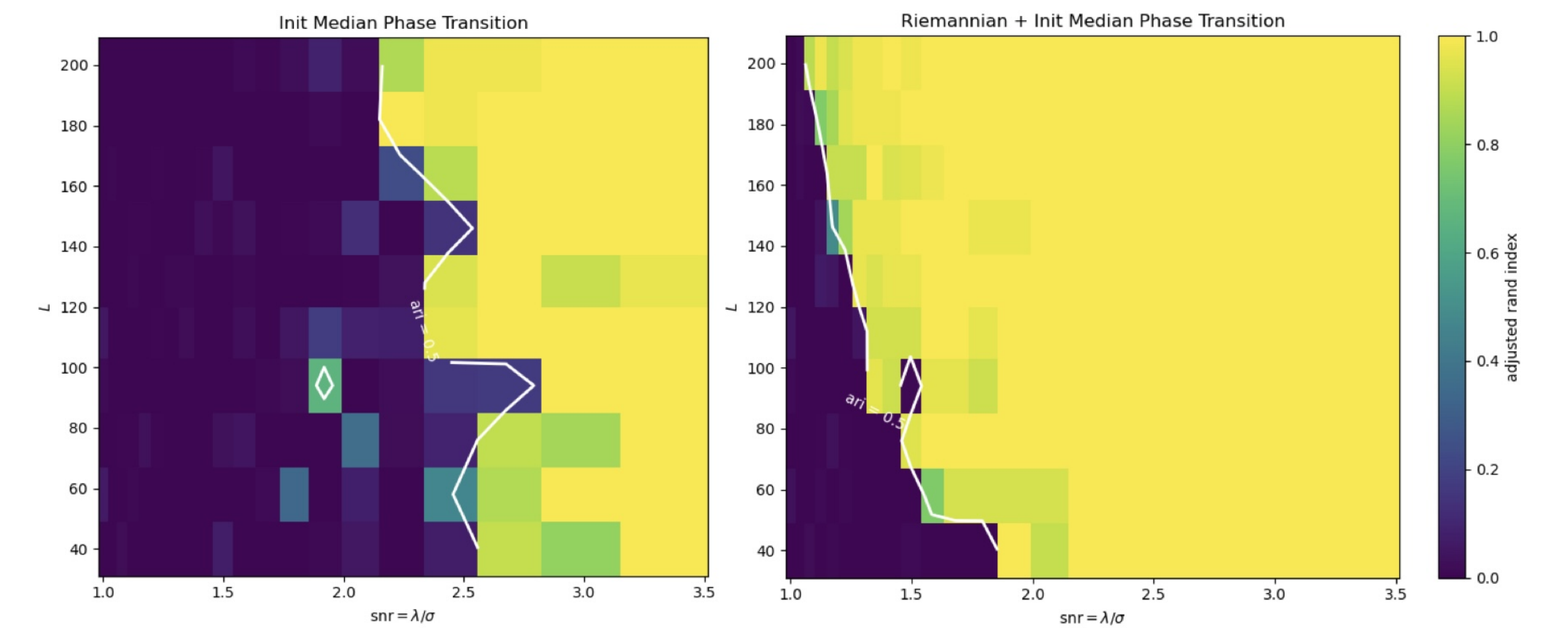


Fig. 2: Phase Transition for Initialization and Riemannian Gradient Step

Comparison on Real Data

We consider Brain Heart Lung (BHL) gene profile dataset as analyzed by [2]. We cluster the covariance matrix $\mathbf{S} = \mathbf{M}^\top \mathbf{M}$. Hence, $n = 1124$, $L = 27$ and $r = 1$.

	GMCS (Our)	Ir-Lloyd	DEEM	K-means	SKM	DTC	TBM	EM
Cluster Error	3.70	3.70	7.41	11.11	11.11	18.52	11.11	11.11

Table 1: GMCS: Gradient Mixture of Common Subspace; SKM: sparse K-means ; DTC: dynamic tensor clustering ; TBM: tensor block model (TBM) ; EM: standard EM;

[1] J. Arroyo, A. Athreya, J. Cape, G. Chen, C. E. Priebe, and J. T. Vogelstein, "Inference for multiple heterogeneous networks with a common invariant subspace," *Journal of Machine Learning Research*, vol. 22, no. 142, pp. 1–49, 2021.

[2] Z. Lyu and D. Xia, "Optimal clustering by Lloyd's algorithm for low-rank mixture model," *Journal of the Royal Statistical Society Series B: Statistical Methodology*, vol. 88, no. 1, pp. 43–65, 2026.