

Mixed Membership Graph Clustering via Systematic Edge Query

Shahana Ibrahim and Xiao Fu

School of Electrical Engineering and Computer Science
Oregon State University
Corvallis, OR 97331, USA

July 13, 2021

Abstract

This work considers clustering nodes of a largely incomplete graph. Under the problem setting, only a small amount of queries about the edges can be made, but the entire graph is not observable. This problem finds applications in large-scale data clustering using limited annotations, community detection under restricted survey resources, and graph topology inference under hidden/removed node interactions. Prior works tackled this problem from various perspectives, e.g., convex programming-based low-rank matrix completion and active query-based clique finding. Nonetheless, many existing methods are designed for estimating the single-cluster membership of the nodes, but nodes may often have mixed (i.e., multi-cluster) membership in practice. Some query and computational paradigms, e.g., the random query patterns and nuclear norm-based optimization advocated in the convex approaches, may give rise to scalability and implementation challenges. This work aims at learning mixed membership of nodes using queried edges. The proposed method is developed together with a systematic query principle that can be controlled and adjusted by the system designers to accommodate implementation challenges—e.g., to avoid querying edges that are physically hard to acquire. Our framework also features a lightweight and scalable algorithm with membership learning guarantees. Real-data experiments on crowdclustering and community detection are used to showcase the effectiveness of our method.

1 Introduction

Graph clustering (GC) is of broad interest in data science, which aims at associating the nodes of a graph with different clusters in an unsupervised manner [1]. The GC technique is well-motivated, since network data frequently arise in various applications (e.g., in social network analysis, brain signal processing, and biological/ecological data mining).

In the past two decades, much effort was invested onto dealing with the GC problem; see, e.g., [2–6]. Nonetheless, some major challenges arise in the era of big data. Notably, many graph data are highly incomplete for various reasons. For example, a large social network follower-followee data could contain billions of nodes, which translates to $\approx 10^{18}$ edges. Data acquisition at such a scale is a highly nontrivial task. In many cases, instead of collecting edge information of the entire network, data analysts have to *sample* edges of interest, and use the sampled network to perform GC [7, 8]. Graph sampling (or, edge query)-based network analysis is also well-motivated for other reasons—e.g., in community detection of networks where edges are intentionally removed or hidden (e.g., terrorist networks or radical group networks) [9] and in biological/ecological networks where acquiring the complete edge information is too resource-consuming [10, 11].

A number of works considered GC under incomplete edge observation. For example, the works in [12–14] used uniformly random edge queries and proposed convex low-rank matrix completion-based algorithms for GC. Another line of work [15–17] proposed active edge query-based methods in which the learning algorithms query the edges on-the-fly, in order to reduce the number of queries needed.

The approaches in [12–17] are insightful, and showed that edge query-based GC is a largely viable task, in both theory and practice. Nonetheless, a number of challenges remain. First, the methods in [12–17] were not designed to estimate multiple membership of the nodes. However, each node is often associated with different clusters in real-world networks (e.g., a person in a co-author network could belong to the machine learning and statistics communities simultaneously). Second, the theoretical guarantees in [12–17] were also established under single membership models—and oftentimes based on random query strategies (see, e.g., [12–14]). Note that in certain applications (e.g., field surveys based graph/network analysis [18]), systematic edge query strategies may be more preferred. Third, the existing methods are not always scalable. In particular, the convex optimization approaches in [12–14] recast the edge query-based graph clustering problem as a nuclear norm minimization problem, which entails N^2 (where N is the number of nodes) optimization variables—making it hard to handle real-world large-scale graphs.

Contributions. In this work, we offer an alternative framework for learning the node membership from an incomplete graph. Some notable features of our framework are as follows:

- **Systematic Edge Query Strategy.** Our algorithm design comes with a *systematic* edge query scheme, under which the node membership can be provably learned. As discussed, systematic query may be more preferable than random query strategies in a variety of applications.
- **Guaranteed Mixed Membership Identification.** Unlike existing provable edge-query based graph clustering methods in [12–17] whose goals are to learn single membership of nodes belonging to disjoint clusters, we model the undirected adjacency graphs using a mixed membership model that allows nodes to be associated with multiple overlapping clusters. Our model is reminiscent of the *mixed membership stochastic block* (MMSB) model in overlapped community detection [2]. Accordingly, we offer mixed membership identification guarantees using queried edges. Another contribution is that we derive a new sufficient condition for mixed membership identification. Compared to existing MMSB identifiability conditions, e.g., those in [19–22], our new condition is more closely tied to key characterizations of the GC problem, e.g., the Dirichlet prior of node membership, the graph size, and the cluster-cluster interaction intensity.
- **Scalable and Lightweight Algorithm.** We propose a simple procedure that only consists of the truncated *singular value decomposition* (SVD) of small matrices and a Gram–Schmidt-type greedy algorithm for *simplex-structured matrix factorization* (SSMF). This procedure is computationally very economical, especially when compared to the convex optimization approaches in [12–14, 23]. We also provide sample complexity and noise robustness analyses for the proposed framework.
- **Evaluation on Real-World Datasets.** We conduct a series of experiments on real-world datasets. First, we consider the query-based crowdclustering task [12, 13], using annotator-labeled graph data to assist image clustering. The data used in our evaluation was uploaded to the *Amazon Mechanical Turk* (AMT) platform and labeled by unknown human annotators. The acquired AMT data has been made publicly available¹. Second, we also use co-authorship network datasets, *Digital Bibliography & Library Project* (DBLP) and *Microsoft Academic Graph* (MAG), to evaluate our method on community detection tasks.

An initial and limited version of this work was published in the proceedings of IEEE ICASSP 2021 [24]. In this journal version, we additionally include (1) detailed identifiability analysis of the proposed algorithm (Theorem 1 and its proof); (2) the integrated performance characterization of the SVD and SSMF stages (Theorem 2 and its proof); (3) a new application, namely, crowdclustering; (4) a series of new real-data experiments; (5) and a newly acquired dataset for crowdclustering that is made publicly available.

Notation. $x, \mathbf{x}, \mathbf{X}$ denote a scalar, vector and matrix, respectively. $\kappa(\mathbf{X})$ denotes the condition number of the matrix $\mathbf{X}$ and is given by $\kappa(\mathbf{X}) = \frac{\sigma_{\max}(\mathbf{X})}{\sigma_{\min}(\mathbf{X})}$ where $\sigma_{\max}$ and $\sigma_{\min}$ are the largest and smallest nonzero singular values of $\mathbf{X}$, respectively. $\mathbf{X} \geq \mathbf{0}$ denotes that all the elements of $\mathbf{X}$ are nonnegative. $\|\mathbf{X}\|_2$ represents the 2-norm of matrix $\mathbf{X}$, i.e., $\|\mathbf{X}\|_2 = \sigma_{\max}(\mathbf{X})$. $\|\mathbf{X}\|_F$ denotes the Frobenius norm of $\mathbf{X}$. $\|\mathbf{X}\|_*$

¹<https://github.com/shahanaibrahimosu/mixed-membership-graph-clustering>

denotes the nuclear norm of $\mathbf{X}$. $\text{svd}_K(\mathbf{X})$ denotes the top- K singular value decomposition (SVD) of $\mathbf{X}$. $\text{range}(\mathbf{X})$ denotes the subspace spanned by the columns of $\mathbf{X}$. $\mathbf{X}(i, j)$ and $x_{i,j}$ both denote the element of $\mathbf{X}$ from its i th row and j th column. $\mathbf{X}(:, j)$ and $\mathbf{x}_j$ both denote the j th column of $\mathbf{X}$. $\|\mathbf{x}\|_2$ and $\|\mathbf{x}\|_1$ denote ℓ_2 and ℓ_1 norms of vector $\mathbf{x}$, respectively. $^\top$ and † denote transpose and pseudo-inverse, respectively. $|\mathcal{C}|$ denotes the cardinality of set $\mathcal{C}$. $\cup$ denotes the union operator for sets. $[T]$ represents the set of positive integers up to T , i.e., $[T] = \{1, \dots, T\}$. $\text{Diag}(\mathbf{x})$ is a diagonal matrix that holds the entries of $\mathbf{x}$ as its diagonal elements.

2 Problem Statement

Consider N data entities (e.g., persons in a social network or image data in a sample-sample similarity network) that are from K clusters. We allow a node to belong to multiple clusters; i.e., we consider the case where the clusters have overlaps and the nodes admits *mixed membership* [2]. Assume that the n th entity belongs to cluster k with probability $m_{k,n}$, where

$$\sum_{k=1}^K m_{k,n} = 1, m_{k,n} \geq 0. \quad (1)$$

Then, the vector $\mathbf{m}_n = [m_{1,n}, \dots, m_{K,n}]^\top$ is referred to as the *membership vector* of data entity n , since it reflects the association of entity n with different clusters. All such vectors together constitute the *membership matrix* $\mathbf{M} = [\mathbf{m}_1, \dots, \mathbf{m}_N] \in \mathbb{R}^{K \times N}$.

In GC, the entities are the nodes of the graph, and the relationship between the nodes are represented as the edges of the graph. In this work, we consider undirected graphs that are represented as symmetric adjacency matrices, i.e., $\mathbf{A} \in \{0, 1\}^{N \times N}$, where each entry $\mathbf{A}(i, j)$ encodes the pairwise relationship between nodes i and j in terms of binary values (i.e., $\mathbf{A}(i, j) \in \{0, 1\}$). Our goal is to learn the membership vector $\mathbf{m}_n$ for each node from limited edge queries about the adjacency matrix $\mathbf{A}$ (i.e., the number of edges queried is much smaller than total number of edges, i.e., $N(N-1)/2$).

2.1 Motivating Examples

Our problem setting is motivated by a couple of important applications, namely, the crowdclustering problem and incomplete graph based overlapping community detection.

- **Example 1: Crowdclustering.** *Crowdclustering* is a technique that clusters data samples with the assistance from crowdsourced annotators [5, 12, 13, 23, 25, 26] [15–17]. In crowdclustering, the data entities are presented to the annotators in pairs. The annotators determine if the pair of entities are from the same category or not; i.e., the annotators apply a rule such that $\mathbf{A}(i, j) = 1$ if entities i, j ($i < j$) are believed to be from the same cluster, and $\mathbf{A}(i, j) = 0$ otherwise. Instead of asking the annotators to determine the exact membership of the n th data entity (i.e., the $\mathbf{m}_n$ vector), the above mentioned rule only asks the annotators to output binary labels (i.e., if two entities are similar or not). This annotation paradigm is arguably more accurate than exact membership labeling, since it can work under much less expertise.

Crowdclustering amounts to learning $\mathbf{m}_n$ from $\mathbf{A}$, which is a GC problem. The challenge is that annotating the full adjacency graph requires $O(N^2)$ pairs to be compared and labeled—a heavy workload if large data sets are considered. Hence, edge query-based GC techniques are natural to handle the crowdclustering task. This way, only a subset of the $O(N^2)$ sample pairs needs to be compared and annotated.

- **Example 2: Edge Sampling-Based Overlapping Community Detection.** The mixed membership modeling in (1) and the learning problem of interest can also be applied to *overlapping community detection* (OCD) [27]—since the mixed membership setting implies that the communities have overlaps. Classic OCD algorithms work with statistical generative models, e.g., the MMSB [2, 22] and the *Bayesian nonnegative matrix factorization model* [28]. The mixed membership identifiability guarantees were established under full observation of $\mathbf{A}$ in [19–22]. Mixed membership identification for OCD under sampled edges is of great

interest for applications like field survey based community analysis [18] or community detection in link-removed/hidden networks (e.g., terrorist networks) [9]. In both cases, one may not be able to observe the entire $\mathbf{A}$ due to reasons such as resource limitations and difficulty of edge acquisition. However, theoretical guarantees and provable algorithms for this problem have been elusive.

2.2 Prior Work

A theoretical challenge associated with our problem is as follows: Is it possible to learn the membership vector $\mathbf{m}_n$ from the binary matrix $\mathbf{A}$? This gives rise to a membership *identifiability* problem. If $\mathbf{A}$ is fully observed, answering the identifiability question amounts to understanding theoretical guarantees for similarity graph-based membership learning techniques, e.g., MMSB identification and spectral clustering—which has been thoroughly studied in the machine learning community; see, e.g., [19–22, 29]. However, when $\mathbf{A}$ is only partially observed, the identifiability problem has not been fully understood.

To handle the GC problem with partial observations, the work in [12, 13] models the generating process of $\mathbf{A}$ using the SBM followed by a random edge query stage for data acquisition. The SBM can be summarized as follows. Assume that $\mathbf{B} \in \mathbb{R}^{K \times K}$ represents a cluster-cluster similarity matrix which is symmetric and $B(p, q)$ represents the probability that cluster p is connected with cluster q . In addition, assume that the nodes are connected through their membership identities. Then, the probability that $\mathbf{A}(i, j) = 1$ is $P(i, j) = \mathbf{m}_i^\top \mathbf{B} \mathbf{m}_j$, i.e., $\mathbf{A}(i, j) \sim \text{Bernoulli}(\mathbf{m}_i^\top \mathbf{B} \mathbf{m}_j)$ where $i < j$ —the adjacency matrix $\mathbf{A}$ is sampled from Bernoulli distributions specified by the entries of the matrix $\mathbf{P} = \mathbf{M}^\top \mathbf{B} \mathbf{M}$. Note that under SBM, the membership vector $\mathbf{m}_i$ is always a unit vector; i.e., only single membership and disjoint clusters are considered under SBM.

To recover $\mathbf{M}$ from partially observed $\mathbf{A}$ under SBM, the work in [12–14, 30] use the following convex program and its variants:

$$\begin{aligned} & \underset{\mathbf{L} \in [0,1]^{N \times N}, \mathbf{S}}{\text{minimize}} \quad \|\mathbf{L}\|_* + \lambda \|\mathbf{S}\|_1 \\ & \text{subject to} \quad \mathbf{L}(i, j) + \mathbf{S}(i, j) = \mathbf{A}(i, j), \quad \forall (i, j) \in \Omega. \end{aligned} \quad (2)$$

In the above, $\mathbf{L}$ is used to model the low-rank component of $\mathbf{A}$, $\mathbf{S}$ is a sparse error matrix, and Ω denotes the set of pairs (i, j) such that the corresponding $\mathbf{A}(i, j)$'s are observed. Notably, it was shown in [12–14] that if $\mathbf{A}$ follows the SBM, then the above convex program (and its variants) recovers $\mathbf{A}$ up to certain errors with provable guarantees using *random* edge queries. The results from [12–14] have opened many doors for edge query-based graph clustering, but several caveats exist. In particular, the recoverability is established based on single membership models, but the mixed membership case is of more interest. In addition, the optimization problem involves $O(N^2)$ variables (memory $\approx 3,000\text{GB}$ when $N = 10^6$), and thus is hardly scalable. Also, the provable guarantees in [12–14] rely on random edge query, which may not be easy to implement under some scenarios, as we discussed.

Query-based GC was also studied in another line of work [15–17]. There, the focus is developing active/online edge query-based methods that can discover the clusters on-the-fly, while attaining certain information-theoretic query lower bounds. Unlike the methods in [12–14] that take a latent model identification perspective, the algorithms in [15–17] are more from a successive clique (i.e., fully connected subgraph) finding perspective (although some connections to the SBM were also made). Similar to the convex approaches, the methods in [15–17] are designed to learn single membership.

3 Proposed Approach

In this work, we relax the SBM assumption on $\mathbf{P} = \mathbf{M}^\top \mathbf{B} \mathbf{M}$ by allowing the nodes to have mixed membership, as advocated in [2, 22, 31]. Hence, the constraints on $\mathbf{M}$ becomes

$$\mathbf{1}^\top \mathbf{M} = \mathbf{1}^\top, \quad \mathbf{M} \geq \mathbf{0}; \quad (3)$$

i.e., $\mathbf{m}_n$ resides in the probability simplex, instead of being the vertices of the simplex as in the SBM. Under the mixed membership assumption in (3), the following Bernoulli model, i.e.,

$$\begin{cases} \mathbf{A}(i, j) \sim \text{Bernoulli}(\mathbf{m}_i^\top \mathbf{B} \mathbf{m}_j), & i < j, \\ \mathbf{A}(i, j) = \mathbf{A}(j, i), & i > j, \\ \mathbf{A}(i, i) = 0, & i \in [N], \end{cases} \quad (4)$$

is adopted in our generative model for the complete adjacency matrix. Note that $\mathbf{m}_i^\top \mathbf{B} \mathbf{m}_j$ physically corresponds to the probability that nodes i and j admit an edge (i.e., $\Pr(\mathbf{A}(i, j) = 1) = \mathbf{m}_i^\top \mathbf{B} \mathbf{m}_j$), and thus this Bernoulli model makes intuitive sense; see more discussions in [2, 22, 31].

3.1 Systematic Edge Query

Our goal is to learn $\mathbf{M}$ from controlled edge sampling strategies. To proceed, we first divide the nodes into L disjoint groups $\mathcal{S}_1, \mathcal{S}_2, \dots, \mathcal{S}_L$ such that $\mathcal{S}_1 \cup \dots \cup \mathcal{S}_L = [N]$. Let $\mathbf{A}_{\ell, m} \in \mathbb{R}^{|\mathcal{S}_\ell| \times |\mathcal{S}_m|}$ denote the adjacency matrix between groups of nodes indexed by $\mathcal{S}_\ell$ and $\mathcal{S}_m$. We propose an edge query principle as described in the following:

Edge Query Principle (EQP). Query the blocks of edges such that the sampled $\mathbf{A}_{\ell, m}$'s satisfy the following two conditions:

- For every $\ell \in [L]$, $K \leq |\mathcal{S}_\ell|$ holds.
- Let $m_r \in [L]$ and $\{\ell_r\}_{r=1}^L = [L]$. For every ℓ_r , there exists a pair of indices m_r and ℓ_{r+1} where $\ell_{r+1} \neq \ell_r$ such that the edges from the blocks $\mathbf{A}_{\ell_r, m_r}$ and $\mathbf{A}_{\ell_{r+1}, m_r}$ are queried.

Note that since $\mathbf{A}$ is symmetric, when we select a block $\mathbf{A}_{\ell, m}$ to be queried, it also covers $\mathbf{A}_{m, \ell}$ since $\mathbf{A}_{m, \ell} = \mathbf{A}_{\ell, m}^\top$. A couple of remarks are in order:

Remark 1 The proposed EQP covers a large variety of query ‘masks’—as shown in Fig. 1. Since the query pattern can be by design instead of random, this entails the flexibility to avoid querying edges that are known a priori hard to acquire, e.g., edges that may have been intentionally removed to conceal information or edges that correspond to interactions between groups that are hard to survey due to various reasons, e.g., long geographical distance.

Remark 2 Instead of sampling individual edges, we sample blocks of edges under the proposed EQP. As one will see, this simple block query pattern allows us to design a provable and lightweight algorithm for mixed membership learning. The proposed query strategy can be easily implemented when the query patterns are controlled by the network analysts. For example, in crowdclustering, one can dispatch blocks of sample pairs for similarity annotation. In community analysis, field surveys can be designed using EQP such that the blocks which are easier to be queried are selected.

3.2 Connections to Matrix Completion

From a matrix completion viewpoint, our EQP is most closely related to those in [32, 33] that consider recoverable sampling patterns of matrix completion (MC). The work in [32] relates the recoverability of a low-rank matrix to the structural properties of a bipartite graph by considering the rows and columns as nodes. It establishes the “finite” (not unique) completability of the matrix based on combinatorial conditions on the observed edges in its subgraphs. In a similar spirit, the work in [33] shows that the range space of a low-rank matrix can be uniquely identified under the existence of certain patterns formed by the sampled entries.

From an MC perspective, our block sampling pattern can be considered as special cases covered by those in [32, 33]—with important distinguishing features. First, the work in [32, 33] established recoverability under the sampling pattern, but had no tractable algorithms. Our EQP features a polynomial-time

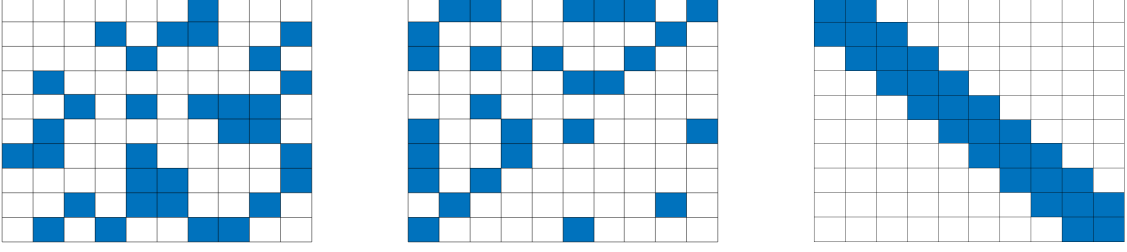


Figure 1: Some example patterns for A following EQP with $N = 1000$, $K = 5$ and $L = 10$. The shaded blue region represents the blocks queried.

lightweight algorithm via exploiting the special block sampling pattern. Second, the recoverability results in [32, 33] only cover sampled continuous real-valued low-rank matrices without noise. Our approach can provably estimate the range space of binary matrices (or, highly noisy low-rank matrices) under reasonable conditions.

3.3 Algorithm Design

In this section, we propose an algorithm under the EQP. Specifically, under the considered model in (4), we develop an algorithm that consists of simple SVD operations to estimate $\text{range}(M^\top)$ and a subsequent SSF stage to ‘extract’ M from the estimated range space.

3.3.1 Main Idea – A Toy Example

To shed some light on how our algorithm approximately identifies $\text{range}(M^\top)$, let us consider the ideal case where $A_{\ell,m} = P_{\ell,m} = M_\ell^\top B M_m$. We start by analyzing a toy example with $L = 3$ and $K \leq N/3$; see Fig. 2. We assume that the blocks $P_{1,2} \in \mathbb{R}^{N/3 \times N/3}$, $P_{2,2} \in \mathbb{R}^{N/3 \times N/3}$, and $P_{3,1} \in \mathbb{R}^{N/3 \times N/3}$ are queried. Hence, by symmetry, the following blocks are known:

$$P_{1,2} = M_1^\top B M_2, P_{2,2} = M_2^\top B M_2, \quad (5)$$

$$P_{2,1} = M_2^\top B M_1, P_{3,1} = M_3^\top B M_1, \quad (6)$$

Define $C_1 := [P_{1,2}^\top, P_{2,2}^\top]^\top$ and $C_2 := [P_{2,1}^\top, P_{3,1}^\top]^\top$. The top- K SVD of C_1 and C_2 can be represented as follows:

$$C_1 = [U_1^\top, U_2^\top]^\top \Sigma V_2^\top, C_2 = [\bar{U}_2^\top, \bar{U}_3^\top]^\top \bar{\Sigma} \bar{V}_1^\top. \quad (7)$$

Combining (5)-(7), and under the assumption that $\text{rank}(M_\ell) = \text{rank}(B) = K$ and $K \leq |S_\ell|$, for all ℓ , one can express the bases of $\text{range}(M_1^\top)$, $\text{range}(M_2^\top)$ and $\text{range}(M_3^\top)$ as $U_1 = M_1^\top G$, $U_2 = M_2^\top G$, $\bar{U}_3 = M_3^\top \bar{G}$, respectively, where $G \in \mathbb{R}^{K \times K}$ and $\bar{G} \in \mathbb{R}^{K \times K}$ are certain nonsingular matrices. Our hope is to “stitch” the bases above to have

$$\begin{aligned} \text{range}(U) &= \text{range}([U_1^\top, U_2^\top, U_3^\top]^\top) \\ &= \text{range}([M_1, M_2, M_3]^\top), \end{aligned} \quad (8)$$

with U_1 and U_2 in (7) and a certain U_3 . Note that $\bar{U}_3$ cannot be directly combined with U_1 and U_2 to attain the above, since $G = \bar{G}$ does not generally hold. To fix this, we define the following operation: $U_3 := \bar{U}_3 \bar{U}_2^\dagger U_2$. It is not hard to see that

$$\bar{U}_3 \bar{U}_2^\dagger U_2 = M_3^\top \bar{G} \times (M_2^\top \bar{G})^\dagger \times M_2^\top G = M_3^\top G,$$

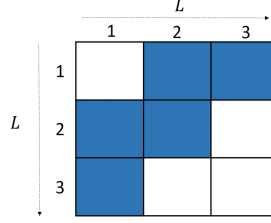


Figure 2: An illustrative case for subspace identifiability analysis. The shaded blue region represents the blocks queried.

Algorithm 1 Proposed Algorithm: BeQuec

Require: $L, K, \{\mathbf{A}_{\ell_r, m_r}\}_{r=1}^L, \{\mathbf{A}_{\ell_{r+1}, m_r}\}_{r=1}^{L-1}$;

- 1: $T \leftarrow \lfloor L/2 \rfloor$;
- 2: $\mathbf{C}_T \leftarrow [\mathbf{A}_{\ell_T, m_T}^\top, \mathbf{A}_{\ell_{T+1}, m_T}^\top]^\top$;
- 3: $[\mathbf{U}_{\ell_T}^\top, \mathbf{U}_{\ell_{T+1}}^\top]^\top \Sigma_T \mathbf{V}_{m_T}^\top \leftarrow \text{svd}_K(\mathbf{C}_T)$;
- 4: $[\mathbf{U}_{\text{ref2}}^\top, \mathbf{U}_{\text{ref1}}^\top]^\top \leftarrow [\mathbf{U}_{\ell_T}^\top, \mathbf{U}_{\ell_{T+1}}^\top]^\top$;
- 5: **for each** $r = T + 1 : L - 1$ **do** ▷ Init. temp. variables
- 6: $\mathbf{U}_{\ell_{r+1}} \leftarrow \text{PairStitch}(\mathbf{A}_{\ell_r, m_r}, \mathbf{A}_{\ell_{r+1}, m_r}, \mathbf{U}_{\text{ref1}})$ ▷ Est. $\{\mathbf{U}_{\ell_r}\}_{r=T+2}^L$
- 7: $\mathbf{U}_{\text{ref1}} \leftarrow \mathbf{U}_{\ell_{r+1}}$;
- 8: **end for each**
- 9: **for each** $r = T - 1 : 2$ **do** ▷ Est. $\{\mathbf{U}_{\ell_r}\}_{r=1}^{T-1}$
- 10: $\mathbf{U}_{\ell_{r-1}} \leftarrow \text{PairStitch}(\mathbf{A}_{\ell_{r-1}, m_{r-1}}, \mathbf{A}_{\ell_r, m_{r-1}}, \mathbf{U}_{\text{ref2}})$
- 11: $\mathbf{U}_{\text{ref2}} \leftarrow \mathbf{U}_{\ell_{r-1}}$;
- 12: **end for each**
- 13: $\hat{\mathbf{U}} \leftarrow [\mathbf{U}_1^\top, \dots, \mathbf{U}_L^\top]^\top$;
- 14: $\hat{\mathbf{G}} \leftarrow \text{SPA}(\hat{\mathbf{U}}, K)$ [35]; ▷ Est. $\mathbf{G}$ under (9)
- 15: $\hat{\mathbf{M}} \leftarrow \hat{\mathbf{G}}^{-\top} \hat{\mathbf{U}}^\top$ (or constrained least squares);
- 16: **return** $\hat{\mathbf{M}}$

which leads to (8). The idea conveyed by this simple example can be recursively applied to cover a general L block case under the proposed EQP, and the range space estimation accuracy can be guaranteed even when $\mathbf{A}$ is a noisy (binary) version of $\mathbf{P}$ —which will be detailed in the next part. After $\mathbf{U}^\top = \mathbf{G}^\top \mathbf{M}$ is obtained, extracting $\mathbf{M}$ from $\mathbf{U}^\top$ is the so-called SSMF problem, which has been extensively studied [34]; see a brief discussion in Sec. 3.4.2 and the supplementary material (Sec. L).

3.3.2 Proposed Algorithm

In a similar spirit as alluded in the toy example, we propose Algorithm 1 (which is referred to as the *block edge query-based graph clustering* (BeQuec) algorithm). One can see that the BeQuec algorithm only consists of top- K SVD and least squares, which can be applied to very large graphs as long as K is of a moderate size. The SSMF stage (carried out by the *successive projection algorithm* (SPA) [35]) is a Gram–Schmidt type algorithm that is also scalable (see Algorithm 3 in the supplementary material in Sec. L). Note that in Algorithm 1, $\mathbf{U}_{\text{ref1}}$ and $\mathbf{U}_{\text{ref2}}$ are temporary variables defined to pass the subspace estimated in the current iteration to the subsequent iteration. We further explain the two major stages of BeQuec, namely, the range space estimation stage and the membership extraction stage as follows:

3.4 Key Ingredients of BeQuec

3.4.1 Range Space Estimation (Lines 1–13)

In a nutshell, BeQuec uses the idea from the toy example in an iterative fashion on the queried blocks $\mathbf{A}_{\ell_r, m_r}$ and $\mathbf{A}_{\ell_{r+1}, m_r}$ for $r = 1, \dots, L - 1$ —to estimate $\mathbf{U}^\top = \mathbf{G}^\top \mathbf{M}$, i.e., the range space of $\mathbf{M}^\top$. We start the

Algorithm 2 PairStitch

Require: $\mathbf{A}_{\ell,m}, \mathbf{A}_{\ell',m}, \mathbf{U}_{\text{ref}}$

- 1: $\mathbf{C} \leftarrow [\mathbf{A}_{\ell,m}^\top, \mathbf{A}_{\ell',m}^\top]^\top$;
 - 2: $[\bar{\mathbf{U}}_\ell^\top, \bar{\mathbf{U}}_{\ell'}^\top]^\top \Sigma \mathbf{V}_m^\top \leftarrow \text{svd}_K(\mathbf{C})$;
 - 3: $\mathbf{U}_\ell \leftarrow \bar{\mathbf{U}}_\ell \bar{\mathbf{U}}_{\ell'}^\dagger \mathbf{U}_{\text{ref}}$ $\triangleright$ “Stitch” $\mathbf{U}_\ell$ and $\mathbf{U}_{\text{ref}}$
 - 4: **return** $\mathbf{U}_\ell$.
-

iterations from $r = \lfloor L/2 \rfloor$ and perform the subspace stitching of the blocks in the ascending and descending orders, respectively—which helps reduce the subspace estimation error from the overall procedure, as will be seen in the analysis of Theorem 2.

3.4.2 SSF for Membership Extraction (Lines 14–15)

If $\mathbf{U}$ is perfectly estimated, the second stage estimates $\mathbf{M}$ from the following SSF model:

$$\mathbf{U}^\top = \mathbf{G}^\top \mathbf{M}, \quad \mathbf{M} \geq \mathbf{0}, \quad \mathbf{1}^\top \mathbf{M} = \mathbf{1}^\top, \quad (9)$$

where $\mathbf{G} \in \mathbb{R}^{K \times K}$ is nonsingular. The SSF problem is the cornerstone of a number of core tasks in machine learning and signal processing, e.g., hyperspectral unmixing and topic modeling, which admits a plethora of algorithms for provably estimating $\mathbf{M}$ under certain conditions [34]. To be precise, line 14 of Algorithm 1 invokes an SSF algorithm, namely, SPA [35]. Under the model in (9), SPA guarantees the identifiability of $\mathbf{M}$ utilizing the so-called *anchor node assumption* (ANC). Here, “identifiability” of $\mathbf{M}$ means that learning $\mathbf{M}$ up to a row permutation ambiguity can be guaranteed; see [36–38].

In essence, the ANC is identical to the so-called *separability* condition [39] from the SSF and nonnegative matrix factorization (NMF) literature:

Definition 1 (ANC/Seperability) [21, 39, 40] *The nonnegative matrix $\mathbf{M} \in \mathbb{R}^{K \times N}$ is said to satisfy the separability condition/ANC if there exists an index set $\Lambda = \{q_1, \dots, q_K\}$ such that $\mathbf{M}(:, q_k) = \mathbf{e}_k$ for all $k \in [K]$. When there exists $\Lambda = \{q_1, \dots, q_K\}$ such that $\|\mathbf{M}(:, q_k) - \mathbf{e}_k\|_2 \leq \varepsilon$ for all $k \in [K]$, then $\mathbf{M}$ is said to satisfy the ε -separability condition.*

In the context of GC, the ANC means that there exists a node q_k (i.e., an anchor node) that only belongs to cluster k for $k \in [K]$. Under the ANC, one can see that if Λ is known, then $\mathbf{G} = \mathbf{U}(\Lambda, :)$ can be easily “read out” from the rows of $\mathbf{U}$, which also enables estimating $\mathbf{M}$ using (constrained) least squares as well (see line 15 in Algorithm 1). The SPA algorithm identifies Λ using a Gram-Schmidt-like procedure in K iterations and its per-iteration complexity is $O(NK^2)$ operations [34, 35]; also see Sec. 4 in the supplementary material for details.

Remark 3 *A number of MMSB learning algorithms rely on the ANC to establish identifiability of $\mathbf{M}$; see, e.g., [19–21]. However, ANC does not reflect many important aspects of the GC problem, other than the existence of some anchor nodes. In the next subsection, we will present a sufficient condition that has a different flavor, and is arguably more intuitive in the context of GC (e.g., graph size and the cluster-cluster interaction pattern). Another remark is that SPA is not the only algorithm for separable SSF/NMF; see more options in [41–46]. Nonetheless, SPA strikes a good balance between complexity and accuracy.*

3.5 Performance Analysis

3.5.1 Algorithm Complexity

Note that Algorithm 1 only consists of top- K SVD on small blocks that have a size of $2(N/L) \times (N/L)$ assuming $|S_\ell| = N/L$ for all ℓ , which has a complexity of $O((N/L)K^2)$ flops—a similar complexity order of

the SPA stage. In terms of memory, the process never needs to instantiate any $N \times N$ matrix that is used in the convex programming based methods [cf. (2)]. Instead, the variables involved are with the size of $N \times K$, thereby being economical in terms of memory (since the number of clusters K is often much smaller than N).

3.5.2 Accuracy

We first use the ideal case where $\mathbf{A}(i, j) = \mathbf{P}(i, j)$ to analyze the key idea behind our proposed approach. Building on the understanding to the ideal noiseless case, the binary observation case, i.e., the model in (4), will be analyzed by treating it as a noisy version of the ideal case, with extensive care paid to the noise induced by the Bernoulli observation process.

To be specific, we have the following subspace identification result:

Theorem 1 (Ideal Case) Assume that $\mathbf{A}_{\ell,m} = \mathbf{P}_{\ell,m} = \mathbf{M}_\ell^\top \mathbf{B} \mathbf{M}_m \in \mathbb{R}^{|\mathcal{S}_\ell| \times |\mathcal{S}_m|}$ holds true for all $\ell, m \in [L]$ and $\text{rank}(\mathbf{M}_\ell) = \text{rank}(\mathbf{B}) = K$. Suppose that the $\mathbf{A}_{\ell,m}$'s are queried according to the proposed EQP. Then, the output $\hat{\mathbf{U}}$ by Algorithm 1 satisfies $\text{range}(\hat{\mathbf{U}}) = \text{range}(\mathbf{M}^\top)$.

The proof of the theorem is relegated to Appendix A. The proof showcases how the range spaces of $\mathbf{M}_\ell^\top$'s are 'stitched' together to recover $\text{range}(\mathbf{M}^\top)$, if $\mathbf{A}_{\ell,m} = \mathbf{P}_{\ell,m}$.

However, the Bernoulli parameters (i.e., the $\mathbf{P}_{\ell,m}$'s) are not observed in practice. Instead, one observes $\mathbf{A}_{\ell,m}$'s with binary values, following the Bernoulli observation model in (4). The Bernoulli observation process can be understood as a noisy data acquisition process. To characterize the performance under this noisy case, we consider the following assumption that is also reminiscent of the classic MMSB model [2]:

Assumption 1 The membership vectors $\mathbf{m}_n \forall n \in [N]$ are independently drawn from the Dirichlet distribution with parameter $\boldsymbol{\nu} = [\nu_1, \dots, \nu_K]^\top$, i.e., $\mathbf{m}_n \sim \text{Dir}(\boldsymbol{\nu})$ for $n \in [N]$.

We first connect the Dirichlet model with ANC/separability using the following proposition:

Proposition 1 Let $\mu > 0, \varepsilon > 0$. Under Assumption 1, assume that the number of nodes N satisfies

$$N = \Omega \left((G(\varepsilon, \boldsymbol{\nu}))^{-2} \log(K/\mu) \right), \quad (10)$$

where

$$G(\varepsilon, \boldsymbol{\nu}) = \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} \min_k \frac{\varepsilon^{\sum_{i \neq k} \nu_i}}{\prod_{i \neq k} \nu_i}, \quad (11)$$

and the gamma function $\Gamma(z) = \int_0^\infty x^{z-1} \exp(-x) dx$ for $z > 0$. Then, with probability of at least $1 - \mu$, $\mathbf{M}$ satisfies the ε -separability condition.

The proof can be found in Appendix B. This result translates the ANC to a condition that relies on the graph size N . To explain the condition in (10), it is critical to understand $G(\varepsilon, \boldsymbol{\nu})$. The value of $G(\varepsilon, \boldsymbol{\nu})$ is affected by how well $\mathbf{m}_n$'s are spread in the probability simplex. To illustrate this, consider a case with $N = 1000$, $K = 3$ and $\varepsilon = 0.1$. Table 1 tabulates the value of $G(\varepsilon, \boldsymbol{\nu})$ under this case for different Dirichlet parameters $\boldsymbol{\nu}$, namely, $[0.5, 0.5, 0.5]^\top$, $[2.0, 0.5, 0.5]^\top$, and $[3.0, 3.0, 3.0]^\top$. The corresponding $\mathbf{m}_n$'s distributions in the probability simplex are shown in Fig. 3. One can see from Table 1 and (10) that $\boldsymbol{\nu} = [0.5, 0.5, 0.5]^\top$ has a larger G -function value and thus requires a smaller N to satisfy the ε -separability condition. This is consistent with the illustration in Fig. 3.

Under Assumption 1, if $|\mathcal{S}_\ell| \geq K$, then $\text{rank}(\mathbf{M}_\ell) = K$ holds with probability one [47, 48]. This means that the following also holds with probability one:

$$\gamma := \max_\ell \kappa(\mathbf{M}_\ell) < \infty.$$

To proceed, we adopt the following definition of node degree that is widely used in network analysis [49–51]:

Table 1: The value of the function $G(\varepsilon, \boldsymbol{\nu})$ for different Dirichlet parameter $\boldsymbol{\nu}$ fixing $K = 3$ and $\varepsilon = 0.1$.

$\boldsymbol{\nu}$	$G(\varepsilon, \boldsymbol{\nu})$
$[0.5, 0.5, 0.5]^\top$	0.045
$[2.0, 0.5, 0.5]^\top$	8.4×10^{-4}
$[3.0, 3.0, 3.0]^\top$	7.0×10^{-5}

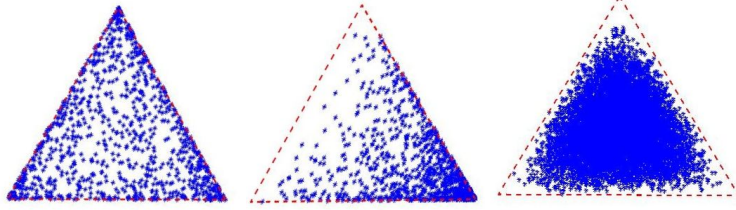


Figure 3: The columns of $\mathbf{M} \in \mathbb{R}^{3 \times 1000}$ generated via the Dirichlet distribution with $\boldsymbol{\nu} = [0.5, 0.5, 0.5]^\top$ (left), $\boldsymbol{\nu} = [2.0, 0.5, 0.5]^\top$ (middle) and $\boldsymbol{\nu} = [3.0, 3.0, 3.0]^\top$ (right) spread in the 2-dimensional probability simplex.

Definition 2 The degree of node i is the number of “similar nodes” it has in the adjacency graph; i.e., $\text{degree}(i) = \sum_{j=1}^N \mathbf{A}(i, j)$, where $\mathbf{A} \in \{0, 1\}^{N \times N}$.

Using Assumption 1, Proposition 1, the notations γ and $\text{degree}(i)$, we state the main theorem:

Theorem 2 (Binary Observation Case) Assume that the matrix $\mathbf{A}$ is generated following (3) and (4) and $\text{rank}(\mathbf{B}) = K$. Also assume that $\mathbf{A}_{\ell, m}$ ’s are queried following the EQP, with $|\mathcal{S}_\ell| = N/L$ for every ℓ , where N/L is an integer. Let $\rho := \max_{i,j} \mathbf{P}(i, j)$ be the maximal entry of $\mathbf{P}$. Suppose that $\rho = \Omega(L \log(N/L)/N)$ and $L = O(\rho N/d)$ where d is the maximal degree of all the nodes. Also assume that

$$N = \Omega \left(\max \left(L^2, \frac{\log(NK/L^2)}{G(\varepsilon, \boldsymbol{\nu})^2}, \frac{(K\gamma^2)^L \rho \kappa^2(\mathbf{B})}{\sigma_{\min}^2(\mathbf{B})} \right) \right) \quad (12)$$

for a certain constant $0 < \varepsilon = O\left(\frac{1}{K\gamma^3}\right)$. Then, the output $\widehat{\mathbf{U}}$ and $\widehat{\mathbf{M}}$ by Algorithm 1 satisfy the following with probability of at least $1 - O(L^2/N)$:

$$\begin{aligned} \|\widehat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_{\text{F}} &= \zeta = O \left(\frac{(K\gamma^2)^{L/2} \kappa(\mathbf{B}) \sqrt{\rho}}{\sigma_{\min}(\mathbf{B}) \sqrt{N/L}} \right), \\ \min_{\Pi} \|\widehat{\mathbf{M}} - \Pi \mathbf{M}\|_{\text{F}} &= O \left(K^2 \sigma_{\max}(\mathbf{M}) \gamma^3 \max(\varepsilon, \zeta) \right), \end{aligned}$$

where $\mathbf{U}$ is an orthogonal basis of $\text{range}(\mathbf{M}^\top)$, $\mathbf{O} \in \mathbb{R}^{K \times K}$ is an orthogonal matrix and $\Pi \in \mathbb{R}^{K \times K}$ is a permutation matrix.

The proof can be found in Appendix C which is an integration of (i) range space estimation accuracy under Bernoulli observations, (ii) least squares perturbation analysis, and (iii) noise robustness analysis of SPA. In particular, the conditions on ρ and d are needed for range space estimation under binary observations, while the conditions on N are used across parts (i)-(iii).

Theorem 2 can be understood as follows. First, under the presumed model, the proposed method favors larger graph sizes (i.e., larger N ’s) and larger query blocks (i.e., smaller L ’s). Particularly, the nested condition on N in (12) means that N should be larger than a certain threshold, since $N \geq E_1 \log(E_2 N)$ for

Table 2: The subspace distance between $\hat{U}$ and U for BeQuec in ideal and binary observation cases; $K = 5$; $L = 10$; $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5}]^\top$.

Graph Size (N)	Ideal Case	Binary Obs. Case
	Dist	Dist
2000	2.48×10^{-14}	0.3214
4000	4.86×10^{-14}	0.2080
8000	2.63×10^{-14}	0.1412
10000	1.28×10^{-13}	0.1257

positive constants E_1 and E_2 holds if N is sufficiently large. Second, if the nodes tend to be more focused on a single cluster (i.e., ν_k is small for every k), the G -function’s value in (11) becomes larger (cf. Fig. 3 and Table 1). Then, our method can reach a good accuracy level under smaller N ’s and larger L ’s. Third, the cluster-cluster interaction matrix B also matters: If the intra-cluster interactions dominate all node interactions and the cluster activities are balanced (i.e., $B(k, k) \approx B(j, j)$ for $j \neq k \in [K]$), then the condition number of B is expected to be small, thereby making membership learning easier.

Theorem 2 also sheds light on the impact of L . Ideally, one hopes to increase L since using a larger L means that the EQP uses fewer queries. With a proper L , the number of queries under the proposed EQP, i.e., $\Omega(N^2/L)$, can approach NK , which is the number of parameters to characterize any (unstructured) K -dimensional subspace in an N -dimensional space. However, L cannot be too large, as a large L may lead to the EQP condition $K \leq |S_\ell| = N/L$ being violated. In addition, a larger L makes the error bound in Theorem 2 looser, since the factor $(K\gamma^2)^{L/2}$ increases with L . Hence, the theorem reveals an important trade-off that is of interest to query designers.

Remark 4 Given the underlying SSMF structure of each block $A_{\ell,m}$, a natural question is as follows: Instead of first estimating $\text{range}(M^\top)$ and then estimating M , why not estimate M_ℓ (or M_m) from $A_{\ell,m}$ (using SSMF algorithms on the range space of $A_{\ell,m}$, i.e., $U_{\ell,m} \approx M_\ell^\top G_m$) and then stitch them together to recover M ? This is not infeasible—but is prone to failure. One reason is that this route needs that every block M_ℓ satisfies the ANC (at least approximately). This is much more unrealistic than requiring the entire M to satisfy the ANC (which is also reflected in Proposition 7).

4 Experiments

For both synthetic and real data experiments, we compare our algorithm with a number of state-of-the-art mixed membership learning algorithms, namely, *geometric intuition-based nonnegative matrix factorization* (GeoNMF) [19], *community detection via minimum volume simplex identification* (CD-MVSI) [22], and *community detection via Bayesian nonnegative matrix factorization* (CD-BNMF) [28]. These baselines are applied onto the queried blocks (e.g., $A_{\ell,m}$) to estimate M_ℓ and M_m , as discussed in Remark 4. The estimated M_ℓ ’s are aligned (by a nontrivial aligning procedure) to unify the row permutation ambiguity. This step is necessary since the estimated M_ℓ ’s from different blocks are subject to different row permutation ambiguities; see details in Appendix M.

We have two different sets of real data experiments, i.e., crowdclustering of dog breed image data and community detection in co-author networks. For real data experiments, we additionally apply another mixed membership learning algorithm, namely, *sequential projection after cleaning* (SPACL) [20], two versions of the spectral clustering algorithm, namely, the unnormalized and normalized spectral clustering algorithms—denoted as SC (unnorm.) and SC (norm.), respectively [52] and k -means. The k -means algorithm employed in these methods are repeated with 20 random initializations. We also use a convex optimization-based graph clustering algorithm with edge queries (denoted as Convex-Opt) [12] as a benchmark. We follow the setup in [12] and apply k -means directly on the adjacency graphs [denoted as

Table 3: The MSE, RE and averaged SRC of M for $K = 5$, $L = 10$ and $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5}]^\top$.

Graph Size (N)	Metric	BeQuec	GeoNMF	CD-MVSI	CD-BNMF
2000	MSE	0.0536	0.1060	0.8160	0.2632
	RE	0.2403	0.2603	0.9757	0.5330
	SRC	0.8058	0.7912	0.2060	0.6404
4000	MSE	0.0236	0.0894	0.7874	0.2185
	RE	0.1607	0.2464	0.9795	0.4407
	SRC	0.8486	0.7988	0.2556	0.7095
8000	MSE	0.0116	0.0331	0.7459	0.2006
	RE	0.1141	0.1338	0.9932	0.3492
	SRC	0.8909	0.8905	0.2887	0.7190
10000	MSE	0.0093	0.0070	0.7405	0.1668
	RE	0.1012	0.0827	0.9947	0.2707
	SRC	0.8992	0.9018	0.2916	0.7321

Table 4: The MSE, RE and averaged SRC of M for $K = 5$, $L = 10$ and $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} 1]^\top$.

Graph Size (N)	Metric	BeQuec	GeoNMF	CD-MVSI	CD-BNMF
2000	MSE	0.4302	0.4895	0.9063	0.8228
	RE	0.6060	0.6633	0.8417	0.8156
	SRC	0.5235	0.4739	0.1572	0.1832
4000	MSE	0.3037	0.3275	0.8975	0.7774
	RE	0.4420	0.4906	0.8515	0.8044
	SRC	0.6013	0.5678	0.1614	0.2089
8000	MSE	0.1073	0.4198	0.8969	0.8144
	RE	0.2631	0.5661	0.8533	0.8171
	SRC	0.7304	0.4948	0.1540	0.1794
10000	MSE	0.0863	0.5793	0.9013	0.8174
	RE	0.2442	0.5724	0.8464	0.8165
	SRC	0.7535	0.4082	0.1489	0.1736

k -means (graph)]. Note that in the crowdclustering experiment in Sec. 4.2 (also see Sec. 2.1 for the problem definition), we also use k -means on the original image samples [denoted as k -means (data)], since the raw data samples are available in this particular problem. The k -means (data) method serves as a basic benchmark where no annotator input is used. The source code and the real data used in our experiments are made publicly available (see the footnote in page 2).

4.1 Synthetic Data

We consider N nodes (where $N \in [2000, 10000]$) and $K = 5$ clusters. The membership vectors $\mathbf{m}_n$ are drawn from the Dirichlet distribution with parameter $\nu \in \mathbb{R}^5$. The symmetric cluster-cluster interaction matrix $\mathbf{B} \in \mathbb{R}^{K \times K}$ is generated by randomly sampling its diagonal and non-diagonal entries from uniform distributions between 0.8 to 1 and between 0 to η , respectively. The diagonal entries are chosen to be relatively large because intra-cluster connections are supposed to be more often. The parameter $\eta \in (0, 1]$ decides the level of interactions across different clusters, which is expected to be lower than that of the intra-cluster interactions.

Metrics. To evaluate the performance of the algorithms, we use a number of standard metrics. First, we

use the subspace distance (‘Dist’) between $\text{range}(\hat{U})$ and the ground-truth $\text{range}(M^\top)$, i.e.,

$$\text{Dist} = \|O_M^\perp Q_U\|_2,$$

where $O_M^\perp = I - M^\top(MM^\top)^{-1}M$ and Q_U is the orthogonal basis of $\hat{U}^\top$ [53]. To measure the accuracy of the estimated membership, we define the *mean squared error* (MSE) as follows [54, 55]:

$$\text{MSE} = \min_{\Pi} \frac{1}{K} \sum_{k=1}^K \left\| \frac{\Pi M(k, :)}{\|\Pi M(k, :)\|_2} - \frac{\hat{M}(k, :)}{\|\hat{M}(k, :)\|_2} \right\|_2^2$$

where $\hat{M}$ denotes the estimate of M and $\Pi \in \mathbb{R}^{K \times K}$ is a permutation matrix (due to the intrinsic row permutation of the outputs of SSMF algorithms such as SPA). As an additional reference, we also present the *relative error* (RE), i.e.,

$$\text{RE} = \min_{\Pi} \frac{\|\Pi M - \hat{M}\|_F}{\|M\|_F}.$$

We also use the averaged *Spearman’s rank correlation coefficient* (SRC). The averaged SRC is often used in the mixed membership learning literature (e.g., [19, 22]). The SRC takes values between -1 and 1 and has a higher value if the ranking of the entries in two vectors are more similar—which is desired.

Results. We first test the identifiability claims under the ideal case (i.e., $A = P$). The blocks of the adjacency matrix with the leftmost query pattern in Fig. 1 is used. We choose the number of groups $L = 10$ (percentage of queried edges = 27.92%), $\eta = 0.1$ and $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5}]^\top$. Table 2 shows the subspace estimation accuracy of our method measured using ‘Dist’. The results are averaged from 20 random trials. One can see that the proposed method estimates the subspace of the membership matrix M very accurately in the ideal case, which verifies our subspace identifiability analysis in Theorem 1. The right column of Table 2 presents the ‘Dist’ measure in the binary observation case. It shows that even under the binary “quantization” of P , the proposed method can still recover U to a good accuracy, if N is sufficiently large. This is consistent with Theorem 2.

Next, we consider the rightmost pattern in Fig. 1 to evaluate the proposed algorithm and the baselines in the binary observation case. We fix $L = 10$ and $\eta = 0.1$. We test two different cases of ν , i.e., $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5}]^\top$ and $\nu = [\frac{1}{5} \frac{1}{5} \frac{1}{5} \frac{1}{5} 1]^\top$. The results, which are averaged from 20 random trials, are presented in Tables 3 and 4 respectively. The former case of ν is easier, since many nodes’ membership vectors tend to reside on the boundaries/corners of the probability simplex (cf. Fig. 3)—leading to the existence of many anchor nodes. Hence, every block M_ℓ for $\ell \in [L]$ may satisfy the ANC, and BeQuec and GeoNMF both work reasonably well (cf. discussions in Remark 4). The latter case is more challenging due to $\nu_5 = 1$, which means that there is a cluster that has (much) fewer anchor nodes. This means that many M_ℓ ’s do not satisfy the ANC. Consequently, BeQuec outperforms the baselines, since our method does not need ANC to hold on each block M_ℓ . Particularly, in the latter case, the best baseline GeoNMF outputs high MSE = 0.579 even when $N = 10,000$, while BeQuec’s performance is considerably better under all three evaluation metrics (with MSE = 0.086).

More synthetic data experiments with different values of η and ν can be found in Sec. N of the supplementary material.

4.2 Crowdfunding using AMT Annotations

We consider the task of clustering a set of images of dogs into different breeds with the assistance of self-registered annotators from a crowdsourcing platform, i.e., the AMT platform.

AMT Data. We upload $N = 900$ images of $K = 5$ different breeds of dogs from the Stanford Dogs Dataset [56]. The chosen breeds are Pug (180 images), Ibizan Hound (180 images), Pomeranian (180 images), Norfolk Terrier (172 images) and Newfoundland (188 images). We followed the setup in Sec. 2.1 to design the experiment. Specifically, we showed each annotator a pair of dog images and ask a simple question (query): “Do this pair of dogs belong to the same category?”. If the annotator marks “Yes”, then we fill

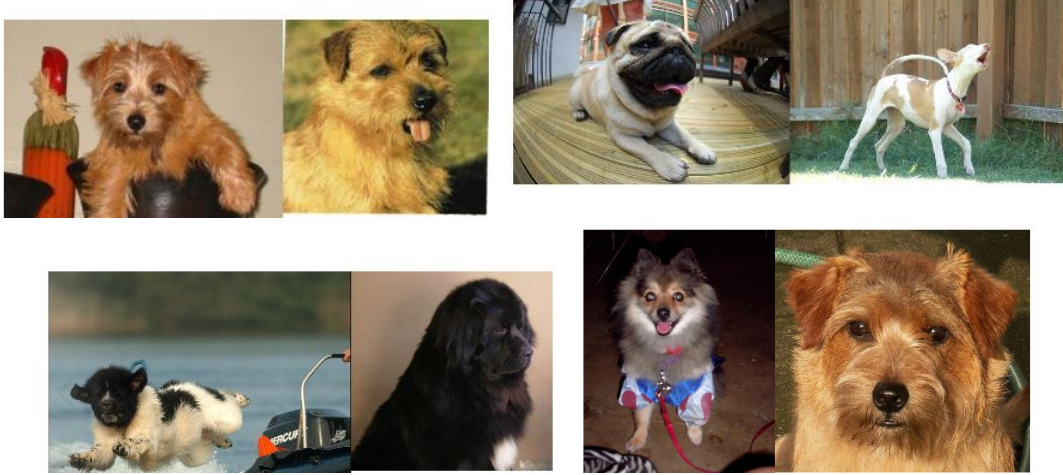


Figure 4: Some examples of the pairs presented to the AMT annotators in the crowdclustering experiment. The image pairs have true similarity labels 1,0,0,1, respectively, in clockwise order. The annotators marked 1,0,1,0, respectively.

the corresponding entry of $\mathbf{A}$ with “1”, and otherwise “0”. In this case, if one hopes to fill all the entries of $\mathbf{A}$, there are approximately 4×10^5 distinct pairs of images to be queried, which is costly. We selected image pairs following the block-diagonal pattern (i.e., the rightmost pattern in Fig. 1) by fixing $L = 15$. As a result, only 19.02% of the pairs (76,950) were queried². Fig. 4 shows some pairs of images queried in this experiment. We obtained the responses from 642 AMT annotators. Each annotator answered (non-overlapping) 120 queries on average. The error rate of the answers given by the annotators was 9.90%.

Metric. To quantify the performance, we use standard metrics for evaluating the clustering algorithms. Specifically, since each image is only labeled with a single cluster label, we round the learned membership vectors to the closest unit vectors. This way, standard metrics, i.e., *clustering accuracy* (ACC) and *normalized mutual information* (NMI) can be used for performance characterization (see definitions of ACC and NMI in [57]). ACC ranges from 0% to 100% where 100% means perfect clustering. NMI takes values in between 0 and 1 with 1 being the best performance.

Results. Table 5 summarizes the results output by the algorithms under test, where the first and the second best results are highlighted. One can see that both the proposed method BeQuec and Convex-Opt output surprisingly high ACC and NMI results. The baselines GeoNMF and SPACL also yield reasonable ACCs/NMIs. It is somewhat expected both SBM and MMSB based methods could work, since the nodes in this case admit single membership. Nonetheless, BeQuec and Convex-Opt’s better clustering performance over other baselines may be due to various reasons. For example, the fact that BeQuec works under less stringent conditions may make it more resilient to modeling mismatches. The all-at-once optimization strategy of Convex-Opt may help avoid error propagation as in greedy methods used in other baselines. In terms of runtime, notably, BeQuec runs much faster (approximately 1500 times faster than Convex-Opt, which strikes a good balance between accuracy and efficiency for this particular task.

Another interesting observation is that BeQuec largely outperforms directly applying clustering methods to the data samples without annotators’ assistance [cf. k -means (data)]. Note that k -means (data) sees all the data and their features (pixels in this case). It does not work well perhaps because the dogs across some breeds are naturally confusing. However, with annotating less than 20% of the data pairs’ similarity (without pointing out the exact breeds), the clustering error approaches zero. This experiment may have articulated the value of this simple annotation strategy.

²As a result, \$81.60 was paid for our annotation task to AMT, but annotating all the pairs of images for this experiment would have cost \$429.02.

Table 5: Graph clustering performance on the dog breed dataset with AMT similarity annotations. $N = 900$, $K = 5$, $L = 15$. Annotation proportion = 19.02%.

Algorithms	ACC (%)	NMI	Time (s)
BeQuec (proposed)	99.22	0.972	0.21
GeoNMF	96.88	0.929	2.80
SPACL	97.11	0.936	0.22
CD-MVSI	38.02	0.199	1.30
CD-BNMF	81.89	0.683	0.28
SC (unnorm.)	22.90	0.032	1.40
SC (norm.)	94.52	0.926	1.05
k -means (graph)	65.73	0.507	1.14
k -means (data)	23.56	0.052	2034.50
Convex-Opt	99.78	0.991	354.22

Table 6: Graph clustering performance on the dog breed dataset for different annotation error rates. $N = 900$, $K = 5$, $L = 15$. Annotation proportion = 19.02%.

Algorithms	annotation error=15%		annotation error=20%	
	ACC(%)	NMI	ACC(%)	NMI
BeQuec (proposed)	95.83	0.941	87.20	0.740
GeoNMF	85.33	0.688	79.83	0.629
SPACL	84.10	0.758	83.71	0.757
CD-MVSI	42.28	0.194	38.57	0.165
CD-BNMF	67.64	0.609	66.89	0.574
SC (unnorm.)	23.50	0.033	23.68	0.036
SC (norm.)	87.27	0.818	80.27	0.732
k -means (graph)	64.96	0.511	62.48	0.475
Convex-Opt	70.98	0.853	50.57	0.582

In order to better understand the impact of annotation errors on crowdclustering, we also manually “inject” some errors to the annotations. Specifically, we randomly flip a certain percentage of the correct annotations given by the AMT annotators, and observe how this affects the performance of the algorithms.

Table 6 presents the clustering performance obtained from two different error rates. The results are averaged over 20 trials, where in each trial the flipped annotations are randomly chosen. One can observe that the proposed BeQuec method offers the most competitive clustering performance. When the annotation error rate increases from 9.9% (Table 5) to 15% and 20%, the performance degradation of BeQuec is much less obvious compared to the baselines. In particular, one can see that NMI of BeQuec decreases from 0.972 to 0.941 when the annotation error rate is changed from 9.9% to 15%. However, GeoNMF’s NMI changes from 0.929 to 0.688, which is a much more drastic drop. This shows BeQuec’s robustness to annotation errors.

4.3 Co-author Network Community Detection

We test the edge query-based graph clustering algorithms using the co-authorship network data from MAG [58] and DBLP [59]. In all the datasets, node-node connection (co-authorship) is represented by a binary value (“1” means that the two nodes have co-authored at least one paper and “0” means otherwise). For

Table 7: Details of the co-author network datasets used in community detection experiments

Dataset	N	K
DBLP ₋₁	15,075	13
DBLP ₋₂	18,106	16
DBLP ₋₃	18,646	16
DBLP ₋₄	18,195	16
DBLP ₋₅	14,982	15
MAG1	37,680	3
MAG2	19,457	3

Table 8: Averaged SRC and runtime in seconds. $L = 10$ (percentage of queried edges = 27.92%).

Datasets	BeQuec (proposed)		GeoNMF		SPACL		CD-MVSI		CD-BNMF	
	SRC	Time(s.)	SRC	Time(s.)	SRC	Time(s.)	SRC	Time(s.)	SRC	Time(s.)
DBLP ₋₁	0.177	0.47	0.075	15.87	0.072	0.73	0.074	0.95	0.058	5.61
DBLP ₋₂	0.166	0.38	0.089	11.13	0.103	1.02	0.055	0.76	0.070	8.20
DBLP ₋₃	0.166	0.46	0.076	17.18	0.094	1.01	0.064	0.98	0.060	8.50
DBLP ₋₄	0.191	0.48	0.098	15.76	0.095	0.95	0.070	0.95	0.079	8.50
DBLP ₋₅	0.195	0.41	0.084	15.65	0.156	0.67	0.073	0.88	0.074	5.85
MAG1	0.125	0.26	0.122	1.79	0.123	0.97	0.089	0.59	0.069	37.29
MAG2	0.441	0.23	0.240	4.66	0.267	0.50	0.249	0.53	0.122	9.92

our experiments, we use the version of these networks published by the authors of [19]. Each network is provided with ground-truth mixed membership of the nodes. In these networks, nodes represent the authors of the research papers from different fields of study. Some nodes exhibit multiple memberships, since some authors publish in different fields.

MAG Data. In our experiment, we use a subset of nodes that have a degree of at least 5. Consequently, two co-author networks, namely, MAG1 and MAG2, have 37,680 and 19,457 authors, respectively. In MAG1 and MAG2, all the authors are from 3 different fields of study, i.e., $K = 3$.

DBLP Data. The 5 DBLP datasets (denoted as DBLP1, ..., DBLP5) used in [19] have 6176, 3145, 2605, 3056, and 6269 nodes, respectively. In addition, the number of clusters present in the DBLP datasets are 6, 3, 3, 3, and 4, respectively. To make the clustering problem more challenging, we combine the smaller original DBLP networks into larger-size networks following the setting in [22]. We use the notation DBLP_{-i} to denote the network where all the 5 DBLP networks except the i -th one are combined. This way, different DBLP_{-i}'s can be produced for testing the algorithms. In addition, each DBLP_{-i}, for $i = 1, \dots, 5$, contains 13-16 clusters (see Table 7 for details), presenting more challenging scenarios.

Metric. The methods are evaluated using the averaged SRC as in the literature [19, 22]. We let all the algorithms to access only part of the network under the block-diagonal query pattern in Fig. 1. We randomize the node order in each of the 20 trials and present the averaged performance.

Results. We first test the mixed membership learning performance by setting $L = 10$ for all the MAG and DBLP datasets under test. The percentage of edges queried in this setting is 27.92%. From Table 8, one can see that the proposed BeQuec algorithm consistently outperforms the baselines in all the graphs under test. In addition, BeQuec's runtime performance is also the most competitive.

Table 9 uses MAG2 dataset to show the performance of BeQuec under various L 's. We also apply the clustering algorithms used in crowdclustering to benchmark BeQuec's performance. The ACC metric is used in this table since the clustering algorithms do not output mixed membership. One can see that the performance of all algorithms decreases along with L 's increase—which is consistent with our analysis

Table 9: Clustering accuracy (%) of MAG2 under various L (percentage of queried edges (%)). $N = 19457$, $K = 3$.

Alorithms	$L = 10$ (27.92%)	$L = 25$ (11.58%)	$L = 50$ (5.82%)	$L = 75$ (3.86%)	$L = 100$ (2.87%)
BeQuec (proposed)	78.70	77.19	67.81	61.85	56.98
GeoNMF	58.16	57.87	56.88	52.68	52.33
SPACL	57.49	58.17	55.25	54.22	53.07
CD-MVSI	53.45	21.82	14.57	13.53	11.71
CD-BNMF	51.19	49.35	49.37	48.99	48.97
SC (unnorm.)	56.97	56.49	56.38	55.88	55.51
SC (norm.)	64.24	62.15	57.15	56.45	55.80
k -means	54.31	53.47	53.87	53.15	51.25

in Theorem 2. For each L used, the corresponding queried percentage of edges is listed in the first row. Notably, when $L = 50$, i.e, only 5.82% of $\mathcal{A}$ is observed, the proposed method still outputs a reasonable clustering accuracy. This demonstrates a promising balance between the query sample complexity and the graph clustering accuracy.

5 Conclusion

In this work, we proposed a graph query scheme that enables provable graph clustering with partially observed edges. Unlike previous works relying on random graph edge query and computationally heavy convex programming, our method features a lightweight algorithm and works with systematic edge query patterns that are arguably more realistic in some applications. Our method also learns mixed membership of nodes from overlapping clusters, which improves upon existing provable graph query-based methods that work with single membership and disjoint clusters. We tested the proposed BeQuec algorithm on two real-world network data analytics problems, namely, crowdclustering and network community detection. BeQuec offers promising results in both experiments.

A potential future direction is to reduce the number of queries by allowing $N/L < K$. This may be achieved by using third-order graph statistics from the sampled blocks and a tensor decomposition framework. Using tensor-based approaches may likely present a challenging optimization problem without polynomial-time algorithms (as opposed to the matrix based framework as in this work). Studying its query complexity and computational cost trade-off presents an interesting research topic.

References

- [1] S. E. Schaeffer, “Graph clustering,” *Computer Science Review*, vol. 1, no. 1, pp. 27 – 64, 2007.
- [2] E. M. Airoldi, D. M. Blei, S. E. Fienberg, and E. P. Xing, “Mixed membership stochastic blockmodels,” *Journal of Machine Learning Research*, vol. 9, no. Sep, pp. 1981–2014, 2008.
- [3] D. Kuang, C. Ding, and H. Park, “Symmetric nonnegative matrix factorization for graph clustering,” *Proceedings of SIAM International Conference on Data Mining*, pp. 106–117, 2012.
- [4] S. Van Dongen, “Graph clustering via a discrete uncoupling process,” *SIAM Journal on Matrix Analysis and Applications*, vol. 30, no. 1, pp. 121–141, 2008.
- [5] A. Saade, M. Lelarge, F. Krzakala, and L. Zdeborová, “Clustering from sparse pairwise measurements,” *IEEE International Symposium on Information Theory*, pp. 780–784, 2016.

- [6] M. T. Schaub, S. Segarra, and J. N. Tsitsiklis, "Blind identification of stochastic block models from dynamical observations," *SIAM Journal on Mathematics of Data Science*, vol. 2, no. 2, pp. 335–367, 2020.
- [7] J. Leskovec and C. Faloutsos, "Sampling from large graphs," *Proceedings of ACM International Conference on Knowledge Discovery and Data Mining*, p. 631–636, 2006.
- [8] P. Hu and W. C. Lau, "A survey and taxonomy of graph sampling," *Computing Research Repository*, 2013.
- [9] M. J. Salganik and D. D. Heckathorn, "Sampling and estimation in hidden populations using respondent-driven sampling," *Sociological Methodology*, vol. 34, no. 1, pp. 193–240, 2004.
- [10] M. Malekinejad, L. Johnston, C. Kendall, L. R. Kerr, M. Rifkin, and G. Rutherford, "Using respondent-driven sampling methodology for HIV biological and behavioral surveillance in international settings: A systematic review," *AIDS and Behavior*, vol. 12, pp. 5–30, 07 2008.
- [11] K. A. Gill and M. O’Neal, "Survey of Soybean Insect Pollinators: Community Identification and Sampling Method Analysis," *Environmental Entomology*, vol. 44, no. 3, pp. 488–498, 02 2015.
- [12] R. Koralakai Vinayak, S. Oymak, and B. Hassibi, "Graph clustering with missing data: Convex algorithms and analysis," *Advances in Neural Information Processing Systems*, vol. 27, pp. 2996–3004, 2014.
- [13] R. Koralakai Vinayak and B. Hassibi, "Crowdsourced clustering: Querying edges vs triangles," *Advances in Neural Information Processing Systems*, vol. 29, pp. 1316–1324, 2016.
- [14] Y. Chen, A. Jalali, S. Sanghavi, and H. Xu, "Clustering partially observed graphs via convex optimization," *Journal of Machine Learning Research*, vol. 15, no. 1, pp. 2213–2238, 2014.
- [15] A. Mazumdar and B. Saha, "Clustering with noisy queries," in *Proceedings of the 31st International Conference on Neural Information Processing Systems*, 2017, p. 5790–5801.
- [16] —, "Clustering via crowdsourcing," *arXiv preprint arXiv:1604.01839*, 2016.
- [17] R. K. Vinayak, "Graph clustering: Algorithms, analysis and query design," Ph.D. dissertation, California Institute of Technology, 2018.
- [18] C. Särndal, B. Swensson, and J. Wretman, *Model Assisted Survey Sampling*, ser. Springer Series in Statistics. Springer-Verlag, 2003.
- [19] X. Mao, P. Sarkar, and D. Chakrabarti, "On mixed memberships and symmetric nonnegative matrix factorizations," in *International Conference on Machine Learning*, 2017, pp. 2324–2333.
- [20] —, "Estimating mixed memberships with sharp eigenvector deviations," *Journal of the American Statistical Association*, pp. 1–13, 2020.
- [21] M. Panov, K. Slavnov, and R. Ushakov, "Consistent estimation of mixed memberships with successive projections," *Complex Networks*, p. 53–64, 2017.
- [22] K. Huang and X. Fu, "Detecting overlapping and correlated communities without pure nodes: Identifiability and algorithm," *Proceedings of International Conference on Machine Learning*, vol. 97, pp. 2859–2868, 2019.
- [23] J. Yi, R. Jin, S. Jain, T. Yang, and A. K. Jain, "Semi-crowdsourced clustering: Generalizing crowd labeling by robust distance metric learning," *Advances in Neural Information Processing Systems*, vol. 25, pp. 1772–1780, 2012.
- [24] S. Ibrahim and X. Fu, "Learning mixed membership from adjacency graph via systematic edge query: Identifiability and algorithm," in *IEEE International Conference on Acoustics, Speech and Signal Processing*, 2021, pp. 5370–5374.

- [25] R. G. Gomes, P. Welinder, A. Krause, and P. Perona, "Crowdclustering," *Advances in Neural Information Processing Systems*, vol. 24, pp. 558–566, 2011.
- [26] J. Yi, R. Jin, A. Jain, and S. Jain, "Crowdclustering with sparse pairwise labels: A matrix completion approach," in *AAAI Workshop on Human Computation*, 2012.
- [27] J. Xie, S. Kelley, and B. K. Szymanski, "Overlapping community detection in networks: The state-of-the-art and comparative study," *ACM Computing Surveys*, vol. 45, no. 4, Aug. 2013.
- [28] I. Psorakis, S. Roberts, M. Ebdon, and B. Sheldon, "Overlapping community detection using Bayesian non-negative matrix factorization," *Physical Review E*, vol. 83, p. 066114, 06 2011.
- [29] A. Y. Ng, M. I. Jordan, and Y. Weiss, "On spectral clustering: Analysis and an algorithm," *Advances in Neural Information Processing Systems*, vol. 14, pp. 849–856, 2002.
- [30] S. Oymak and B. Hassibi, "Finding dense clusters via "low rank + sparse" decomposition," *Computing Research Repository*, 04 2011.
- [31] A. Anandkumar, R. Ge, D. Hsu, and S. M. Kakade, "A tensor approach to learning mixed membership community models," *Journal of Machine Learning Research*, vol. 15, no. 1, pp. 2239–2312, 2014a.
- [32] F. J. Király, L. Theran, and R. Tomioka, "The algebraic combinatorial approach for low-rank matrix completion," *Journal of Machine Learning Research*, vol. 16, no. 41, pp. 1391–1436, 2015.
- [33] D. L. Pimentel-Alarcón, N. Boston, and R. D. Nowak, "A characterization of deterministic sampling patterns for low-rank matrix completion," *IEEE J. Sel. Topics Signal Process.*, vol. 10, no. 4, pp. 623–636, 2016.
- [34] X. Fu, K. Huang, N. D. Sidiropoulos, and W.-K. Ma, "Nonnegative matrix factorization for signal and data analytics: Identifiability, algorithms, and applications," *IEEE Signal Process. Mag.*, vol. 36, no. 2, pp. 59–80, 2019.
- [35] N. Gillis and S. Vavasis, "Fast and robust recursive algorithms for separable nonnegative matrix factorization," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 36, no. 4, pp. 698–714, April 2014.
- [36] T.-H. Chan, C.-Y. Chi, Y.-M. Huang, and W.-K. Ma, "A convex analysis-based minimum-volume enclosing simplex algorithm for hyperspectral unmixing," *IEEE Trans. Signal Process.*, vol. 57, no. 11, pp. 4418–4432, Nov. 2009.
- [37] X. Fu, K. Huang, and N. D. Sidiropoulos, "On identifiability of nonnegative matrix factorization," *IEEE Signal Process. Lett.*, vol. 25, no. 3, pp. 328–332, 2018.
- [38] K. Huang, N. Sidiropoulos, and A. Swami, "Non-negative matrix factorization revisited: Uniqueness and algorithm for symmetric decomposition," *IEEE Trans. Signal Process.*, vol. 62, no. 1, pp. 211–224, 2014.
- [39] D. Donoho and V. Stodden, "When does non-negative matrix factorization give a correct decomposition into parts?" in *Advances in Neural Information Processing Systems*, vol. 16, 2003.
- [40] S. Ibrahim, X. Fu, N. Kargas, and K. Huang, "Crowdsourcing via pairwise co-occurrences: Identifiability and algorithms," *Advances in Neural Information Processing Systems*, vol. 32, pp. 7847–7857, 2019.
- [41] A. Kumar, V. Sindhwani, and P. Kambadur, "Fast conical hull algorithms for near-separable non-negative matrix factorization," *Proceedings of International Conference on Machine Learning*, vol. 28, pp. 231–239, 2012.
- [42] T. Mizutani, "Ellipsoidal rounding for nonnegative matrix factorization under noisy separability," *Journal of Machine Learning Research*, vol. 15, no. 29, pp. 1011–1039, 2014.

- [43] X. Fu and W.-K. Ma, "Robustness analysis of structured matrix factorization via self-dictionary mixed-norm optimization," *IEEE Signal Process. Lett.*, vol. 23, no. 1, pp. 60–64, 2016.
- [44] N. Gillis, "Successive nonnegative projection algorithm for robust nonnegative blind source separation," *SIAM Journal on Imaging Sciences*, vol. 7, no. 2, pp. 1420–1450, 2014.
- [45] X. Fu, W.-K. Ma, T.-H. Chan, and J. M. Bioucas-Dias, "Self-dictionary sparse regression for hyperspectral unmixing: Greedy pursuit and pure pixel search are related," *IEEE J. Sel. Topics Signal Process.*, vol. 9, no. 6, pp. 1128–1141, 2015.
- [46] X. Fu, K. Huang, B. Yang, W.-K. Ma, and N. D. Sidiropoulos, "Robust volume minimization-based matrix factorization for remote sensing and document clustering," *IEEE Trans. Signal Process.*, vol. 64, no. 23, pp. 6254–6268, 2016.
- [47] N. D. Sidiropoulos, E. E. Papalexakis, and C. Faloutsos, "Parallel randomly compressed cubes: A scalable distributed architecture for big tensor decomposition," *IEEE Signal Processing Magazine*, vol. 31, no. 5, pp. 57–70, 2014.
- [48] N. D. Sidiropoulos and A. Kyrillidis, "Multi-way compressed sensing for sparse low-rank tensors," *IEEE Signal Process. Lett.*, vol. 19, no. 11, pp. 757–760, 2012.
- [49] U. Luxburg, "A tutorial on spectral clustering," *Statistics and Computing*, vol. 17, no. 4, p. 395–416, Dec. 2007.
- [50] J. Scott, "Social network analysis," *Sociology*, vol. 22, no. 1, pp. 109–127, 1988.
- [51] J. Lei and A. Rinaldo, "Consistency of spectral clustering in stochastic block models," *Annals of Statistics*, vol. 43, no. 1, pp. 215–237, 02 2015.
- [52] Jianbo Shi and J. Malik, "Normalized cuts and image segmentation," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 22, no. 8, pp. 888–905, 2000.
- [53] G. H. Golub and C. F. Van Loan, *Matrix computations*. JHU Press, 2012, vol. 3.
- [54] X. Fu, W.-K. Ma, K. Huang, and N. D. Sidiropoulos, "Blind separation of quasi-stationary sources: Exploiting convex geometry in covariance domain," *IEEE Trans. Signal Process.*, vol. 63, no. 9, pp. 2306–2320, 2015.
- [55] X. Fu, S. Ibrahim, H.-T. Wai, C. Gao, and K. Huang, "Block-randomized stochastic proximal gradient for low-rank tensor factorization," *IEEE Trans. Signal Process.*, vol. 68, pp. 2170–2185, 2020.
- [56] A. Khosla, N. Jayadevaprakash, B. Yao, and L. Fei-Fei, "Novel dataset for fine-grained image categorization," in *First Workshop on Fine-Grained Visual Categorization, IEEE Conference on Computer Vision and Pattern Recognition*, Colorado Springs, CO, June 2011.
- [57] C. Lu, J. Tang, M. Lin, L. Lin, S. Yan, and Z. Lin, "Correntropy induced l2 graph for robust subspace clustering," *IEEE International Conference on Computer Vision*, pp. 1801–1808, 2013.
- [58] A. Sinha, Z. Shen, Y. Song, H. Ma, D. Eide, B. Hsu, and K. Wang, "An overview of microsoft academic service (MAS) and applications," *Proceedings of International Conference on World Wide Web*, p. 243–246, 2015.
- [59] M. Ley, "The DBLP computer science bibliography: Evolution, research issues, perspectives," *Proceedings of International Symposium on String Processing and Information Retrieval*, pp. 1–10, 2002.
- [60] A. B. Frigyik, A. Kapila, and M. Gupta, "Introduction to the dirichlet distribution and related processes," *Dept. Elect. Eng., Univ. Washington, Seattle, WA, USA*, 2010.

- [61] A. Gouda and T. Szántai, “On numerical calculation of probabilities according to dirichlet distribution,” *Annals of Operations Research*, vol. 177, pp. 185–200, 2010.
- [62] M. J. Wainwright, *High-dimensional statistics: A non-asymptotic viewpoint*. Cambridge University Press, 2019, vol. 48.
- [63] Y. Yu, T. Wang, and R. J. Samworth, “A useful variant of the Davis–Kahan theorem for statisticians,” *Biometrika*, vol. 102, no. 2, pp. 315–323, 04 2014.
- [64] P. Wedin, “Perturbation theory for pseudo-inverses,” *BIT Numerical Mathematics*, vol. 13, pp. 217–232, 1973.
- [65] R. A. Horn and C. R. Johnson, *Matrix Analysis*, 2nd ed. USA: Cambridge University Press, 2012.
- [66] U. Araújo, B. Saldanha, R. Galvão, T. Yoneyama, H. Chame, and V. Visani, “The successive projections algorithm for variable selection in spectroscopic multicomponent analysis,” *Chemometrics and Intelligent Laboratory Systems*, vol. 57, no. 2, pp. 65–73, 2001.
- [67] R. Jonker and T. Volgenant, “Improving the hungarian assignment algorithm,” *Operations Research Letters*, vol. 5, no. 4, pp. 171–175, 1986.

A Proof of Theorem 1

Algorithm 1 aims to output $\hat{U}$ such that $\text{range}(\hat{U}) = \text{range}([M_1, \dots, M_L]^\top)$. It implies that $\hat{U} = [U_1^\top, \dots, U_L^\top]^\top = [M_1, M_2, \dots, M_L]^\top G = M^\top G$, for a certain nonsingular $G \in \mathbb{R}^{K \times K}$. To show this, we start by considering the below four blocks for any $r \in [L - 2]$:

$$P_{\ell_r, m_r} = M_{\ell_r}^\top B M_{m_r}, \quad (13a)$$

$$P_{\ell_{r+1}, m_r} = M_{\ell_{r+1}}^\top B M_{m_r}, \quad (13b)$$

$$P_{\ell_{r+1}, m_{r+1}} = M_{\ell_{r+1}}^\top B M_{m_{r+1}}, \quad (13c)$$

$$P_{\ell_{r+2}, m_{r+1}} = M_{\ell_{r+2}}^\top B M_{m_{r+1}}. \quad (13d)$$

Algorithm 1 defines the following matrices in r th and $(r + 1)$ th iterations, respectively (recall that we have assumed $A_{\ell, m} = P_{\ell, m} = M_\ell^\top B M_m$ for any $\ell, m \in [L]$ in the ideal case):

$$C_r := [P_{\ell_r, m_r}^\top, P_{\ell_{r+1}, m_r}^\top]^\top, \quad C_{r+1} := [P_{\ell_{r+1}, m_{r+1}}^\top, P_{\ell_{r+2}, m_{r+1}}^\top]^\top.$$

The top- K SVDs of C_r and C_{r+1} can be represented as follows:

$$C_r = [U_{\ell_r}^\top, U_{\ell_{r+1}}^\top]^\top \Sigma_r V_{m_r}^\top, \quad (14a)$$

$$C_{r+1} = [\bar{U}_{\ell_{r+1}}^\top, \bar{U}_{\ell_{r+2}}^\top]^\top \Sigma_{r+1} V_{m_{r+1}}^\top. \quad (14b)$$

Combining (13) and (14), and using the assumption that $\text{rank}(M_\ell) = \text{rank}(B) = K$ and $K \leq |S_\ell|$ for every ℓ , the following equations hold:

$$U_{\ell_r} = M_{\ell_r}^\top G_r, \quad U_{\ell_{r+1}} = M_{\ell_{r+1}}^\top G_r, \quad (15a)$$

$$\bar{U}_{\ell_{r+1}} = M_{\ell_{r+1}}^\top \bar{G}_{r+1}, \quad \bar{U}_{\ell_{r+2}} = M_{\ell_{r+2}}^\top \bar{G}_{r+1}, \quad (15b)$$

where $G_r \in \mathbb{R}^{K \times K}$ and $\bar{G}_{r+1} \in \mathbb{R}^{K \times K}$ are certain nonsingular matrices. The equations in (15) imply that U_{ℓ_r} , $U_{\ell_{r+1}}$ and $\bar{U}_{\ell_{r+2}}$ are the bases of $\text{range}(M_{\ell_r}^\top)$, $\text{range}(M_{\ell_{r+1}}^\top)$ and $\text{range}(M_{\ell_{r+2}}^\top)$, respectively. Since $G_r = \bar{G}_{r+1}$ does not generally hold, the basis $\bar{U}_{\ell_{r+2}}$ cannot be directly combined with the bases U_{ℓ_r} and

$U_{\ell_{r+1}}$ to obtain $\text{range}([M_{\ell_r}, M_{\ell_{r+1}}, M_{\ell_{r+2}}]^\top)$. Therefore, we are interested in obtaining the basis defined as below:

$$U_{\ell_{r+2}} := M_{\ell_{r+2}}^\top G_r. \quad (16)$$

In order to identify $U_{\ell_{r+2}}$ as defined in (16), we use the second equality in (15a) and the first equality in (15b) to construct the following:

$$\bar{U}_{\ell_{r+1}}^\dagger U_{\ell_{r+1}} = \bar{G}_{r+1}^{-1} G_r. \quad (17)$$

Combining the second equality in (15b) with (17), we have

$$\bar{U}_{\ell_{r+2}} \bar{U}_{\ell_{r+1}}^\dagger U_{\ell_{r+1}} = (M_{\ell_{r+2}}^\top \bar{G}_{r+1}) (\bar{G}_{r+1}^{-1} G_r) = M_{\ell_{r+2}}^\top G_r.$$

Hence, the following holds:

$$U_{\ell_{r+2}} = \bar{U}_{\ell_{r+2}} \bar{U}_{\ell_{r+1}}^\dagger U_{\ell_{r+1}} = M_{\ell_{r+2}}^\top G_r, \quad (18)$$

where $\bar{U}_{\ell_{r+2}}$ and $\bar{U}_{\ell_{r+1}}$ are obtained from the top- K SVD of C_{r+1} and $U_{\ell_{r+1}}$ is obtained from the top- K SVD of C_r .

Similarly, $U_{\ell_{r-1}}$ can be identified using the top- K SVDs of C_r and C_{r-1} by the following:

$$U_{\ell_{r-1}} = \bar{U}_{\ell_{r-1}} \bar{U}_{\ell_r}^\dagger U_{\ell_r} = M_{\ell_{r-1}}^\top G_r, \quad (19)$$

where U_{ℓ_r} is obtained from the top- K SVD of C_r and $\bar{U}_{\ell_{r-1}}$ and $\bar{U}_{\ell_r}$ are obtained from the top- K SVD of $C_{r-1} := [P_{\ell_{r-1}, m_{r-1}}^\top, P_{\ell_r, m_{r-1}}^\top]^\top$.

Algorithm 1 first applies top- K SVD for C_T where $T = \lfloor L/2 \rfloor$ and the bases U_{ℓ_T} and $U_{\ell_{T+1}}$ are identified (lines 2-3). By letting $G_T = G$, we thus obtain

$$U_{\ell_T} = M_{\ell_T}^\top G \quad \text{and} \quad U_{\ell_{T+1}} = M_{\ell_{T+1}}^\top G. \quad (20)$$

Next, we use induction to show how the iterations in Algorithm 1 (lines 5-8) identify $U_{\ell_{T+2}}, \dots, U_L$. The induction hypothesis is that before any iteration at $r = T+1, T+2, \dots, L-1$, the basis U_{ℓ_r} is identified from the previous iteration such that $U_{\ell_r} = M_{\ell_r}^\top G$. Note that the induction hypothesis holds true for $r = T+1$ from lines 2-3 of Algorithm 1 by performing top- K SVD of C_T . From (18), we have seen that given U_{ℓ_r} , one can identify $U_{\ell_{r+1}}$ using the matrices resulted from the top- K SVD of C_r . Therefore, by induction, we can establish that Algorithm 1 identifies

$$U_{\ell_r} = M_{\ell_r}^\top G, \text{ for every } r \in \{T+2, \dots, L\}. \quad (21)$$

Similarly, using induction via employing the results in (19) and (20), we can establish that lines 9-12 of Algorithm 1 identifies

$$U_{\ell_r} = M_{\ell_r}^\top G, \text{ for every } r \in \{1, \dots, T-1\}. \quad (22)$$

Combining (20), (21), and (22), one can see that Algorithm 1 identifies $\hat{U} = [U_1^\top, U_2^\top, \dots, U_L^\top]^\top$ such that $\text{range}(\hat{U}) = \text{range}([M_1, M_2, \dots, M_L]^\top)$.

B Proof of Proposition 1

B.1 Preliminaries of Dirichlet Distribution

Let $\mathcal{P}^K = \{z \in \mathbb{R}^K | z^\top \mathbf{1} = 1, z \geq 0\}$ denote the $(K-1)$ -dimensional probability simplex. The Dirichlet distribution with parameter $\nu = [\nu_1, \dots, \nu_K]^\top$ has the following probability density function (PDF):

$$f(z_{-k}; \nu) = \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} (1 - \sum_{i \neq k} z_i)^{\nu_k - 1} \prod_{i \neq k} z_i^{\nu_i - 1}, \quad (23)$$

for any $k \in [K]$ where $\mathbf{z}_{-k} = [z_1, \dots, z_{k-1}, z_{k+1}, \dots, z_K]^\top$ and $[z_1, \dots, z_K]^\top \in \mathcal{P}^K$ [60]. The choice of k in (23) is arbitrary since $\sum_{i \neq k} z_i = 1 - z_k$ holds for any k . The cumulative distribution function (CDF) of the Dirichlet distribution is given by

$$\begin{aligned} F(\mathbf{z}_{-k}; \boldsymbol{\nu}) &= \int_0^{z_1} \cdots \int_0^{z_{k-1}} \int_0^{z_{k+1}} \cdots \int_0^{z_K} f(\mathbf{z}_{-k}; \boldsymbol{\nu}) d\mathbf{z}_{-k} \\ &= \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} L_k(\mathbf{m}, \boldsymbol{\nu}) \prod_{i \neq k} \frac{z_i^{\nu_i}}{\nu_i}, \end{aligned} \quad (24)$$

where the function $L_k(\mathbf{m}, \boldsymbol{\nu})$ is known as the Lauricella function of the first kind and $L_k(\mathbf{m}, \boldsymbol{\nu}) > 1$ for any k [61]. The CDF $F(\mathbf{z}_{-k}; \boldsymbol{\nu})$ measures the cumulative probability starting from the k -th vertex of the probability simplex till each coordinate $i \neq k$ reaches z_i .

B.2 Proof of the Proposition

According to Assumption 1, $\mathbf{m}_n \in \mathcal{P}^K$ holds for all $n \in [N]$ and the vectors $\mathbf{m}_n$ are sampled randomly from a Dirichlet distribution with parameter $\boldsymbol{\nu}$. Using Assumption 1 we aim to analyze the conditions under which the matrix $\mathbf{M} \in [\mathbf{m}_1, \dots, \mathbf{m}_N]$ satisfies ε -separability (see Definition 1), i.e., there exists $\boldsymbol{\Lambda} = \{q_1, \dots, q_K\}$ such that for $k = 1, \dots, K$,

$$\|\mathbf{m}_{q_k} - \mathbf{e}_k\|_2 \leq \varepsilon, \quad \mathbf{m}_{q_k} \in \mathcal{P}^K. \quad (25)$$

The criterion in (25) means that there exist $\mathbf{m}_{q_k} \in \mathcal{P}_k(\varepsilon)$ for $k = 1, \dots, K$ where

$$\mathcal{P}_k(\varepsilon) = \{\mathbf{z} \in \mathcal{P}^K \mid \|\mathbf{z} - \mathbf{e}_k\|_2 \leq \varepsilon\}. \quad (26)$$

Next, we proceed to check under what conditions there exists at least one $\mathbf{m}_n$ among $\mathbf{m}_1, \dots, \mathbf{m}_N$ such that $\mathbf{m}_n \in \mathcal{P}_k(\varepsilon)$. For this, let us define two random variables as below:

$$E_{n,k} = \begin{cases} 1, & \text{if } \mathbf{m}_n \in \mathcal{P}_k(\varepsilon) \\ 0, & \text{otherwise.} \end{cases} \quad (27)$$

$$H_k = \sum_{n=1}^N E_{n,k}. \quad (28)$$

In order to characterize the probability that there exists at least one $\mathbf{m}_n$ among $\mathbf{m}_1, \dots, \mathbf{m}_N$ such that $\mathbf{m}_n \in \mathcal{P}_k(\varepsilon)$, we need to characterize $\Pr(H_k > 0)$. To achieve this, we invoke Chernoff-Hoeffding bound [62]:

Theorem 3 Let $E_{1,k}, \dots, E_{N,k}$ be independent bounded random variables such that $E_{n,k} \in [a_{n,k}, b_{n,k}]$. Then for any $t > 0$,

$$\Pr\left(\sum_{n=1}^N E_{n,k} - \sum_{n=1}^N \mathbb{E}[E_{n,k}] \leq -t\right) \leq e^{\frac{-2t^2}{\sum_{n=1}^N (b_{n,k} - a_{n,k})^2}}.$$

By letting $t = \sum_{n=1}^N \mathbb{E}[E_{n,k}]$ and using the fact that $E_{n,k} \in [0, 1]$, we get the below from Theorem 3:

$$\Pr(H_k > 0) \geq 1 - e^{-\frac{2\left(\sum_{n=1}^N \mathbb{E}[E_{n,k}]\right)^2}{N}}. \quad (29)$$

Next, we need to characterize the term $\mathbb{E}[E_{n,k}]$ which appears in the R.H.S. of (29). From the definition of the random variable $E_{n,k}$ in (27), we can see that

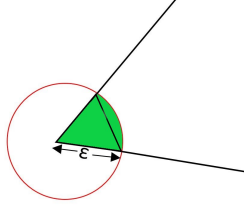


Figure 5: The figure shows any vertex $k \in \{1, 2, 3\}$ of the probability simplex $\mathcal{P}^3$. The green shaded region represents $\mathcal{P}_k(\varepsilon)$ (see (26)) which is the intersection of the probability simplex $\mathcal{P}^3$ and a Euclidean ball of radius ε with its center at e_k . The small triangle near the vertex is 2-dimensional probability simplex with edge lengths ε ; i.e, it intersects the co-ordinate axes at $\frac{\varepsilon}{\sqrt{2}}e_k$. It can be observed that the volume of the green shaded region is larger than the volume of the small triangle.

$$\begin{aligned}
\mathbb{E}[E_{n,k}] &= \Pr(\mathbf{m}_n \in \mathcal{P}_k(\varepsilon)) = \int_{\mathbf{z} \in \mathcal{P}_k(\varepsilon)} f(\mathbf{z}_{-k}; \boldsymbol{\nu}) d\mathbf{z}_{-k}, \\
&\geq \int_{z_1=0}^{\frac{\varepsilon}{\sqrt{2}}} \cdots \int_{z_{k-1}=0}^{\frac{\varepsilon}{\sqrt{2}}} \int_{z_{k+1}=0}^{\frac{\varepsilon}{\sqrt{2}}} \cdots \int_{z_K=0}^{\frac{\varepsilon}{\sqrt{2}}} f(\mathbf{z}_{-k}; \boldsymbol{\nu}) d\mathbf{z}_{-k} \\
&= F(\mathbf{z}_{-k}^*; \boldsymbol{\nu}) \geq \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} \prod_{i \neq k} \frac{\left(\frac{\varepsilon}{\sqrt{2}}\right)^{\nu_i}}{\nu_i} \\
&= \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} \frac{(\varepsilon/\sqrt{2})^{\sum_{i \neq k} \nu_i}}{\prod_{i \neq k} \nu_i} := G_k(\varepsilon, \boldsymbol{\nu}), \tag{30}
\end{aligned}$$

where $\mathbf{z}_{-k}^* = [z_1^*, \dots, z_{k-1}^*, z_{k+1}^*, z_K^*]^\top = [\frac{\varepsilon}{\sqrt{2}}, \dots, \frac{\varepsilon}{\sqrt{2}}]^\top$. The first inequality is obtained by generalizing the geometric property as shown in Fig. 5 for any $K-1$ dimensional probability simplex and the second inequality is obtained by combining (24) and the property that $L_k(\mathbf{m}, \boldsymbol{\alpha}) > 1$ for any k [61].

Applying (30) in (29), we get the probability such that there exists at least one $\mathbf{m}_n$ such that $\mathbf{m}_n \in \mathcal{P}_k(\varepsilon)$ as below:

$$\Pr(H_k > 0) \geq 1 - e^{-2N(G_k(\varepsilon, \boldsymbol{\nu}))^2}.$$

To characterize the corresponding probability for all the K vertices, we can apply the union bound to obtain the below

$$\Pr\left(\bigcap_{k=1}^K (H_k > 0)\right) \geq 1 - \sum_{k=1}^K e^{-2N(G_k(\varepsilon, \boldsymbol{\nu}))^2} \tag{31}$$

$$\geq 1 - K e^{-2N(G(\varepsilon, \boldsymbol{\nu}))^2}, \tag{32}$$

where $G(\varepsilon, \boldsymbol{\nu}) = \min_k G_k(\varepsilon, \boldsymbol{\nu})$.

Eq. (31) implies that with probability greater than or equal to $1 - K e^{-2N(G(\varepsilon, \boldsymbol{\nu}))^2}$, the matrix $\mathbf{M} = [\mathbf{m}_1, \dots, \mathbf{m}_N]$ satisfies ε -separability. By letting $\mu = K e^{-2N(G(\varepsilon, \boldsymbol{\nu}))^2}$, one can see that if $N \geq \log\left(\frac{K}{\mu}\right) / (2G(\varepsilon, \boldsymbol{\nu}))^2$, the matrix $\mathbf{M} = [\mathbf{m}_1, \dots, \mathbf{m}_N]$ satisfies ε -separability with probability greater than $1 - \mu$ where

$$G(\varepsilon, \boldsymbol{\nu}) = \frac{\Gamma(\sum_{i=1}^K \nu_i)}{\prod_{i=1}^K \Gamma(\nu_i)} \min_k \frac{(\varepsilon/\sqrt{2})^{\sum_{i \neq k} \nu_i}}{\prod_{i \neq k} \nu_i}.$$

C Proof of Theorem 2

The analysis for Theorem 2 consists of three major steps. We provide an overview to the proof as follows:

(A) Range Space Estimation. In this step, lines 1-13 of Algorithm 1 that learn the matrix U such that $\text{range}(U) = \text{range}(M^\top)$ is analyzed.

- **A1)** Under the binary observation model, the top- K SVD's are performed on the adjacency blocks $A_{\ell,m}$'s not on blocks $P_{\ell,m}$'s. The noise associated with this is characterized by using Lemmas 1 and 2
- **A2)** Each matrix block U_r is estimated using the sub-algorithm 2. Lemmas 3-5 are derived to characterize this basis stitching operation. A number of lemmas (Lemmas 10-13) are employed to assist this step, for e.g., to characterize the condition numbers of the latent matrices.
- **A3)** The final estimation error bound for the matrix $U = [U_1^\top, \dots, U_L^\top]^\top$ is obtained by recursively applying the error bound for each U_r and by employing Lemma 6.

(B) Membership Estimation. In this step, the error bound for the membership matrix M (see lines 14-15 of Algorithm 1) is analyzed.

- **B1)** The conditions for M to approximately satisfy the ANC is characterized by utilizing Proposition 1
- **B2)** The final estimation error bound for M is derived using Lemma 7. Lemma 7 characterizes the combined noise in the model which is contributed by both the noise in the estimate of U and the deviation from the ANC.

(C) Putting Together. In this step, various conditions used in Lemma 1-7 and Proposition 1 are combined to derive the final conditions of Theorem 2.

C.1 Range Space Estimation

In Theorem 1, the subspace identifiability is established via successively applying the top- K SVD on $C_r := [P_{\ell_r, m_r}^\top, P_{\ell_{r+1}, m_r}^\top]^\top$, i.e., $C_r = [\bar{U}_{\ell_r}^\top, \bar{U}_{\ell_{r+1}}^\top]^\top \Sigma_r V_{m_r}^\top$. However, since C_r is not observed in practice, Algorithm 1 performs the top- K SVD on its estimates $\hat{C}_r$, which is defined as $\hat{C}_r := [A_{\ell_r, m_r}^\top, A_{\ell_{r+1}, m_r}^\top]^\top$. The factors extracted from the top- K SVD of $\hat{C}_r$ are denoted as $\hat{\bar{U}}_{\ell_r}$ and $\hat{\bar{U}}_{\ell_{r+1}}$, which correspond to the 'noiseless versions' $\bar{U}_{\ell_r}$ and $\bar{U}_{\ell_{r+1}}$, respectively. To proceed, we invoke the following lemma to bound the estimation accuracy of $\hat{\bar{U}}_{\ell_r}$ and $\hat{\bar{U}}_{\ell_{r+1}}$.

Lemma 1 Denote $\sigma_1 := \max_{r \in [L-1]} \sigma_{\max}(C_r)$ and $\sigma_K := \min_{r \in [L-1]} \sigma_{\min}(C_r)$. Assume that

$$\text{rank}(C_r) = K, \quad \|\hat{C}_r - C_r\|_2 \leq \|C_r\|_2. \quad (33)$$

Then, there exists an orthogonal matrix O_r such that

$$\|\hat{\bar{U}}_{\ell_r} - \bar{U}_{\ell_r} O_r\|_F \leq \frac{6\sqrt{2K}\sigma_1 \|\hat{C}_r - C_r\|_2}{\sigma_K^2}. \quad (34)$$

In addition, $\|\hat{\bar{U}}_{\ell_{r+1}} - \bar{U}_{\ell_{r+1}} O_r\|_F$ admits the same upper bound in (34).

The proof of the lemma is given in Sec. D in the supplemental material.

In order to utilize the bound in (34), we proceed to characterize the term $\|\hat{C}_r - C_r\|_2$. To this end, we present the following lemma:

Lemma 2 Let $\rho := \max_{i,j \in [N]} P(i,j)$. Assume there exists a positive constant c_0 such that $L \leq \frac{4\rho N}{d}$ and $\rho \geq \frac{Lc_0 \log(3N/L)}{3N}$. Then, there exists another positive constant $\tilde{c}$ (which is a function of c_0) such that

$$\|\hat{\mathbf{C}}_r - \mathbf{C}_r\|_2 \leq \tilde{c}\sqrt{\rho N/L}, \quad \forall r \in [L-1] \quad (35)$$

holds with probability of at least $1 - \frac{L}{2N}$.

The proof of the lemma is given in Sec. E of the supplementary material.

Since $\{\ell_r\}_{r=1}^L = [L]$ according to EQP, without loss of any generality (w.l.o.g.), in the following sections of the proof, we can fix $\ell_r = r$ for every r . Combining Lemmas 1 and 2 the inequalities, i.e.,

$$\|\hat{\mathbf{U}}_r - \bar{\mathbf{U}}_r \mathbf{O}_r\|_F \leq \phi \text{ and } \|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_F \leq \phi \quad (36)$$

hold with probability of at least $1 - \frac{L}{N}$ where $\phi = \frac{6\tilde{c}\sqrt{2K\rho N}\sigma_1}{\sqrt{L}\sigma_K^2}$.

Algorithm 1 (lines 2-3) first estimates $\mathbf{U}_T$ and $\mathbf{U}_{T+1}$ by directly performing the top- K SVD to $\hat{\mathbf{C}}_T = [\mathbf{A}_{T,m_T}^\top, \mathbf{A}_{T+1,m_T}^\top]^\top$. Recall that we use the notation $\mathbf{C}_T = [\mathbf{P}_{T,m_T}^\top, \mathbf{P}_{T+1,m_T}^\top]^\top = [\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top \boldsymbol{\Sigma}_T \mathbf{V}_{m_T}^\top$. Therefore, by (36) and denoting $\mathbf{O}_T := \mathbf{O}$, we have

$$\|\hat{\mathbf{U}}_T - \mathbf{U}_T \mathbf{O}\|_F \leq \phi \text{ and } \|\hat{\mathbf{U}}_{T+1} - \mathbf{U}_{T+1} \mathbf{O}\|_F \leq \phi \quad (37)$$

holding with probability of at least $1 - \frac{L}{N}$, where $\hat{\mathbf{U}}_T$ and $\hat{\mathbf{U}}_{T+1}$ are the factors resulted from the top- K SVD of $\hat{\mathbf{C}}_T$.

Next, the iterative steps (lines 5-8) in Algorithm 1 (which includes the sub-algorithm PairStitch) perform the below operation to estimate the basis $\mathbf{U}_{r+1}$ for $r = T+1, \dots, L$ (where $T = \lfloor L/2 \rfloor$):

$$\hat{\mathbf{U}}_{r+1} = \hat{\mathbf{U}}_{r+1} \hat{\mathbf{U}}_r^\dagger \hat{\mathbf{U}}_r. \quad (38)$$

In order to analyze $\hat{\mathbf{U}}_{r+1}$, we first define the following notations:

$$\alpha_{\max} := \max_r \sigma_{\max}(\mathbf{M}_r), \quad \alpha_{\min} := \min_r \sigma_{\min}(\mathbf{M}_r).$$

The above implies that $\frac{\alpha_{\max}}{\alpha_{\min}} = \max_r \kappa(\mathbf{M}_r) = \gamma$. Under Assumption 1, $\text{rank}(\mathbf{M}_r) = K$ holds with probability 1. Hence, both $\alpha_{\max} < \infty$ and $\alpha_{\min} > 0$ also hold with probability one.

Using the above notations and the assumptions of Theorem 2 we have the following three lemmas:

Lemma 3 There exists constants $\beta_{\min} > 0$ and $\beta_{\max} < \infty$ such that for every $r \in [L-1]$, we have

$$\begin{aligned} \sigma_{\max}([\mathbf{M}_r, \mathbf{M}_{r+1}]) &\leq \beta_{\max} = \sqrt{2K} \alpha_{\max}, \\ \sigma_{\min}([\mathbf{M}_r, \mathbf{M}_{r+1}]) &\geq \beta_{\min} = \alpha_{\min}. \end{aligned}$$

Lemma 4 Let $\hat{\mathbf{U}}_r$ denote the noisy estimate of $\bar{\mathbf{U}}_r \mathbf{O}_r$ obtained from the top- K SVD of $\hat{\mathbf{C}}_r$. Suppose that $\|\hat{\mathbf{U}}_r - \bar{\mathbf{U}}_r \mathbf{O}_r\|_2 \leq \frac{\alpha_{\min}}{2\beta_{\max}}$. Then, we have $\|\hat{\mathbf{U}}_r^\dagger - (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 \leq \frac{2\sqrt{2}\beta_{\max}^2}{\alpha_{\min}^2} \|\hat{\mathbf{U}}_r - (\bar{\mathbf{U}}_r \mathbf{O}_r)\|_2$.

Lemma 5 Consider the relation in (38). Assume that $\|\hat{\mathbf{U}}_r - \bar{\mathbf{U}}_r \mathbf{O}_r\|_2 \leq \frac{\alpha_{\min}}{2\beta_{\max}}$ and $\|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 \leq \frac{\alpha_{\min}}{2\beta_{\max}}$, then we have

$$\begin{aligned} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1} \mathbf{O}\|_2 &\leq \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 + 2 \left(\frac{\alpha_{\max}}{\beta_{\min}} \right)^2 \|\hat{\mathbf{U}}_r^\dagger - (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 \\ &\quad + 4 \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\hat{\mathbf{U}}_r - \mathbf{U}_r \mathbf{O}\|_2. \end{aligned} \quad (39)$$

The proofs of Lemmas 3, 5 can be found Sec. F, G, and H of the supplementary material, respectively. By applying Lemma 4 to the second term in (39) of Lemma 5, we get

$$\begin{aligned} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_2 &\leq \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1}\mathbf{O}_r\|_2 + 4\sqrt{2} \left(\frac{\alpha_{\max}\beta_{\max}}{\beta_{\min}\alpha_{\min}} \right)^2 \|\hat{\mathbf{U}}_r - (\bar{\mathbf{U}}_r\mathbf{O}_r)\|_2 \\ &\quad + 4 \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_2. \end{aligned} \quad (40)$$

Employing Lemma 3, we have

$$\frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} = \frac{\sqrt{2K}\alpha_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\alpha_{\min}} = \sqrt{2K}\gamma^2. \quad (41)$$

Combining (40) and (41) and defining $\bar{\kappa} = \sqrt{2K}\gamma^2$, we obtain

$$\|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_2 \leq \bar{\kappa} \|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1}\mathbf{O}_r\|_2 + 4\sqrt{2}\bar{\kappa}^2 \|\hat{\mathbf{U}}_r - (\bar{\mathbf{U}}_r\mathbf{O}_r)\|_2 + 4\bar{\kappa} \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_2. \quad (42)$$

Hence, we obtain the following inequalities with probability of at least $1 - \frac{L}{N}$:

$$\begin{aligned} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_{\text{F}} &\leq \sqrt{K} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_2 \\ &\leq \sqrt{K}\bar{\kappa} \|\hat{\mathbf{U}}_{r+1} - \bar{\mathbf{U}}_{r+1}\mathbf{O}_r\|_2 + 4\sqrt{2}\sqrt{K}\bar{\kappa}^2 \|\hat{\mathbf{U}}_r - (\bar{\mathbf{U}}_r\mathbf{O}_r)\|_2 \\ &\quad + 4\sqrt{K}\bar{\kappa} \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_2 \\ &\leq 5\sqrt{2K}\bar{\kappa}^2\phi + 4\bar{\kappa}\sqrt{K} \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_{\text{F}}, \end{aligned}$$

where $\phi = \frac{6\tilde{c}\sqrt{2K\rho N}\sigma_1}{\sqrt{L}\sigma_K^2}$. The first inequality is obtained by using norm equivalence, the second inequality is obtained by applying (42) and the last inequality is by applying (36) and using the fact that $\bar{\kappa} \geq 1$.

By recursively applying the above result for any $r \in \{T+1, \dots, L-1\}$, we further have

$$\begin{aligned} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_{\text{F}} &\leq 5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa} \left(5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa} \dots \right. \\ &\quad \left. (5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa} \|\hat{\mathbf{U}}_{T+1} - \mathbf{U}_{T+1}\mathbf{O}\|_{\text{F}}) \right) \\ &\leq 5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa} \left(5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa} \dots \right. \\ &\quad \left. (5\sqrt{2K}\bar{\kappa}^2\phi + 4\sqrt{K}\bar{\kappa}\phi) \right) \\ &= 5\sqrt{2K}\bar{\kappa}^2\phi \left(1 + 4\sqrt{K}\bar{\kappa} + (4\sqrt{K}\bar{\kappa})^2 + \dots + (4\sqrt{K}\bar{\kappa})^{r-T} \right), \end{aligned}$$

where we have applied (37) for obtaining the last inequality. By applying the geometric sum formula, for any $r \in \{T+1, \dots, L-1\}$, we have

$$\begin{aligned} \|\hat{\mathbf{U}}_{r+1} - \mathbf{U}_{r+1}\mathbf{O}\|_{\text{F}} &\leq \frac{5\sqrt{2K}\bar{\kappa}^2\phi \left((4\sqrt{K}\bar{\kappa})^{r-T+1} - 1 \right)}{4\sqrt{K}\bar{\kappa} - 1} \\ &\leq \frac{5\sqrt{2K}\bar{\kappa}^2\phi (4\sqrt{K}\bar{\kappa})^{r-T+1}}{4\sqrt{K}\bar{\kappa} - 1}. \end{aligned} \quad (43)$$

Following the same derivation, one can show that for $r = 2, \dots, T$, we have

$$\|\hat{\mathbf{U}}_{r-1} - \mathbf{U}_{r-1}\mathbf{O}\|_{\text{F}} \leq \frac{5\sqrt{2K}\bar{\kappa}^2\phi (4\sqrt{K}\bar{\kappa})^{T-r+2}}{4\sqrt{K}\bar{\kappa} - 1}. \quad (44)$$

Consequently, we get

$$\begin{aligned}
\|\hat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_F &\leq \sum_{r=1}^L \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_F \\
&= \sum_{r=1}^{T-1} \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_F + \|\hat{\mathbf{U}}_T - \mathbf{U}_T\mathbf{O}\|_F + \|\hat{\mathbf{U}}_{T+1} - \mathbf{U}_{T+1}\mathbf{O}\|_F + \sum_{r=T+2}^L \|\hat{\mathbf{U}}_r - \mathbf{U}_r\mathbf{O}\|_F \\
&\leq \frac{5\sqrt{2K}\bar{\kappa}^2\phi(4\sqrt{K}\bar{\kappa})^{T+1}}{(4\sqrt{K}\bar{\kappa} - 1)^2} + \frac{5\sqrt{2K}\bar{\kappa}^2\phi(4\sqrt{K}\bar{\kappa})^{T+1}}{(4\sqrt{K}\bar{\kappa} - 1)^2} \\
&= \frac{10\sqrt{2K}\bar{\kappa}^2\phi(4\sqrt{K}\bar{\kappa})^{L/2+1}}{(4\sqrt{K}\bar{\kappa} - 1)^2},
\end{aligned}$$

where the first inequality is by the triangle inequality, the second inequality is obtained by (37), (43), (44), and by the geometric sum formula.

Substituting $\phi = \frac{6\tilde{c}\sqrt{2K\rho N}\sigma_1}{\sqrt{L}\sigma_K^2}$ into the above result, we have the following result with probability at least $1 - L(L-1)/N$:

$$\|\hat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_F \leq \frac{15K\tilde{c}\bar{\kappa}^2(4\sqrt{K}\bar{\kappa})^{(L/2+1)}\sigma_1\sqrt{\rho N}}{2(\sqrt{K}\bar{\kappa} - 1)^2\sigma_K^2\sqrt{L}}. \quad (45)$$

Note that we have applied union bound to reach the probability $1 - L(L-1)/N$, since the bound (36) derived from Lemma 2 are invoked for every $r \in [L-1]$. Next, we bound σ_1 and σ_K using the following lemma:

Lemma 6 Assume that the columns of $\mathbf{M}$ are generated from any continuous distribution and there exist positive constants c and C , where $c \leq C$ depending only on distribution of the columns of $\mathbf{M}$. Also assume that $(N/L) \geq \frac{4}{c} \log(NK/L)$. Then, the following holds true with probability of at least $1 - (3L(L-1)/N)$:

$$\sigma_K \geq \sqrt{2}c\sigma_{\min}(\mathbf{B})N/L, \quad (46)$$

$$\sigma_1 \leq \sqrt{2}C\sigma_{\max}(\mathbf{B})N/L. \quad (47)$$

The proof can be found in Sec. J of the supplementary material.

Applying Lemma 6 in (45) and substituting $\bar{\kappa} = \sqrt{2K}\gamma^2$, we get the following with probability at least $1 - (4L(L-1)/N)$:

$$\|\hat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_F \leq \frac{15K^2\tilde{c}C\gamma^4(4\sqrt{2K}\gamma^2)^{(L/2+1)}\kappa(\mathbf{B})\sqrt{\rho}}{c^2\sigma_{\min}(\mathbf{B})(\sqrt{2K}\gamma^2 - 1)^2\sqrt{2(N/L)}} := \zeta. \quad (48)$$

C.2 Membership Estimation

Note that in the $\mathbf{A} = \mathbf{P}$ case, $\text{range}(\mathbf{U}) = \text{range}(\mathbf{M}^\top)$. Therefore, there exists a nonsingular matrix $\mathbf{G} \in \mathbb{R}^{K \times K}$ such that $\mathbf{U} = \mathbf{M}^\top \mathbf{G}$. In the binary observation case, in order to estimate the membership matrix $\mathbf{M}$, Algorithm 1 applies SPA to the estimated $\hat{\mathbf{U}}$. To analyze this step, we start by considering the following noisy NMF model:

$$\hat{\mathbf{U}}^\top = \mathbf{O}^\top \mathbf{G}^\top \mathbf{M} + \mathbf{N}, \quad (49)$$

where $\mathbf{O} \in \mathbb{R}^{K \times K}$ is the orthogonal matrix, $\mathbf{M}$ satisfies simplex constraints, i.e., $\mathbf{M} \geq \mathbf{0}$, $\mathbf{1}^\top \mathbf{M} = \mathbf{1}^\top$ and $\mathbf{N} \in \mathbb{R}^{K \times N}$ is the noise matrix, whose Euclidean norm upper bound was obtained in the previous subsection by (48) such that $\|\mathbf{N}\|_F \leq \zeta$. From the noisy observation $\hat{\mathbf{U}}$, one can get an estimate for $\mathbf{M}$ employing

SPA. However, SPA can provably work only when $\mathbf{M}$ satisfies ε -separability. For this, we use Proposition 1 to get the condition for ε -separability for $\mathbf{M}$. Then, combining the noise robustness analysis of SPA in [35] and Proposition 1, we show the following:

Lemma 7 Consider the noisy model in (49) and assume that $\|\mathbf{N}\|_F \leq \zeta$ and that the assumptions in Proposition 1 hold true. Also, suppose that

$$\zeta \leq \frac{1}{4\sqrt{2}K\alpha_{\max}\tilde{\kappa}(\mathbf{M})} \quad \text{and} \quad \varepsilon \leq \frac{1}{4\sqrt{2}K\gamma\tilde{\kappa}(\mathbf{M})}. \quad (50)$$

Then, Algorithm 1 outputs $\widehat{\mathbf{M}}$ such that

$$\min_{\Pi} \|\widehat{\mathbf{M}} - \Pi\mathbf{M}\|_F \leq \sigma_{\max}(\mathbf{M})\sqrt{2}K\gamma\tilde{\kappa}(\mathbf{M})(3\varepsilon + 4\zeta), \quad (51)$$

where $\tilde{\kappa}(\mathbf{M}) = 1 + 160K\gamma^2$.

The proof is provided in Sec K of the supplementary material.

C.3 Putting Together

To make the bounds for $\mathbf{U}$ and $\mathbf{M}$ given by (48) and (51), respectively, legitimate, we need the assumptions in Lemmas 1, 7 hold true.

The assumptions on $\mathbf{C}_r$ in Lemma 1 given by the equalities in (33) are satisfied if

$$\text{rank}(\mathbf{M}_r) = \text{rank}(\mathbf{B}) = K, \quad \forall r \in [L] \quad \text{and} \quad (52)$$

$$\tilde{c}\sqrt{\rho N/L} \leq \sigma_K. \quad (53)$$

Under Assumption 1, $\text{rank}(\mathbf{M}_r) = K$ hold with probability one [47, 48]. Using (46), The condition (53) can be further written as

$$N/L \geq \frac{\tilde{c}^2\rho}{2\sigma_{\min}^2(\mathbf{B})c^2}. \quad (54)$$

The assumptions in Lemma 4 and Lemma 5 can be written as:

$$\begin{aligned} \phi &= \frac{6\tilde{c}\sqrt{2K\rho N}\sigma_1}{\sqrt{L}\sigma_K^2} \leq \frac{6\tilde{c}C\kappa(\mathbf{B})\sqrt{LK\rho}}{c^2\sigma_{\min}(\mathbf{B})\sqrt{N}} \\ &\leq \frac{\alpha_{\min}}{2\beta_{\max}} = \frac{\alpha_{\min}}{2\sqrt{2}K\alpha_{\max}} = \frac{1}{2\sqrt{2}K\gamma}, \end{aligned}$$

where we have utilized Lemma 6 to obtain the first inequality and Lemma 3 to obtain the second equality. Hence, we obtain the condition such that

$$N/L \geq \frac{(6\sqrt{2}\tilde{c}C)^2\kappa^2(\mathbf{B})K^2\rho\gamma^2}{c^4\sigma_{\min}^2(\mathbf{B})}. \quad (55)$$

We can see that the condition in (54) is satisfied if (55) holds.

Also, the below conditions from Lemmas 2 and 6 have to be satisfied:

$$L \leq \frac{4\rho N}{d}, \quad \rho \geq \frac{Lc_0 \log(3N/L)}{3N} \quad (56a)$$

$$(N/L) \geq \frac{4}{c} \log(NK/L). \quad (56b)$$

Next, we proceed to analyze the conditions in Lemma 7. The conditions in Proposition 1 has to be satisfied by Lemma 7 as well which implies that the below has to hold:

$$N = \Omega((G(\varepsilon, \nu))^{-2} \log(K/\mu)), \quad (57)$$

where $G(\varepsilon, \nu)$ is given by (11).

Lemma 7 has a condition on ζ given by (50). By substituting the value of ζ given by (48), the condition on ζ in (50) can be re-written as:

$$N/L \geq \frac{(30K^3 \tilde{c} C \gamma^4 \kappa(\mathbf{B}) \tilde{\kappa}(\mathbf{M}) \alpha_{\max})^2 \rho(4\sqrt{2}K\gamma^2)^{L+2}}{(c^2 \sigma_{\min}(\mathbf{B})(K\gamma^2 - 1)^2)^2}. \quad (58)$$

Note that the condition on N given by (55) is satisfied if the above condition on N holds true.

Hence, if the conditions (57) and (58) are satisfied together with (52) and (56), the bounds for $\mathbf{U}$ and $\mathbf{M}$ given by Theorem 2 hold.

Next, we proceed to derive the probability with which the final bounds for $\mathbf{U}$ and $\mathbf{M}$ hold true. For this, we summarize the probabilistic events appeared in the proof as follows:

1. The first event is that the subspace perturbation error associated with every adjacency matrix blocks $\mathbf{C}_r$, $r \in [L-1]$ is bounded. This happens with probability of at least $1 - L(L-1)/N$ (see Lemma 2);
2. The second event is that the singular values σ_1 and σ_K associated with the adjacency matrix blocks $\mathbf{C}_r$, $r \in [L-1]$ are bounded. This holds with probability of at least $1 - (3L(L-1)/N)$ (see Lemma 6);
3. The third event is that there exist anchor nodes in the membership matrix $\mathbf{M}$ up to an error bounded by ε . This happens with probability of at least $1 - \mu$ according to Proposition 1. By substituting $\mu = L(L-1)/N$ in (10) of Proposition 1, we can get that the result in Proposition 1 holds with probability of at least $1 - L(L-1)/N$.

Hence, by applying the union bound to the three event probabilities listed above, the final bounds for $\mathbf{U}$ and $\mathbf{M}$ hold true with probability of at least $1 - (5L(L-1)/N)$.

D Proof of Lemma 1

We invoke the following lemma from [63] to characterize the estimation accuracy of the factors resulting from the top- K SVD of $\hat{\mathbf{C}}_r$:

Lemma 8 [63] *Let $\mathbf{C} \in \mathbb{R}^{m \times n}$ and $\hat{\mathbf{C}} \in \mathbb{R}^{m \times n}$ have singular values $\alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_{\min(m,n)}$ and $\hat{\alpha}_1 \geq \hat{\alpha}_2 \geq \dots \geq \hat{\alpha}_{\min(m,n)}$, respectively. Fix $1 \leq t \leq s \leq \text{rank}(\mathbf{C})$ and assume that $\min(\alpha_{t-1}^2 - \alpha_t^2, \alpha_s^2 - \alpha_{s+1}^2) > 0$, where $\alpha_0^2 := \infty$ and $\alpha_{\text{rank}(\mathbf{C})+1}^2 := 0$. Let $q := s - t + 1$ and let $\mathbf{U} = [\mathbf{u}_t, \mathbf{u}_{t+1}, \dots, \mathbf{u}_s] \in \mathbb{R}^{m \times q}$ and $\hat{\mathbf{U}} = [\hat{\mathbf{u}}_t, \hat{\mathbf{u}}_{t+1}, \dots, \hat{\mathbf{u}}_s] \in \mathbb{R}^{m \times q}$ have orthonormal columns satisfying $\mathbf{C}^\top \mathbf{u}_j = \alpha_j \mathbf{v}_j$ and $\hat{\mathbf{C}}^\top \hat{\mathbf{u}}_j = \hat{\alpha}_j \hat{\mathbf{v}}_j$ for $j = t, t+1, \dots, s$ and let $\mathbf{V} = [\mathbf{v}_t, \mathbf{v}_{t+1}, \dots, \mathbf{v}_s] \in \mathbb{R}^{n \times q}$ and $\hat{\mathbf{V}} = [\hat{\mathbf{v}}_t, \hat{\mathbf{v}}_{t+1}, \dots, \hat{\mathbf{v}}_s] \in \mathbb{R}^{n \times q}$ have orthonormal columns satisfying $\mathbf{C} \mathbf{v}_j = \alpha_j \mathbf{u}_j$ and $\hat{\mathbf{C}} \hat{\mathbf{v}}_j = \hat{\alpha}_j \hat{\mathbf{u}}_j$ for $j = t, t+1, \dots, s$. Then there exists an orthogonal matrix $\mathbf{O} \in \mathbb{R}^{q \times q}$ such that*

$$\|\hat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_F \leq \frac{2^{3/2}(2\alpha_1 + \|\hat{\mathbf{C}} - \mathbf{C}\|_2) \min(q^{1/2} \|\hat{\mathbf{C}} - \mathbf{C}\|_2, \|\hat{\mathbf{C}} - \mathbf{C}\|_F)}{\min(\alpha_{t-1}^2 - \alpha_t^2, \alpha_s^2 - \alpha_{s+1}^2)}$$

and the same upper bound holds for $\|\hat{\mathbf{V}} - \mathbf{V}\mathbf{O}\|_F$.

By letting $t = 1$ and $s = K$, using the assumption that $\text{rank}(\mathbf{C}_r) = K$ and applying Lemma 8, we have

$$\begin{aligned}\|\widehat{\mathbf{U}}_{\ell_r} - \overline{\mathbf{U}}_{\ell_r} \mathbf{O}_r\|_F &= \frac{2^{3/2} \sqrt{K} (2\sigma_{\max}(\mathbf{C}_r) + \|\widehat{\mathbf{C}}_r - \mathbf{C}_r\|_2) \|\widehat{\mathbf{C}}_r - \mathbf{C}_r\|_2}{\sigma_{\min}(\mathbf{C}_r)^2} \\ &\leq \frac{2^{3/2} \sqrt{K} 3\sigma_{\max}(\mathbf{C}_r) \|\widehat{\mathbf{C}}_r - \mathbf{C}_r\|_2}{\sigma_{\min}(\mathbf{C}_r)^2},\end{aligned}$$

where $\mathbf{O}_r \in \mathbb{R}^{K \times K}$ is orthogonal, the first equality is due to the fact that for any matrix $\mathbf{X}$ of rank at least K , $\|\mathbf{X}\|_F \leq \sqrt{\text{rank}(\mathbf{X})} \|\mathbf{X}\|_2$ and $\text{rank}(\mathbf{X}) \geq K$ and the last inequality is by using the assumption that $\|\widehat{\mathbf{C}}_r - \mathbf{C}_r\|_2 \leq \|\mathbf{C}_r\|_2 = \sigma_{\max}(\mathbf{C}_r)$.

Therefore, using the defined σ_1 and σ_K , we have for any $r \in [L-1]$,

$$\|\widehat{\mathbf{U}}_{\ell_r} - \overline{\mathbf{U}}_{\ell_r} \mathbf{O}_r\|_F \leq \frac{6\sqrt{2K} \sigma_1 \|\widehat{\mathbf{C}}_r - \mathbf{C}_r\|_2}{\sigma_K^2}. \quad (59)$$

Note that upper bound in (59) is valid for $\|\widehat{\mathbf{U}}_{\ell_{r+1}} - \overline{\mathbf{U}}_{\ell_{r+1}} \mathbf{O}_r\|_F$ as well.

E Proof of Lemma 2

Consider the below lemma:

Lemma 9 [51] *Let $\mathbf{A} \in \mathbb{R}^{N \times N}$ be the symmetric binary adjacency matrix of a random graph on N nodes in which the edges $\mathbf{A}(i, j)$ occur independently. Let $\overline{\mathbf{P}}$ be a matrix of size $N \times N$ such that the entries $\overline{\mathbf{P}}(i, j) \in [0, 1]$. Assume that $\mathbf{A}(i, j) \sim \text{Bernoulli}(\overline{\mathbf{P}}(i, j))$ for $i < j$, $\mathbf{A}(j, i) = \mathbf{A}(i, j)$ for $j > i$ and $\mathbf{A}(i, i) = 0$ for every $i \in [N]$. Set $\mathbb{E}[\mathbf{A}] = \overline{\mathbf{P}}$ and assume that $d \geq \max(\rho N, c_0 \log N)$ where $c_0 > 0$. Then, for any $t > 0$, there exists a constant $c_t = c(t, c_0)$ such that*

$$\|\mathbf{A} - \overline{\mathbf{P}}\|_2 \leq c_t \sqrt{d},$$

with probability greater than $1 - N^{-t}$.

To utilize Lemma 9 in our case, we analyze diagonal pattern (rightmost pattern in Fig. 1) and nondiagonal patterns (for e.g., the first and second patterns in Fig. 1) of EQP separately. For the diagonal pattern where $\ell_r = r$ and $\ell_{r+1} = m_r = r + 1$ for every r , we consider the below sub-adjacency matrix for any $r \in [L-1]$:

$$\tilde{\mathbf{A}}_r = \begin{bmatrix} \mathbf{A}_{r,r} & \mathbf{A}_{r,r+1} \\ \mathbf{A}_{r+1,r} & \mathbf{A}_{r+1,r+1} \end{bmatrix},$$

where $\tilde{\mathbf{A}}_r \in \mathbb{R}^{2N/L \times 2N/L}$.

Now, let us define

$$\tilde{\mathbf{P}}_r = \begin{bmatrix} \mathbf{P}_{r,r} & \mathbf{P}_{r,r+1} \\ \mathbf{P}_{r+1,r} & \mathbf{P}_{r+1,r+1} \end{bmatrix}.$$

It can be noted that $\mathbb{E}[\tilde{\mathbf{A}}_r] = \tilde{\mathbf{P}}_r - \text{diag}(\tilde{\mathbf{P}}_r)$. Denoting $\overline{\mathbf{P}}_r := \tilde{\mathbf{P}}_r - \text{diag}(\tilde{\mathbf{P}}_r)$ and $\rho := \max_{i,j \in [N]} \mathbf{P}(i, j)$, we have

$$\begin{aligned}\|\tilde{\mathbf{A}}_r - \tilde{\mathbf{P}}_r\|_2 &= \|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r + \overline{\mathbf{P}}_r - \tilde{\mathbf{P}}_r\|_2 \\ &= \|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r - \text{diag}(\tilde{\mathbf{P}}_r)\|_2 \\ &\leq \|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r\|_2 + \|\text{diag}(\tilde{\mathbf{P}}_r)\|_2 \\ &\leq \|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r\|_2 + \max_i \tilde{\mathbf{P}}_r(i, i) \\ &\leq \|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r\|_2 + \rho.\end{aligned} \quad (60)$$

In order to bound the first term $\|\tilde{\mathbf{A}}_r - \overline{\mathbf{P}}_r\|_2$, we can utilize Lemma 9. Hence, by assuming

$$d \geq \max\left(\frac{2\rho N}{L}, c_0 \log(2N/L)\right), \quad (61)$$

we apply Lemma 9 and (60) and obtain $\|\tilde{\mathbf{A}}_r - \tilde{\mathbf{P}}_r\|_2 \leq c_t \sqrt{d} + \rho$ with probability greater than $1 - (2N/L)^{-t}$. Next, we will rewrite the condition (61) in a simpler form. Suppose

$$\frac{2\rho N}{L} \geq c_0 \log(2N/L), \quad (62)$$

then the condition (61) becomes $d \geq \frac{2\rho N}{L}$. If we choose L to be small enough, then we can satisfy this condition on d , i.e., if we have

$$L \leq \frac{4\rho N}{d}, \quad (63)$$

then the condition $d \geq \frac{2\rho N}{L}$ gets automatically satisfied. Therefore, by using the assumptions (62)-(63) and fixing $t = 1$, we can apply Lemma 9 to obtain

$$\begin{aligned} \|\tilde{\mathbf{A}}_r - \tilde{\mathbf{P}}_r\|_2 &\leq c_1 \sqrt{d} + \rho \\ &\leq 2c_1 \sqrt{\rho N/L} + \rho \\ &\leq (2c_1 + 1) \sqrt{\rho N/L} \end{aligned} \quad (64)$$

with probability at least $1 - (2N/L)^{-1}$. In the above, the first inequality is by using the relation (60), the second inequality is by employing the assumed bound on L given in (63) and the third inequality is by using the fact that $\rho \leq 1$.

Then, we have

$$\begin{aligned} \|\hat{\mathbf{C}}_r - \mathbf{C}_r\|_2 &= \left\| [\mathbf{A}_{r,r+1}^\top, \mathbf{A}_{r+1,r+1}^\top]^\top - [\mathbf{P}_{r,r+1}^\top, \mathbf{P}_{r+1,r+1}^\top]^\top \right\|_2 \\ &\leq \|\tilde{\mathbf{A}}_r - \tilde{\mathbf{P}}_r\|_2 \leq (2c_1 + 1) \sqrt{\rho N/L} \end{aligned} \quad (65)$$

holds with probability greater than $1 - (2N/L)^{-1}$, where we have applied (64) to obtain the last inequality.

In order to apply Lemma 9 for non-diagonal patterns in EQP, we start by defining $\mathbf{A}_r^*$ and $\mathbf{P}_r^*$ in the following way:

$$\mathbf{A}_r^* = \begin{bmatrix} \mathbf{0} & \mathbf{0} & \mathbf{A}_{\ell_r, m_r} \\ \mathbf{0} & \mathbf{0} & \mathbf{A}_{\ell_{r+1}, m_r} \\ \mathbf{A}_{\ell_r, m_r}^\top & \mathbf{A}_{\ell_{r+1}, m_r}^\top & \mathbf{0} \end{bmatrix},$$

$$\mathbf{P}_r^* = \begin{bmatrix} \mathbf{0} & \mathbf{0} & \mathbf{P}_{\ell_r, m_r} \\ \mathbf{0} & \mathbf{0} & \mathbf{P}_{\ell_{r+1}, m_r} \\ \mathbf{P}_{\ell_r, m_r}^\top & \mathbf{P}_{\ell_{r+1}, m_r}^\top & \mathbf{0} \end{bmatrix}$$

where $\mathbf{0}$ represents a zero matrix which is of size $N/L \times N/L$. With this definition, we can see that $\mathbf{A}_r^*(i, j) \sim \text{Bernoulli}(\mathbf{P}_r^*(i, j))$ for $i < j$, $\mathbf{A}_r^*(j, i) = \mathbf{A}_r^*(i, j)$ for $j > i$, $\mathbf{A}_r^*(i, i) = \mathbf{0}$ for every i and $\mathbb{E}[\mathbf{A}_r^*] = \mathbf{P}_r^*$. Therefore, we can apply Lemma 9 by fixing $t = 1$ to result in

$$\|\mathbf{A}_r^* - \mathbf{P}_r^*\|_2 \leq c_1 \sqrt{d}, \quad (66)$$

with probability greater than $1 - (3N/L)^{-1}$. Similar as before, we can derive the conditions in order to obtain (66). It can be easily shown that the sufficient conditions to be satisfied are

$$\frac{3\rho N}{L} \geq c_0 \log(3N/L), \quad L \leq \frac{6\rho N}{d}. \quad (67)$$

Then, for non-diagonal EQP patterns, we have

$$\begin{aligned}\|\hat{C}_r - C_r\|_2 &= \left\| [\mathbf{A}_{\ell_r, m_r}^\top, \mathbf{A}_{\ell_{r+1}, m_r}^\top]^\top - [\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_{r+1}, m_r}^\top]^\top \right\|_2 \\ &\leq \|\mathbf{A}_r^* - \mathbf{P}_r^*\|_2 \leq c_1 \sqrt{d} \leq \sqrt{6} c_1 \sqrt{\rho N / L},\end{aligned}\quad (68)$$

where we obtained the last inequality by applying the assumed bound on L as given by the second condition in (67).

Hence, from (65) and (68), we can see that $\|\hat{C}_r - C_r\|_2$ is bounded for any patterns in EQP with probability of at least $1 - (2N/L)^{-1}$ if the below conditions are satisfied.

$$\frac{3\rho N}{L} \geq c_0 \log(3N/L), \quad L \leq \frac{4\rho N}{d}.$$

F Proof of Lemma 3

Consider the below relation for any $r \in [L - 1]$:

$$\begin{aligned}\sigma_{\max}^2([\mathbf{M}_r, \mathbf{M}_{r+1}]) &= \|[\mathbf{M}_r, \mathbf{M}_{r+1}]\|_2^2 \leq \|[\mathbf{M}_r, \mathbf{M}_{r+1}]\|_{\text{F}}^2 \\ &= \|\mathbf{M}_r\|_{\text{F}}^2 + \|\mathbf{M}_{r+1}\|_{\text{F}}^2 \\ &\leq K\|\mathbf{M}_r\|_2^2 + K\|\mathbf{M}_{r+1}\|_2^2 \leq 2K\alpha_{\max}^2,\end{aligned}$$

where we have utilized the norm equivalence for the first and second inequalities and applied the definition of $\alpha_{\max}$ for the last inequality. Hence, we have

$$\sigma_{\max}([\mathbf{M}_r, \mathbf{M}_{r+1}]) \leq \sqrt{2K}\alpha_{\max} := \beta_{\max}.$$

Next, consider the below for any $\mathbf{x} \in \mathbb{R}^K$ such that $\|\mathbf{x}\|_2 = 1$:

$$\begin{aligned}\|[\mathbf{M}_r, \mathbf{M}_{r+1}]^\top \mathbf{x}\|_2^2 &= \|\mathbf{M}_r^\top \mathbf{x}\|_2^2 + \|\mathbf{M}_{r+1}^\top \mathbf{x}\|_2^2 \\ \Rightarrow \min_{\mathbf{x}} \|[\mathbf{M}_r, \mathbf{M}_{r+1}]^\top \mathbf{x}\|_2^2 &= \min_{\mathbf{x}} \|\mathbf{M}_r^\top \mathbf{x}\|_2^2 + \min_{\mathbf{x}} \|\mathbf{M}_{r+1}^\top \mathbf{x}\|_2^2 \\ \Rightarrow \min_{\mathbf{x}} \|[\mathbf{M}_r, \mathbf{M}_{r+1}]^\top \mathbf{x}\|_2^2 &\geq \min_{\mathbf{x}} \|\mathbf{M}_r^\top \mathbf{x}\|_2^2 \\ \Rightarrow \sigma_{\min}([\mathbf{M}_r, \mathbf{M}_{r+1}]) &\geq \sigma_{\min}(\mathbf{M}_r) \geq \alpha_{\min} := \beta_{\min},\end{aligned}$$

where we have utilized the definition of $\alpha_{\min}$ for the last inequality.

G Proof of Lemma 4

We utilize the following classic result to characterize the pseudo-inverse:

Lemma 10 [64] Consider any matrices $\mathbf{Y}, \mathbf{Z}, \mathbf{E} \in \mathbb{R}^{m \times n}$ such that $\mathbf{Z} = \mathbf{Y} + \mathbf{E}$. Suppose $\text{rank}(\mathbf{Z}) = \text{rank}(\mathbf{Y})$. Then, we have

$$\|\mathbf{Z}^\dagger - \mathbf{Y}^\dagger\|_2 \leq \sqrt{2}\|\mathbf{Y}^\dagger\|_2\|\mathbf{Z}^\dagger\|_2\|\mathbf{E}\|_2.$$

By fixing $\mathbf{Z} := \hat{\mathbf{U}}_r^\dagger$ and $\mathbf{Y} = (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger$, we can apply Lemma 10 to obtain

$$\begin{aligned}\|\hat{\mathbf{U}}_r^\dagger - (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 &\leq \sqrt{2}\|(\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2\|\hat{\mathbf{U}}_r^\dagger\|_2\|\hat{\mathbf{U}}_r - (\bar{\mathbf{U}}_r \mathbf{O}_r)\|_2.\end{aligned}\quad (69)$$

Regarding the first term in the R.H.S. of the above equation, we have

$$\begin{aligned}\|(\overline{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 &= \sigma_{\max}((\overline{\mathbf{U}}_r \mathbf{O}_r)^\dagger) \\ &= \frac{1}{\sigma_{\min}(\overline{\mathbf{U}}_r \mathbf{O}_r)} = \frac{1}{\sigma_{\min}(\overline{\mathbf{U}}_r)},\end{aligned}\quad (70)$$

where the last equality is due to the orthogonality of $\mathbf{O}_r$.

Similarly, we have

$$\|\widehat{\mathbf{U}}_r^\dagger\|_2 = \sigma_{\max}(\widehat{\mathbf{U}}_r^\dagger) = \frac{1}{\sigma_{\min}(\widehat{\mathbf{U}}_r)}.\quad (71)$$

To proceed, we aim to get a lower bound for $\sigma_{\min}(\widehat{\mathbf{U}}_r)$. For this, we have the following result:

Lemma 11 *The below relation holds true if $\|\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r\|_2 \leq \sigma_{\min}(\overline{\mathbf{U}}_r)/2$:*

$$\sigma_{\min}(\widehat{\mathbf{U}}_r) \geq \sigma_{\min}(\overline{\mathbf{U}}_r)/2.$$

Proof: Consider the below set of relations for any vector $\mathbf{x} \in \mathbb{R}^K$ satisfying $\|\mathbf{x}\| = 1$:

$$\begin{aligned}\|\widehat{\mathbf{U}}_r \mathbf{x}\|_2 &= \|\overline{\mathbf{U}}_r \mathbf{O}_r \mathbf{x} + (\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r) \mathbf{x}\|_2 \\ &\geq \|\overline{\mathbf{U}}_r \mathbf{O}_r \mathbf{x}\|_2 - \|(\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r) \mathbf{x}\|_2, \\ \Rightarrow \min_{\mathbf{x}} \|\widehat{\mathbf{U}}_r \mathbf{x}\|_2 &\geq \min_{\mathbf{x}} \|\overline{\mathbf{U}}_r \mathbf{O}_r \mathbf{x}\|_2 - \max_{\mathbf{x}} \|(\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r) \mathbf{x}\|_2, \\ \Rightarrow \sigma_{\min}(\widehat{\mathbf{U}}_r) &\geq \sigma_{\min}(\overline{\mathbf{U}}_r) - \|\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r\|_2,\end{aligned}$$

where the first inequality is by applying the triangle inequality.

Using the assumption that $\|\widehat{\mathbf{U}}_r - \overline{\mathbf{U}}_r \mathbf{O}_r\|_2 \leq \sigma_{\min}(\overline{\mathbf{U}}_r)/2$, we get $\sigma_{\min}(\widehat{\mathbf{U}}_r) \geq \sigma_{\min}(\overline{\mathbf{U}}_r)/2$. $\square$

Applying (70), (71) and Lemma 11 in (69), we get

$$\|\widehat{\mathbf{U}}_r^\dagger - (\overline{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 \leq \frac{2\sqrt{2}}{\sigma_{\min}^2(\overline{\mathbf{U}}_r)} \|\widehat{\mathbf{U}}_r - (\overline{\mathbf{U}}_r \mathbf{O}_r)\|_2.\quad (72)$$

The final step is to characterize $\sigma_{\min}(\overline{\mathbf{U}}_r)$. For this, we utilize the following result:

Lemma 12 *For every r , we have*

$$\begin{aligned}\sigma_{\min}(\overline{\mathbf{U}}_r) &\geq \frac{\sigma_{\min}(\mathbf{M}_r)}{\sigma_{\max}([\mathbf{M}_r, \mathbf{M}_{r+1}])} = \frac{\alpha_{\min}}{\beta_{\max}}, \\ \sigma_{\max}(\overline{\mathbf{U}}_r) &\leq \frac{\sigma_{\max}(\mathbf{M}_r)}{\sigma_{\min}([\mathbf{M}_r, \mathbf{M}_{r+1}])} = \frac{\alpha_{\max}}{\beta_{\min}}.\end{aligned}$$

Proof: Recall the below set of relations (assuming $\ell_r = r$ for every r):

$$\mathbf{C}_r = [\mathbf{P}_{r,m_r}^\top, \mathbf{P}_{r+1,m_r}^\top]^\top = [\mathbf{M}_r, \mathbf{M}_{r+1}]^\top \mathbf{B} \mathbf{M}_{m_r}.\quad (73)$$

The top- K SVD of $\mathbf{C}_r$ results in the below relation:

$$\mathbf{C}_r = [\overline{\mathbf{U}}_r^\top, \overline{\mathbf{U}}_{r+1}^\top]^\top \boldsymbol{\Sigma}_r \mathbf{V}_{m_r}^\top.\quad (74)$$

From (73) and (74), we get that there exists a nonsingular matrix $\overline{\mathbf{G}}_r$ such that

$$[\overline{\mathbf{U}}_r^\top, \overline{\mathbf{U}}_{r+1}^\top]^\top = [\mathbf{M}_r, \mathbf{M}_{r+1}]^\top \overline{\mathbf{G}}_r,\quad (75)$$

where the matrix $[\overline{\mathbf{U}}_r^\top, \overline{\mathbf{U}}_{r+1}^\top]^\top$ is semi-orthogonal. We invoke the below fact to proceed further.

Fact 1 Consider the equation $\bar{\mathbf{U}} = \bar{\mathbf{M}}^\top \bar{\mathbf{G}}$ where the matrix $\bar{\mathbf{U}}$ is tall and is semi-orthogonal. Then the below holds:

$$\sigma_{\max}(\bar{\mathbf{G}}) = 1/\sigma_{\min}(\bar{\mathbf{M}}) \quad \text{and} \quad \sigma_{\min}(\bar{\mathbf{G}}) = 1/\sigma_{\max}(\bar{\mathbf{M}}).$$

Proof: Since $\bar{\mathbf{U}}$ is a semi-orthogonal matrix, we have

$$\begin{aligned} \bar{\mathbf{U}}^\top \bar{\mathbf{U}} &= \mathbf{I}, \\ \Rightarrow \bar{\mathbf{G}}^\top \bar{\mathbf{M}} \bar{\mathbf{M}}^\top \bar{\mathbf{G}} &= \mathbf{I}, \\ \Rightarrow \bar{\mathbf{M}} \bar{\mathbf{M}}^\top &= \bar{\mathbf{G}}^{-\top} \bar{\mathbf{G}}^{-1} = (\bar{\mathbf{G}} \bar{\mathbf{G}}^\top)^{-1}. \end{aligned}$$

The above relation implies that $\sigma_{\max}(\bar{\mathbf{G}}) = 1/\sigma_{\min}(\bar{\mathbf{M}})$, $\sigma_{\min}(\bar{\mathbf{G}}) = 1/\sigma_{\max}(\bar{\mathbf{M}})$ hold true. $\square$

Due to Fact 1, from (75), we get

$$\sigma_{\max}(\bar{\mathbf{G}}_r) = \frac{1}{\sigma_{\min}([\mathbf{M}_r, \mathbf{M}_{r+1}])}, \quad (76a)$$

$$\sigma_{\min}(\bar{\mathbf{G}}_r) = \frac{1}{\sigma_{\max}([\mathbf{M}_r, \mathbf{M}_{r+1}])}. \quad (76b)$$

Next we consider the below relation obtained from (75):

$$\bar{\mathbf{U}}_r = \mathbf{M}_r^\top \bar{\mathbf{G}}_r.$$

Since $\mathbf{M}_r$ is full row-rank, we have

$$\begin{aligned} \sigma_{\min}(\bar{\mathbf{U}}_r) &= \min_{\|\mathbf{x}\|_2=1} \|\mathbf{M}_r^\top \bar{\mathbf{G}}_r \mathbf{x}\|_2 \geq \min_{\|\mathbf{x}\|_2=1} \sigma_{\min}(\mathbf{M}_r) \|\bar{\mathbf{G}}_r \mathbf{x}\|_2 \\ &= \sigma_{\min}(\mathbf{M}_r) \min_{\|\mathbf{x}\|_2=1} \|\bar{\mathbf{G}}_r \mathbf{x}\|_2 = \sigma_{\min}(\mathbf{M}_r) \sigma_{\min}(\bar{\mathbf{G}}_r) \\ &= \frac{\sigma_{\min}(\mathbf{M}_r)}{\sigma_{\max}([\mathbf{M}_r, \mathbf{M}_{r+1}])} \geq \frac{\alpha_{\min}}{\beta_{\max}}, \end{aligned}$$

where we have applied (76) to obtain the last equality and used Lemma 3 for the last inequality.

Similarly, we can easily have,

$$\begin{aligned} \sigma_{\max}(\bar{\mathbf{U}}_r) &\leq \sigma_{\max}(\bar{\mathbf{G}}_r) \sigma_{\max}(\mathbf{M}_r) \\ &= \frac{\sigma_{\max}(\mathbf{M}_r)}{\sigma_{\min}([\mathbf{M}_r, \mathbf{M}_{r+1}])} \leq \frac{\alpha_{\max}}{\beta_{\min}}. \end{aligned}$$

$\square$

Combining Lemma 12 with (72), we have

$$\|\hat{\bar{\mathbf{U}}}_r^\dagger - (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 \leq \frac{2\sqrt{2}\beta_{\max}^2}{\alpha_{\min}^2} \|\hat{\bar{\mathbf{U}}}_r - (\bar{\mathbf{U}}_r \mathbf{O}_r)\|_2 \quad (77)$$

H Proof of Lemma 5

For simpler notations, let us define $\mathbf{G}_1 := \bar{\mathbf{U}}_{r+1} \mathbf{O}_r$, $\mathbf{G}_2 := (\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger$ and $\mathbf{G}_3 := \mathbf{U}_r \mathbf{O}$. Also we define $\hat{\mathbf{G}}_1 := \hat{\bar{\mathbf{U}}}_{r+1}$, $\hat{\mathbf{G}}_2 := (\hat{\bar{\mathbf{U}}}_r)^\dagger$ and $\mathbf{G}_3 := \hat{\mathbf{U}}_r \mathbf{O}$. With these, consider the below:

$$\begin{aligned}
\|\hat{\mathbf{G}}_1 \hat{\mathbf{G}}_2 \hat{\mathbf{G}}_3 - \mathbf{G}_1 \mathbf{G}_2 \mathbf{G}_3\|_2 &= \|\hat{\mathbf{G}}_1 \hat{\mathbf{G}}_2 \hat{\mathbf{G}}_3 - \mathbf{G}_1 \mathbf{G}_2 \mathbf{G}_3 - \hat{\mathbf{G}}_1 \mathbf{G}_2 \mathbf{G}_3 + \hat{\mathbf{G}}_1 \mathbf{G}_2 \mathbf{G}_3\|_2 \\
&= \|(\hat{\mathbf{G}}_1 - \mathbf{G}_1) \mathbf{G}_2 \mathbf{G}_3 + \hat{\mathbf{G}}_1 (\hat{\mathbf{G}}_2 \hat{\mathbf{G}}_3 - \mathbf{G}_2 \mathbf{G}_3)\|_2 \\
&\leq \|(\hat{\mathbf{G}}_1 - \mathbf{G}_1) \mathbf{G}_2 \mathbf{G}_3\|_2 + \|\hat{\mathbf{G}}_1 (\hat{\mathbf{G}}_2 \hat{\mathbf{G}}_3 - \mathbf{G}_2 \mathbf{G}_3)\|_2 \\
&= \|(\hat{\mathbf{G}}_1 - \mathbf{G}_1) \mathbf{G}_2 \mathbf{G}_3\|_2 \\
&\quad + \|\hat{\mathbf{G}}_1 (\hat{\mathbf{G}}_2 - \mathbf{G}_2) \mathbf{G}_3 + \hat{\mathbf{G}}_1 \hat{\mathbf{G}}_2 (\hat{\mathbf{G}}_3 - \mathbf{G}_3)\|_2 \\
&\leq \|\mathbf{G}_2\|_2 \|\mathbf{G}_3\|_2 \|\hat{\mathbf{G}}_1 - \mathbf{G}_1\|_2 + \|\hat{\mathbf{G}}_1\|_2 \|\mathbf{G}_3\|_2 \|\hat{\mathbf{G}}_2 - \mathbf{G}_2\|_2 \\
&\quad + \|\hat{\mathbf{G}}_1\|_2 \|\hat{\mathbf{G}}_2\|_2 \|\hat{\mathbf{G}}_3 - \mathbf{G}_3\|_2, \tag{78}
\end{aligned}$$

where we have applied triangle inequality and the fact that $\|\mathbf{Y}\mathbf{Z}\|_2 \leq \|\mathbf{Y}\|_2 \|\mathbf{Z}\|_2$ to obtain the first and last inequalities, respectively.

We introduce the following lemma and then proceed to characterize $\|\hat{\mathbf{G}}_1\|_2$, $\|\mathbf{G}_2\|_2$, $\|\hat{\mathbf{G}}_2\|_2$ and $\|\mathbf{G}_3\|_2$.

Lemma 13 *The below relation holds for every r :*

$$\sigma_{\max}(\mathbf{U}_r) \leq \frac{\alpha_{\max}}{\beta_{\min}}.$$

The proof of the lemma is provided in Sec. 4.

- **Bound for $\|\hat{\mathbf{G}}_1\|_2 = \|\hat{\bar{\mathbf{U}}}_{r+1}\|_2$**

$$\begin{aligned}
\|\hat{\mathbf{G}}_1\|_2 &= \|\hat{\bar{\mathbf{U}}}_{r+1}\|_2 = \|\hat{\bar{\mathbf{U}}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r + \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 \\
&\leq \|\hat{\bar{\mathbf{U}}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 + \|\bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 \\
&\leq \sigma_{\max}(\bar{\mathbf{U}}_{r+1}) + \sigma_{\max}(\bar{\mathbf{U}}_{r+1}) \\
&= 2\sigma_{\max}(\bar{\mathbf{U}}_{r+1}) \leq 2\frac{\alpha_{\max}}{\beta_{\min}}.
\end{aligned}$$

where we have used triangle inequality for the first inequality, used the assumption that $\|\hat{\bar{\mathbf{U}}}_{r+1} - \bar{\mathbf{U}}_{r+1} \mathbf{O}_r\|_2 \leq \sigma_{\min}(\bar{\mathbf{U}}_{r+1})/2$ for the second inequality and invoked Lemma 12 for the last inequality.

- **Bound for $\|\mathbf{G}_2\|_2 = \|(\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2$**

$$\|\mathbf{G}_2\|_2 = \|(\bar{\mathbf{U}}_r \mathbf{O}_r)^\dagger\|_2 = 1/\sigma_{\min}(\bar{\mathbf{U}}_r) \leq \frac{\beta_{\max}}{\alpha_{\min}},$$

where we have applied Lemma 12 for the last inequality.

- **Bound for $\|\hat{\mathbf{G}}_2\|_2 = \|\hat{\bar{\mathbf{U}}}_r^\dagger\|_2$**

$$\|\hat{\mathbf{G}}_2\|_2 = \|\hat{\bar{\mathbf{U}}}_r^\dagger\|_2 = 1/\sigma_{\min}(\hat{\bar{\mathbf{U}}}_r) \leq 2/\sigma_{\min}(\bar{\mathbf{U}}_r) \leq 2\frac{\beta_{\max}}{\alpha_{\min}},$$

where we have used applied Lemma 11 by using the assumption that $\|\hat{\bar{\mathbf{U}}}_r - \bar{\mathbf{U}}_r \mathbf{O}_r\|_2 \leq \sigma_{\min}(\bar{\mathbf{U}}_r)/2$ for the first inequality and invoked Lemma 12 for the last inequality.

- **Bound for $\|\mathbf{G}_3\|_2 = \|\mathbf{U}_r \mathbf{O}\|_2$**

$$\|\mathbf{G}_3\|_2 = \|\mathbf{U}_r \mathbf{O}\|_2 = \|\mathbf{U}_r\|_2 \leq \frac{\alpha_{\max}}{\beta_{\min}},$$

where we have invoked Lemma 13 for the inequality.

Applying these upper bounds in (78), we finally have

$$\begin{aligned} \|\widehat{\mathbf{G}}_1 \widehat{\mathbf{G}}_2 \widehat{\mathbf{G}}_3 - \mathbf{G}_1 \mathbf{G}_2 \mathbf{G}_3\|_2 &\leq \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\widehat{\mathbf{G}}_1 - \mathbf{G}_1\|_2 + 2 \frac{\alpha_{\max}}{\beta_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\widehat{\mathbf{G}}_2 - \mathbf{G}_2\|_2 \\ &\quad + 4 \frac{\beta_{\max}}{\alpha_{\min}} \frac{\alpha_{\max}}{\beta_{\min}} \|\widehat{\mathbf{G}}_3 - \mathbf{G}_3\|_2. \end{aligned}$$

I Proof of Lemma 13

Theorem 1 shows that there exists a certain nonsingular $\mathbf{G}$ such that for each true basis $\mathbf{U}_r$ for $r \in [L]$, we have:

$$\mathbf{U}_r = \mathbf{M}_r^\top \mathbf{G}. \quad (79)$$

Recall that, under noiseless case, Algorithm 1 first directly estimates $\mathbf{U}_T$ and $\mathbf{U}_{T+1}$ by performing the top- K SVD to $\mathbf{C}_T = [\mathbf{P}_{T,m_T}^\top, \mathbf{P}_{T+1,m_T}^\top]^\top$ (assuming $\ell_r = r$ for every r) where

$$\mathbf{C}_T = [\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top \boldsymbol{\Sigma}_T \mathbf{V}_{m_T}^\top. \quad (80)$$

The above implies that,

$$[\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top = [\mathbf{M}_T, \mathbf{M}_{T+1}]^\top \mathbf{G}$$

holds where $[\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]$ is semi-orthogonal. By utilizing Fact 1, we get $\sigma_{\max}(\mathbf{G}) = 1/\sigma_{\min}([\mathbf{M}_T, \mathbf{M}_{T+1}])$ and $\sigma_{\min}(\mathbf{G}) = 1/\sigma_{\max}([\mathbf{M}_T, \mathbf{M}_{T+1}])$. Using this result, from (79), we have

$$\begin{aligned} \sigma_{\max}(\mathbf{U}_r) &\leq \sigma_{\max}(\mathbf{G}) \sigma_{\max}(\mathbf{M}_r) \\ &\leq \frac{\sigma_{\max}(\mathbf{M}_r)}{\sigma_{\min}([\mathbf{M}_T, \mathbf{M}_{T+1}])} \leq \frac{\alpha_{\max}}{\beta_{\min}}, \end{aligned}$$

where we have utilized Lemma 3 to obtain the last inequality.

J Proof of Lemma 6

Let us first consider the below:

$$\mathbf{C}_r = [\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_r+1, m_r}^\top]^\top = [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_r+1}]^\top \mathbf{B} \mathbf{M}_{m_r}.$$

Consider the K -th eigenvalue of the symmetric matrix $\mathbf{C}_r^\top \mathbf{C}_r$:

$$\begin{aligned} \lambda_K(\mathbf{C}_r^\top \mathbf{C}_r) &= \lambda_K(\underbrace{\mathbf{M}_{m_r}^\top \mathbf{B} [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_r+1}] [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_r+1}]^\top}_{\mathbf{W}_1 \in \mathbb{R}^{N/L \times K}} \underbrace{\mathbf{B} \mathbf{M}_{m_r}}_{\mathbf{W}_2 \in \mathbb{R}^{K \times N/L}}), \end{aligned} \quad (81)$$

where λ_K denotes the K th eigenvalue with $\lambda_1 \geq \dots \geq \lambda_K$.

We utilize the following fact to proceed further:

Lemma 14 (Theorem 1.3.22) [65] Consider two matrices $\mathbf{W}_1 \in \mathbb{R}^{N \times K}$ and $\mathbf{W}_2 \in \mathbb{R}^{K \times N}$. Let $\lambda_1, \dots, \lambda_K$ be the nonzero eigenvalues of the matrix $\mathbf{W}_1 \mathbf{W}_2$. Then, the matrix $\mathbf{W}_2 \mathbf{W}_1$ also holds the same set of eigenvalues.

Combining Lemma [14] and (81), we have

$$\lambda_K(\mathbf{C}_r^\top \mathbf{C}_r) = \lambda_K(\mathbf{B} \mathbf{M}_{m_r} \mathbf{M}_{m_r}^\top \mathbf{B} [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}] [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}]^\top). \quad (82)$$

Next, we consider the following matrix

$$\mathbf{H}_r := \mathbf{M}_{m_r} \mathbf{M}_{m_r}^\top = \sum_{n=1}^{N/L} \mathbf{M}_{m_r}(:, n) \mathbf{M}_{m_r}(:, n)^\top. \quad (83)$$

To proceed, we have the lemma from [21] to bound $\lambda_K(\mathbf{H}_r)$ and $\lambda_1(\mathbf{H}_r)$.

Lemma 15 [21] Assume that the columns of $\mathbf{M}$ are generated from any continuous distribution. Then, there exists positive constants c and C depending only on distribution of the columns of $\mathbf{M}$ such that

$$\Pr(\lambda_K(\mathbf{H}_r) \leq cN/L) \leq K \exp(-cN/4L),$$

$$\Pr(\lambda_1(\mathbf{H}_r) \geq CN/L) \leq \frac{K}{2^{CN/L}}.$$

To simplify the bound on the probability, let us find the conditions at which the following is satisfied:

$$K \exp(-cN/4L) \leq (N/L)^{-1}, \quad (84)$$

$$\frac{K}{2^{CN/L}} \leq (N/L)^{-1}. \quad (85)$$

The conditions (84) and (85) are equivalent to the following:

$$(N/L) \geq \frac{4}{c} \log(NK/L), \quad (86)$$

$$(N/L) \geq \frac{\log_2 e}{C} \log(NK/L). \quad (87)$$

We can immediately see that if (86) is satisfied, (87) is also satisfied since $c \leq C$.

Therefore, we get that if $(N/L) \geq \frac{4}{c} \log(NK/L)$, with probability of at least $1 - (N/L)^{-1}$,

$$\lambda_K(\mathbf{M}_{m_r} \mathbf{M}_{m_r}^\top) \geq cN/L, \quad (88a)$$

$$\lambda_1(\mathbf{M}_{m_r} \mathbf{M}_{m_r}^\top) \leq CN/L. \quad (88b)$$

Using the above result, we can also have

$$\lambda_K([\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}] [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}]^\top) \geq 2cN/L, \quad (89a)$$

$$\lambda_1([\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}] [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}]^\top) \leq 2CN/L, \quad (89b)$$

each holding with probability at least $1 - (2N/L)^{-1}$, if $(2N/L) \geq \frac{4}{c} \log(2NK/L)$.

Thus, applying (88) and (89) in (82), we get

$$\begin{aligned} \lambda_K(\mathbf{C}_r^\top \mathbf{C}_r) &\geq \lambda_K^2(\mathbf{B}) \lambda_K([\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}] [\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}]^\top) \lambda_K(\mathbf{M}_{m_r} \mathbf{M}_{m_r}^\top) \\ &\geq \lambda_K^2(\mathbf{B}) c(2N/L) c(N/L) = 2\lambda_K^2(\mathbf{B}) c^2 N^2 / L^2, \\ &= 2\sigma_{\min}^2(\mathbf{B}) c^2 N^2 / L^2, \end{aligned}$$

with probability at least $1 - (2N/3L)^{-1}$ where we have used the facts that $\lambda_1((\mathbf{A}\mathbf{B})^{-1}) \leq \lambda_1(\mathbf{A}^{-1})\lambda_1(\mathbf{B}^{-1})$ and $\lambda_1(\mathbf{A}^{-1}) = 1/\lambda_K(\mathbf{A})$ for obtaining the first inequality and used $\lambda_K(\mathbf{B}) = \sigma_{\min}(\mathbf{B})$ for the last equality.

Similarly, we can get the below with probability at least $1 - (2N/3L)^{-1}$:

$$\begin{aligned}\lambda_1(\mathbf{C}_r^\top \mathbf{C}_r) &\leq \lambda_1^2(\mathbf{B})\lambda_1([\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}][\mathbf{M}_{\ell_r}, \mathbf{M}_{\ell_{r+1}}]^\top)\lambda_1(\mathbf{M}_{m_r}\mathbf{M}_{m_r}^\top) \\ &\leq \lambda_1^2(\mathbf{B})C(2N/L)C(N/L) = 2\lambda_1^2(\mathbf{B})C^2N^2/L^2, \\ &= 2\sigma_{\max}^2(\mathbf{B})C^2N^2/L^2.\end{aligned}$$

Therefore, by taking union bound over every $r \in [L-1]$, with probability at least $1 - (3L(L-1)/N)$, the following equations hold simultaneously for every r :

$$\sigma_{\min}([\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_{r+1}, m_r}^\top]^\top) = \sigma_{\min}(\mathbf{C}_r) = \sqrt{\lambda_K(\mathbf{C}_r^\top \mathbf{C}_r)} \geq \sqrt{2}\sigma_{\min}(\mathbf{B})cN/L, \quad (90)$$

$$\sigma_{\max}([\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_{r+1}, m_r}^\top]^\top) = \sigma_{\max}(\mathbf{C}_r) = \sqrt{\lambda_1(\mathbf{C}_r^\top \mathbf{C}_r)} \leq \sqrt{2}\sigma_{\max}(\mathbf{B})cN/L. \quad (91)$$

The equations (90) and (91) immediately follow that the below holds with probability at least $1 - (3L(L-1)/N)$:

$$\begin{aligned}\sigma_K &= \min_{r \in \{1, \dots, L-1\}} \sigma_{\min}([\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_{r+1}, m_r}^\top]^\top) \\ &\geq \sqrt{2}\sigma_{\min}(\mathbf{B})cN/L, \\ \sigma_1 &= \max_{r \in \{1, \dots, L-1\}} \sigma_{\max}([\mathbf{P}_{\ell_r, m_r}^\top, \mathbf{P}_{\ell_{r+1}, m_r}^\top]^\top) \\ &\leq \sqrt{2}\sigma_{\max}(\mathbf{B})cN/L.\end{aligned}$$

Combining the conditions to obtain (88) and (89), we get that the above results hold true if $(N/L) \geq \frac{4}{c} \log(NK/L)$.

K Proof of Lemma 7

We consider the given noisy NMF model:

$$\hat{\mathbf{U}}^\top = \mathbf{O}^\top \mathbf{G}^\top \mathbf{M} + \mathbf{N}.$$

It is given that the assumptions in Proposition 1 hold true. This immediately means that $\mathbf{M}$ satisfies ε -separability (see Definition 1), i.e., there exists a set of indices $\mathbf{A} := \{q_1, \dots, q_K\}$ such that $\mathbf{M}(:, \mathbf{A}) = \mathbf{I}_K + \mathbf{E}$, where $\mathbf{E} \in \mathbb{R}^{K \times K}$ is the error matrix with $\|\mathbf{E}(:, k)\|_2 \leq \varepsilon$ and $\mathbf{I}_K$ is the identity matrix of size $K \times K$. Without loss of generality, we let $\mathbf{A} = \{1, \dots, K\}$. Therefore, we have

$$\begin{aligned}\hat{\mathbf{U}}^\top &= \mathbf{O}^\top \mathbf{G}^\top \mathbf{M} + \mathbf{N} = \mathbf{O}^\top \mathbf{G}^\top [\mathbf{I}_K + \mathbf{E}, \tilde{\mathbf{M}}] + \mathbf{N} \\ &= \mathbf{O}^\top \mathbf{G}^\top [\mathbf{I}_K, \tilde{\mathbf{M}}] + [\mathbf{O}^\top \mathbf{G}^\top \mathbf{E}, \mathbf{0}] + \mathbf{N},\end{aligned}$$

where the zero matrix $\mathbf{0}$ has the same dimension of $\tilde{\mathbf{M}}$.

Let $\tilde{\mathbf{N}} = [\mathbf{O}^\top \mathbf{G}^\top \mathbf{E}, \mathbf{0}] + \mathbf{N}$. This gives us the noisy NMF model as

$$\hat{\mathbf{U}}^\top = \underbrace{\mathbf{O}^\top \mathbf{G}^\top}_{\mathbf{W}} \underbrace{[\mathbf{I}_K, \tilde{\mathbf{M}}]}_{\mathbf{M}} + \tilde{\mathbf{N}}. \quad (92)$$

Then, the below holds for any $k \in [K]$:

$$\begin{aligned}\|\tilde{\mathbf{N}}(:, k)\|_2 &\leq \|\mathbf{G}\mathbf{O}\|_2 \|\mathbf{E}(:, k)\|_2 + \|\mathbf{N}(:, k)\|_2, \\ &\leq \sigma_{\max}(\mathbf{G})\varepsilon + \zeta := \tilde{\zeta}(\mathbf{G}),\end{aligned} \quad (93)$$

where the first inequality is by applying triangle inequality and the second inequality is by using the orthogonality of $\mathbf{O}$ and also by using the given assumption $\|\mathbf{N}(:, k)\|_2 \leq \zeta$.

We utilize the below lemma from [35] which characterizes the estimation accuracy of SPA under the NMF model (92).

Lemma 16 [35] Consider the noisy model in (92) and suppose that $\mathbf{W}$ is full rank and $\mathbf{M}$ satisfies the simplex constraints, i.e., $\mathbf{M} \geq \mathbf{0}, \mathbf{1}^\top \mathbf{M} = \mathbf{1}^\top$. Let each column of $\tilde{\mathbf{N}}$ satisfies the condition $\|\tilde{\mathbf{N}}(:, k)\|_2 \leq \tilde{\zeta}$ with

$$\tilde{\zeta} \leq \sigma_{\min}(\mathbf{W}) \varrho \tilde{\kappa}^{-1}(\mathbf{W}). \quad (94)$$

Then the SPA Algorithm returns indices $\hat{\Delta} = \{\hat{q}_1, \dots, \hat{q}_K\}$ such that for each k ,

$$\min_{\pi} \|\hat{\mathbf{U}}(\hat{q}_k, :) - \mathbf{W}(:, \pi(k))\|_2 \leq \tilde{\kappa}(\mathbf{W}) \tilde{\zeta}, \quad (95)$$

where $\tilde{\kappa}(\mathbf{W}) := 1 + 80\kappa^2(\mathbf{W})$, $\varrho := \min\left(\frac{1}{2\sqrt{K-1}}, \frac{1}{4}\right)$ and π is certain permutation for $k \in [K]$.

Therefore, we can apply Lemma 16 to the the noisy model in (92) and obtain an estimate $\hat{\mathbf{G}} = \hat{\mathbf{U}}(\hat{\Delta}, :)$ by applying SPA such that

$$\|\hat{\mathbf{G}} - \Pi \mathbf{G} \mathbf{O}\|_F \leq \sqrt{K} \tilde{\kappa}(\mathbf{G}) \tilde{\zeta}(\mathbf{G}). \quad (96)$$

In the above, we applied $\mathbf{W} = \mathbf{O}^\top \mathbf{G}^\top$, used the norm equivalence to convert from ℓ_2 -norm to the Frobenius norm, used the fact that $\sigma_{\min}(\mathbf{G} \mathbf{O}) = \sigma_{\min}(\mathbf{G})$ and $\sigma_{\max}(\mathbf{G} \mathbf{O}) = \sigma_{\max}(\mathbf{G})$ due to the orthogonality of the matrix $\mathbf{O}$ and Π is a permutation matrix such that $\Pi = \arg \min_{\tilde{\Pi}} \|\hat{\mathbf{G}} - \tilde{\Pi} \mathbf{G} \mathbf{O}\|_F$. Also, the condition in Lemma 16 can be represented as

$$\tilde{\zeta}(\mathbf{G}) = \sigma_{\max}(\mathbf{G}) \varepsilon + \zeta \leq \frac{\sigma_{\min}(\mathbf{G}) \varrho}{\tilde{\kappa}(\mathbf{G})}. \quad (97)$$

Next, we proceed to characterize the singular values of $\mathbf{G}$ to express the above bound in more intuitive way. To this end, we have the following lemma:

Lemma 17 The below relations hold:

$$\begin{aligned} \sigma_{\max}(\mathbf{G}) &= 1/\sigma_{\min}([\mathbf{M}_T, \mathbf{M}_{T+1}]) \leq \frac{1}{\alpha_{\min}}, \\ \sigma_{\max}(\mathbf{G}) &\geq \sigma_{\min}(\mathbf{G}) = 1/\sigma_{\max}([\mathbf{M}_T, \mathbf{M}_{T+1}]) \geq \frac{1}{\sqrt{2K} \alpha_{\max}}, \\ \kappa(\mathbf{G}) &\leq \sqrt{2K} \gamma. \end{aligned}$$

Proof 1 Recall that, under noiseless case, Algorithm 1 first directly estimates $\mathbf{U}_T$ and $\mathbf{U}_{T+1}$ by performing the top- K SVD to $\mathbf{C}_T = [\mathbf{P}_{T,m_T}^\top, \mathbf{P}_{T+1,m_T}^\top]^\top$ where

$$\mathbf{C}_T = [\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top \Sigma_T \mathbf{V}_{m_T}^\top. \quad (98)$$

The above relation implies that, $[\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top = [\mathbf{M}_T, \mathbf{M}_{T+1}]^\top \mathbf{G}$ holds where $[\mathbf{U}_T^\top, \mathbf{U}_{T+1}^\top]^\top$ is semi-orthogonal. Applying Fact 7 we can obtain

$$\begin{aligned} \sigma_{\max}(\mathbf{G}) &= 1/\sigma_{\min}([\mathbf{M}_T, \mathbf{M}_{T+1}]) \quad \text{and} \\ \sigma_{\min}(\mathbf{G}) &= 1/\sigma_{\max}([\mathbf{M}_T, \mathbf{M}_{T+1}]). \end{aligned}$$

By applying Lemma 3 we can finally get

$$\begin{aligned} \sigma_{\max}(\mathbf{G}) &\leq 1/\alpha_{\min}, \quad \sigma_{\min}(\mathbf{G}) \geq 1/(\sqrt{2K} \alpha_{\max}) \\ \kappa(\mathbf{G}) &\leq \sqrt{2K} \gamma. \end{aligned}$$

Applying Lemma 17, we have

$$\begin{aligned}\tilde{\kappa}(\mathbf{G}) &= 1 + 80\kappa^2(\mathbf{G}) \leq 1 + 160K\gamma^2 := \tilde{\kappa}(\mathbf{M}), \\ \tilde{\zeta}(\mathbf{G}) &= \sigma_{\max}(\mathbf{G})\varepsilon + \zeta \leq \frac{\varepsilon}{\alpha_{\min}} + \zeta.\end{aligned}$$

Hence the bound in (96) can be written as

$$\|\hat{\mathbf{G}} - \mathbf{\Pi G O}\|_F \leq \sqrt{K}\tilde{\kappa}(\mathbf{M})(\varepsilon/\alpha_{\min} + \zeta), \quad (99)$$

and the condition (97) can be rewritten as:

$$\frac{\varepsilon}{\alpha_{\min}} + \zeta \leq \frac{\varrho}{\sqrt{2K}\alpha_{\max}\tilde{\kappa}(\mathbf{M})}. \quad (100)$$

Using the given assumptions on ε and ζ , we can see that (100) gets satisfied.

Now that we have $\mathbf{G}$ estimated as given by (99), the next step is to estimate the matrix $\mathbf{M}$. Algorithm 1 estimates $\mathbf{M}$ via least squares. To characterize this step, we consider the following lemma:

Lemma 18 [64] Consider the noiseless model $\mathbf{X} = \mathbf{W}\mathbf{H}$. Suppose that $\hat{\mathbf{X}}$ and $\hat{\mathbf{W}}$ are the noisy estimates of $\mathbf{X}$ and $\mathbf{W}$, respectively. Assume that there exist constants $\phi_1, \phi_2 > 0$ such that $\|\hat{\mathbf{W}} - \mathbf{W}\|_2 \leq \phi_1\|\mathbf{W}\|_2$, $\|\hat{\mathbf{X}} - \mathbf{X}\|_2 \leq \phi_2\|\mathbf{X}\|_2$ and $\phi_1\kappa(\mathbf{W}) < 1$. Then we have

$$\frac{\|\hat{\mathbf{H}} - \mathbf{H}\|_2}{\|\mathbf{H}\|_2} \leq \frac{\kappa(\mathbf{W})(\phi_1 + \phi_2)}{1 - \phi_1\kappa(\mathbf{W})} + \kappa(\mathbf{W})\phi_1,$$

where $\hat{\mathbf{H}}$ is the least squares solution to the problem $\hat{\mathbf{X}} = \hat{\mathbf{W}}\mathbf{H}$.

By fixing $\mathbf{X} := \mathbf{O}^\top \mathbf{U}^\top$, $\mathbf{W} := \mathbf{O}^\top \mathbf{G}^\top \mathbf{\Pi}^\top$ and $\mathbf{H} := \mathbf{\Pi M}$, we can apply Lemma 18. By applying the matrix norm equivalence and considering the fact that the matrix $\mathbf{O}$ is orthogonal, the constants ϕ_1 and ϕ_2 in Lemma 18 takes the following values:

$$\phi_1 = \sqrt{K}\tilde{\kappa}(\mathbf{M})(\varepsilon/\alpha_{\min} + \zeta)/\sigma_{\max}(\mathbf{G}), \quad \phi_2 = \zeta/\sigma_{\max}(\mathbf{U}),$$

where the first equality is from (99), the second equality is by using the given assumption that $\|\hat{\mathbf{U}} - \mathbf{U}\mathbf{O}\|_F = \|\mathbf{N}\|_F \leq \zeta$, and $\tilde{\kappa}(\mathbf{M}) = 1 + 160\gamma^2$.

Then, by applying Lemma 18, the estimation bound for $\mathbf{M}$ can be expressed as

$$\frac{\|\hat{\mathbf{M}} - \mathbf{\Pi M}\|_2}{\|\mathbf{M}\|_2} \leq \frac{\kappa(\mathbf{G})(\phi_1 + \phi_2)}{1 - \phi_1\kappa(\mathbf{G})} + \kappa(\mathbf{G})\phi_1. \quad (101)$$

Using the given assumptions on ε and ζ , we can see that the denominator term $1 - \phi_1\kappa(\mathbf{G})$ is lower bounded:

$$\begin{aligned}1 - \phi_1\kappa(\mathbf{G}) &= 1 - \left(\frac{\sqrt{K}\tilde{\kappa}(\mathbf{M})\kappa(\mathbf{G})(\varepsilon/\alpha_{\min} + \zeta)}{\sigma_{\max}(\mathbf{G})} \right), \\ &= 1 - \left(\frac{\sqrt{K}\tilde{\kappa}(\mathbf{M})(\varepsilon/\alpha_{\min} + \zeta)}{\sigma_{\min}(\mathbf{G})} \right) \\ &\geq 1 - \left(\sqrt{2K}\alpha_{\max}\tilde{\kappa}(\mathbf{M})(\varepsilon/\alpha_{\min} + \zeta) \right) \\ &\geq 1/2.\end{aligned} \quad (102)$$

where we have used the assumptions on ε and ζ for the last inequality. Note the result in (102) also makes sure that the condition $\phi_1\kappa(\mathbf{G}) < 1$ in Lemma 18 is satisfied.

By applying (102) in the bound (101), we get

$$\begin{aligned} \frac{\|\widehat{\mathbf{M}} - \Pi \mathbf{M}\|_2}{\|\mathbf{M}\|_2} &\leq 2\kappa(\mathbf{G})(\phi_1 + \phi_2) + \kappa(\mathbf{G})\phi_1 = \kappa(\mathbf{G})(3\phi_1 + 2\phi_2) \\ &= \kappa(\mathbf{G}) \left(\frac{3\sqrt{K}\tilde{\kappa}(\mathbf{M})(\varepsilon/\alpha_{\min} + \zeta)}{\sigma_{\max}(\mathbf{G})} + \frac{2\zeta}{\sigma_{\max}(\mathbf{U})} \right). \end{aligned} \quad (103)$$

Consider the below to lower bound $\sigma_{\max}(\mathbf{U})$ where $\mathbf{U}^\top = \mathbf{G}^\top \mathbf{M}$:

$$\begin{aligned} \sigma_{\max}(\mathbf{U}) &\geq \sigma_{\min}(\mathbf{G})\sigma_{\max}(\mathbf{M}) \\ &= \frac{\sigma_{\max}(\mathbf{M})}{\sigma_{\max}([\mathbf{M}_T, \mathbf{M}_{T+1}])} \geq 1, \end{aligned} \quad (104)$$

where we used Lemma 17 to obtain the first equality and used the fact that $\sigma_{\max}(\mathbf{M}) \geq \sigma_{\max}([\mathbf{M}_T, \mathbf{M}_{T+1}])$ for the last inequality.

Applying Lemma 17 and (104), we have

$$\begin{aligned} \|\widehat{\mathbf{M}} - \Pi \mathbf{M}\|_2 &\leq \sigma_{\max}(\mathbf{M}) \left(3\sqrt{K}\tilde{\kappa}(\mathbf{M}) \left(\frac{\varepsilon}{\alpha_{\min}} + \zeta \right) \sqrt{2K}\alpha_{\max} + 2\sqrt{2K}\gamma\zeta \right) \\ &= \sigma_{\max}(\mathbf{M}) \left(3\sqrt{2K}\tilde{\kappa}(\mathbf{M})\gamma\varepsilon + 3\sqrt{2K}\tilde{\kappa}(\mathbf{M})\alpha_{\max}\zeta + 2\sqrt{2K}\gamma\zeta \right) \\ &\leq \sigma_{\max}(\mathbf{M}) \left(3\sqrt{2K}\tilde{\kappa}(\mathbf{M})\gamma\varepsilon + 4\sqrt{2K}\tilde{\kappa}(\mathbf{M})\gamma\zeta \right) \\ &= \sigma_{\max}(\mathbf{M})\sqrt{2K}\gamma\tilde{\kappa}(\mathbf{M}) (3\varepsilon + 4\zeta), \end{aligned}$$

where $\tilde{\kappa}(\mathbf{M}) = 1 + 160K\gamma^2$.

L Successive Projection Algorithm (SPA)

In this section, we detail the successive projection algorithm (SPA) [35, 45, 66] presented in Algorithm 3.

Assume that the separability condition holds. For the model $\mathbf{U}^\top = \mathbf{G}^\top \mathbf{M} \in \mathbb{R}^{K \times N}$, SPA estimates the anchor nodes $\Lambda = \{q_1, \dots, q_K\}$ through a series of steps as follows:

$$q_1 = \arg \max_{n \in [N]} \|\mathbf{U}(n, :)\|_2^2, \quad (105a)$$

$$q_k = \arg \max_{n \in [N]} \left\| \mathbf{P}_{\overline{\mathbf{G}}(:, 1:k-1)}^\perp \mathbf{U}(n, :)\right\|_2^2, \quad k > 1, \quad (105b)$$

where $\overline{\mathbf{G}}(:, 1:k-1) = [\mathbf{U}^\top(:, q_1), \dots, \mathbf{U}^\top(:, q_{k-1})]$ and $\mathbf{P}_{\overline{\mathbf{G}}(:, 1:k-1)}^\perp$ is a projector onto the orthogonal complement of $\text{range}(\overline{\mathbf{G}}(:, 1:k-1))$.

Eq. (105a) indicates that the first anchor node q_1 can be identified by finding the index of the row of $\mathbf{U}$ that has the maximum ℓ_2 norm. This can be understood by considering the below for any $n \in [N]$,

$$\begin{aligned} \|\mathbf{U}(n, :)\|_2 &= \left\| \sum_{k=1}^K \mathbf{G}(k, :)\mathbf{M}(k, n) \right\|_2 \leq \sum_{k=1}^K \|\mathbf{G}(k, :)\mathbf{M}(k, n)\|_2 \\ &= \sum_{k=1}^K \mathbf{M}(k, n) \|\mathbf{G}(k, :)\|_2 \leq \max_{k=1, \dots, K} \|\mathbf{G}(k, :)\|_2, \end{aligned}$$

where the first inequality is due to triangle inequality, the second equality is due to the non-negativity of $\mathbf{M}$ and the last inequality is by the sum-to-one condition on the rows of $\mathbf{M}$. In the above, all the equalities hold simultaneously if and only if $\mathbf{M}(:, n) = \mathbf{e}_k^\top$ for a certain k , i.e., $n \in \{q_1, \dots, q_K\}$.

Algorithm 3 Successive Projection Algorithm (SPA) [35]

Require: U, K

- 1: Initialize $\tilde{U} \leftarrow U$;
 - 2: **for each** $k = 1 : K$ **do**
 - 3: $q_k \leftarrow \arg \max_{n \in [N]} \|\tilde{U}(n, :)\|_2^2$; ▷ Select anchor node
 - 4: $\mathbf{u}_k \leftarrow \tilde{U}(q_k, :)$;
 - 5: $\tilde{U} \leftarrow \tilde{U} \left(I_K - \frac{\mathbf{u}_k \mathbf{u}_k^\top}{\|\mathbf{u}_k\|_2^2} \right)$;
 - 6: **end for each**
 - 7: $\mathbf{G} \leftarrow U(\{q_1, \dots, q_K\}, :)$;
 - 8: **return** $\mathbf{G}$.
-

After identifying the first anchor node q_1 , all the rows of U are projected onto the orthogonal complement of $U(q_1, :)$ such that the same index q_1 will not show up in the subsequent steps. From the projected columns, the next anchor node is then identified in a similar way as shown in Eq. (105b). Following this procedure, SPA can identify the anchor node set $\mathbf{A}$ in K steps.

The work in [35] showed that under the following noisy model

$$\hat{U}^\top = \mathbf{G}^\top \mathbf{M} + \mathbf{N}, \quad \mathbf{M} \geq \mathbf{0}, \quad \mathbf{1}^\top \mathbf{M} = \mathbf{1}^\top, \quad (106)$$

where $\hat{U}$ is the noisy estimate of U , SPA can still provably identify $\mathbf{M}$ under the separability condition if the noise $\mathbf{N}$ has bounded norm (see Theorem 16).

M Permutation Fixing Procedure for Baseline Algorithms

For synthetic and real data experiments, we employ the procedure in Algorithm 4 for the baseline algorithms to align the permutations of the estimated $\hat{\mathbf{M}}_\ell$ over different blocks. Note that this aligning procedure is not from any of the baseline papers, but is only presented for benchmarking the baselines under the diagonal EQP pattern. To explain this procedure, let us consider the diagonal EQP pattern as shown in Fig. 6 (a toy example with $L = 3$). In order to learn the membership matrix $\mathbf{M} = [\mathbf{M}_1, \mathbf{M}_2, \mathbf{M}_3]$ using the baselines, we may run the baseline algorithm first on $\mathbf{A}_{1,2}$ and subsequently on $\mathbf{A}_{2,3}$ to obtain the below (assuming no estimation error):

$$\begin{aligned} [\mathbf{M}_1, \mathbf{M}_2] &= \text{BaselineAlgorithm}(\mathbf{A}_{1,2}) \\ [\bar{\mathbf{M}}_2, \bar{\mathbf{M}}_3] &= \text{BaselineAlgorithm}(\mathbf{A}_{2,3}) \end{aligned}$$

where $\bar{\mathbf{M}}_2 = \Pi \mathbf{M}_2$ and $\bar{\mathbf{M}}_3 = \Pi \mathbf{M}_3$. Here, $\Pi \in \{0, 1\}^{K \times K}$ is a row permutation matrix which is unknown. In order to learn the complete membership matrix, i.e., $\mathbf{M} = [\mathbf{M}_1, \mathbf{M}_2, \mathbf{M}_3]$, we need to identify this row permutation matrix Π . This can be achieved by using any permutation fixing algorithms such as Hungarian algorithm [67] by inputting $\mathbf{M}_2$ and $\bar{\mathbf{M}}_2 (= \Pi \mathbf{M}_2)$. Algorithm 4 uses this idea in an iterative fashion to learn $\mathbf{M}$ using the baseline algorithms for general L case.

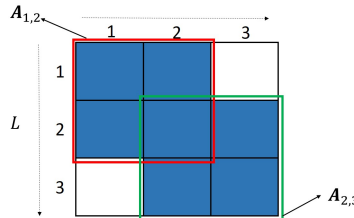


Figure 6: Adjacency matrix observed via diagonal EQP pattern ($L = 3$). The blue shaded region represents the observed blocks in the adjacency matrix.

In Algorithm 4, BaselineAlgorithm can be any graph clustering algorithm, for e.g., GeoNMF, SPACL, CD-MVSI, CD-BNMF, SC (unnorm.), SC (norm.), k -means etc.

Algorithm 4 Permutation Fixing

Require: $L, K, \{\mathbf{A}_{\ell, \ell}\}_{\ell=1}^L, \{\mathbf{A}_{\ell, \ell+1}\}_{\ell=1}^{L-1}$, BaselineAlgorithm.

```

1:  $\mathbf{\Pi} \leftarrow \mathbf{I}_K$ ;
2: for each  $\ell = 1$  to  $L - 1$  do
3:    $\mathbf{H}_\ell \leftarrow \begin{bmatrix} \mathbf{A}_{\ell, \ell} & \mathbf{A}_{\ell, \ell+1} \\ \mathbf{A}_{\ell, \ell+1}^\top & \mathbf{A}_{\ell+1, \ell+1} \end{bmatrix}$ ;
4:    $[\bar{\mathbf{M}}_\ell, \bar{\mathbf{M}}_{\ell+1}] \leftarrow \text{BaselineAlgorithm}(\mathbf{H}_\ell, K)$ ;
5:   if  $\ell > 1$  then
6:      $\mathbf{\Pi} \leftarrow \text{HugarianAlgorithm}(\bar{\mathbf{M}}_\ell, \mathbf{M}_\ell^{(prev)})$  [67];
7:   end if
8:    $[\mathbf{M}_\ell, \mathbf{M}_{\ell+1}] \leftarrow \mathbf{\Pi}^\top [\bar{\mathbf{M}}_\ell, \bar{\mathbf{M}}_{\ell+1}]$ ;
9:    $\mathbf{M}_\ell^{(prev)} \leftarrow \mathbf{M}_{\ell+1}$ ;
10: end for each
11:  $\hat{\mathbf{M}} = [\mathbf{M}_1, \dots, \mathbf{M}_L]$ ;
12: return  $\hat{\mathbf{M}}$ .

```

Table 10: $K = 5, L = 10, \boldsymbol{\nu} = [\frac{1}{5} \frac{1}{2} 1 \frac{1}{2} \frac{1}{5}]^\top$ and $\eta = 0.1$

Graph Size	Metric	BeQuec	GeoNMF	CD-MVSI	CD-BNMF
2000	MSE	0.5992	0.8277	0.7403	0.6893
	RE	0.5378	0.5695	0.8282	0.7194
	SRC	0.4452	0.2807	0.2211	0.3074
4000	MSE	0.3123	0.3704	0.7325	0.6662
	RE	0.6213	0.8930	0.8127	0.7245
	SRC	0.5908	0.5855	0.2165	0.3204
8000	MSE	0.1804	0.1859	0.7090	0.6743
	RE	0.3021	0.3597	0.8215	0.7251
	SRC	0.7465	0.7417	0.2376	0.3019
10000	MSE	0.1454	0.1435	0.7023	0.6626
	RE	0.2865	0.3580	0.8198	0.7210
	SRC	0.7715	0.7585	0.2490	0.3110

Table 11: $K = 5, L = 10, \nu = [\frac{1}{5} \frac{1}{3} \frac{1}{2} \frac{1}{5} 1]^\top$ and $\eta = 0.2$

Graph Size (N)	Metric	BeQuec	GeoNMF	CD-MVSI	CD-BNMF
2000	MSE	0.7742	0.8266	0.8793	0.8217
	RE	0.7198	0.8494	0.8237	0.8104
	SRC	0.2761	0.2737	0.1770	0.1669
4000	MSE	0.5280	0.3613	0.8810	0.8320
	RE	0.5868	0.5263	0.8336	0.8161
	SRC	0.4894	0.5496	0.1655	0.1431
8000	MSE	0.1733	0.2082	0.8875	0.8313
	RE	0.3178	0.3829	0.8439	0.8170
	SRC	0.6886	0.6618	0.1575	0.1329
10000	MSE	0.1385	0.2794	0.8861	0.8344
	RE	0.2869	0.4508	0.8444	0.8184
	SRC	0.7124	0.6216	0.1517	0.1245

Table 12: $K = 5, L = 10, \nu = [\frac{1}{5} \frac{1}{3} \frac{1}{2} \frac{1}{5} 1]^\top$ and $\eta = 0.3$

Graph Size (N)	Metric	BeQuec	GeoNMF	CD-MVSI	CD-BNMF
2000	MSE	0.7838	0.9181	0.8811	0.8494
	RE	0.7136	0.8997	0.8135	0.8182
	SRC	0.2991	0.2246	0.1634	0.1281
4000	MSE	0.5952	0.4236	0.8790	0.8403
	RE	0.5970	0.6097	0.8274	0.8159
	SRC	0.4367	0.5100	0.1537	0.1254
8000	MSE	0.2748	0.3225	0.8643	0.8467
	RE	0.4157	0.5030	0.8181	0.8211
	SRC	0.6332	0.5834	0.1607	0.1025
10000	MSE	0.2005	0.2770	0.8728	0.8456
	RE	0.3385	0.4903	0.8330	0.8214
	SRC	0.6820	0.6131	0.1548	0.1006

N Additional Synthetic Data Experiments

In this section, we present additional synthetic data experiments.

Tables 10, 12 present the results averaged over 20 random trials for different ν 's (i.e., different Dirichlet parameters) and η 's. We let entries of ν to be different, which simulates the scenario where unbalanced clusters are present. The values of η are varied to test different levels of interactions among the clusters. In these tables, we follow the same diagonal EQP pattern as in the manuscript. Note that in most of the cases, the proposed method outperforms the baselines. One can also see that all methods' performance deteriorates when η increases from 0.1 to 0.3. Nonetheless, the proposed method keeps its margin against the baselines under all η 's.