

SIMPLE: Statistical Inference on Membership Profiles in Large Networks *

Jianqing Fan¹, Yingying Fan², Xiao Han³ and Jinchi Lv²

Princeton University¹, University of Southern California²

and University of Science and Technology of China³

September 1, 2021

Abstract

Network data is prevalent in many contemporary big data applications in which a common interest is to unveil important latent links between different pairs of nodes. Yet a simple fundamental question of how to precisely quantify the statistical uncertainty associated with the identification of latent links still remains largely unexplored. In this paper, we propose the method of statistical inference on membership profiles in large networks (SIMPLE) in the setting of degree-corrected mixed membership model, where the null hypothesis assumes that the pair of nodes share the same profile of community memberships. In the simpler case of no degree heterogeneity, the model reduces to the mixed membership model for which an alternative more robust test is also proposed. Both tests are of the Hotelling-type statistics based on the rows of empirical eigenvectors or their ratios, whose asymptotic covariance matrices are very challenging to derive and estimate. Nevertheless, their analytical expressions are unveiled and the unknown covariance matrices are consistently estimated. Under some mild regularity conditions, we establish the exact limiting distributions of the two forms of SIMPLE test statistics under the null hypothesis and contiguous alternative hypothesis. They are the chi-square

*Jianqing Fan is Frederick L. Moore '18 Professor of Finance, Department of Operations Research and Financial Engineering, Princeton University, Princeton, NJ 08544, USA (E-mail: jqfan@princeton.edu). Yingying Fan is Centennial Chair in Business Administration and Professor, Data Sciences and Operations Department, Marshall School of Business, University of Southern California, Los Angeles, CA 90089 (E-mail: fanyingy@marshall.usc.edu). Xiao Han is Professor, International Institute of Finance, Department of Statistics and Finance, University of Science and Technology of China, Hefei, China 230026 (E-mail: xhan011@ustc.edu.cn). Jinchi Lv is Kenneth King Stonier Chair in Business Administration and Professor, Data Sciences and Operations Department, Marshall School of Business, University of Southern California, Los Angeles, CA 90089 (E-mail: jinchilv@marshall.usc.edu). This work was supported by NIH grant R01-GM072611-16, NSF grants DMS-1712591, DMS-1953356, DMS-2052926, and DMS-2052964, NSF CAREER Award DMS-1150318, a grant from the Simons Foundation, Adobe Data Science Research Award, and a grant from NSF of China (No.12001518). Corresponding authors: Jianqing Fan and Xiao Han.

distributions and the noncentral chi-square distributions, respectively, with degrees of freedom depending on whether the degrees are corrected or not. We also address the important issue of estimating the unknown number of communities and establish the asymptotic properties of the associated test statistics. The advantages and practical utility of our new procedures in terms of both size and power are demonstrated through several simulation examples and real network applications.

Running title: SIMPLE

Key words: Network p-values; Statistical inference; Large networks; Clustering; Big data; Random matrix theory; Eigenvectors; Eigenvalues

1 Introduction

Large-scale network data that describes the pairwise relational information among objects is commonly encountered in many applications such as the studies of citation networks, protein-protein interaction networks, health networks, financial networks, trade networks, and social networks. The popularity of such applications has motivated a spectrum of research with network data. Popularly used methods include algorithmic ones and model-based ones, where the former uses algorithms to optimize some carefully designed criteria (e.g., [Newman \(2013a,b\)](#); [Zhang and Moore \(2014\)](#)), and the latter relies on specific structures of some probabilistic models (see, e.g., [Goldenberg et al. \(2010\)](#) for a review). This paper belongs to the latter group. In the literature, a number of probabilistic models have been proposed for modeling network data. As arguably the simplest model with planted community identity, the stochastic block model (SBM) ([Holland et al., 1983](#); [Wang and Wong, 1987](#); [Abbe, 2017](#)) has received a tremendous amount of attention in the last decade. To overcome the limitation and increase the flexibility in the basic stochastic block model, various variants have been proposed. To name a few, the degree-corrected SBM ([Karrer and Newman, 2011](#)) introduces a degree parameter for each node to make the expected degrees match the observed ones. The overlapping SBM, such as the mixed membership model ([Airoldi et al., 2008](#)), allows the communities to overlap by assigning each node a profile of community memberships. See also [Newman and Peixoto \(2015\)](#) for a review of network models.

An important problem in network analysis is to unveil the true latent links between different pairs of nodes, where nodes can be broadly defined such as individuals, economic entities, documents, or medical disorders in social, economic, text, or health networks. There is a growing literature on network analysis with various methods available for clustering the nodes into different communities within which nodes are more densely connected, based on the observed adjacency matrices or the similarity matrices constructed using the node information. These methods focus mainly on the clustering aspect of the problem, outputting subgroups with predicted membership identities. Yet the statistical inference aspect such as quantifying the statistical uncertainty associated with the identification of latent links

has been largely overlooked. This paper aims at filling this crucial gap by proposing new statistical tests for testing whether any given pair of nodes share the same membership profiles, and providing the associated p -values.

Knowing the statistical significance of membership profiles can bring more confidence to practitioners in decision making. Taking the stock market for example, investors often want to form diversified portfolios by including stocks with little or no correlation in their returns. The correlation matrix of stock returns can then be used to construct an affinity matrix, and stocks with relatively highly correlated returns can be regarded as in the same community. Obtaining the pairwise p -values of stocks can help investors form diversified portfolios with statistical confidence. For instance, if one is interested in the Apple stock, then the pairwise p -values of Apple and all other candidate stocks can be calculated, and stocks with the smallest p -values can be included to form portfolios. Another important application is in legislation. For example, illegal logging greatly conflicts with indigenous and local populations, contributing to violence, human rights abuses, and corruption. The DNA sequencing technology has been used to identify the region of logs. In such application, an affinity matrix can be calculated according to the similarity of DNA sequences. Then applying our method, p -values can be calculated and used in court as statistical evidence in convicting illegal logging.

To make the problem concrete, we consider the family of degree-corrected mixed membership models, which includes the mixed membership model and the stochastic block model as special cases. In the degree-corrected mixed membership model, node i is assumed to have a membership profile characterized by a community membership probability vector $\pi_i \in \mathbb{R}^K$, where K is the number of communities and the k th entry of π_i specifies the mixture proportion of node i in community k (Airoldi et al., 2008; Zhao et al., 2012). For example, a book can be 30% liberal and 70% conservative. In addition, each node is allowed to have its own degree. For any given pair of nodes i and j , we investigate whether they have the same membership profile or not by testing the hypothesis $H_0 : \pi_i = \pi_j$ vs. $H_a : \pi_i \neq \pi_j$. Two forms of statistical inference on membership profiles in large networks (SIMPLE) test are proposed. Under the mixed membership model where all nodes have the same degree, we construct the first form of SIMPLE test by resorting to the i th and j th rows of the spiked eigenvector matrix of the observed adjacency matrix. We establish the asymptotic null and alternative distributions of the test statistic, where under the null hypothesis the asymptotic distribution is chi-square with K degrees of freedom and under the alternative hypothesis, the asymptotic distribution is noncentral chi square with a location parameter determined by how distinct the membership profiles of nodes i and j are.

In the more general degree-corrected mixed membership model, where nodes are allowed to have heterogeneous degrees, we build the second form of SIMPLE test based on the ratio statistic proposed in Jin (2015). We show that the asymptotic null distribution is chi-

square with $K - 1$ degrees of freedom, and under the alternative hypothesis and some mild regularity conditions, the test statistic diverges to infinity with asymptotic probability one. We prove that these asymptotic properties continue to hold even with estimated population parameters (including the number of communities K) provided that these parameters can be estimated reasonably well. We then suggest specific estimators of these unknown parameters and show that they achieve the desired estimation precision. These new theoretical results enable us to construct rejection regions that are pivotal to the unknown parameters for each of these two forms of the SIMPLE test, and to calculate p -values explicitly. Our method is more applicable than most existing ones in the community detection literature where K is required to be known. Although the second form of SIMPLE test can be applied to both cases with and without degree heterogeneity, we would like to point out that the first test is empirically more stable since it does not involve any ratio calculations. To the best of our knowledge, this paper is the first in the literature to provide quantified uncertainty levels in community membership estimation and inference.

Our test is most useful when one cares about local information of the network. For instance, if the interest is whether two (or several) nodes belong to the same community with quantified significance level, then SIMPLE can be used. Indeed, our statistics do not rely on any pre-determined membership information. Compared to community detection methods, our work has at least three advantages: 1) we do not need to assign memberships to nodes that are not of interests; 2) our method can provide the level of significance, which can be very important in scientific discoveries; and 3) if partial membership information is known in a network, then the nodes with missing membership information can be recovered with statistical confidence by applying our tests.

Both forms of SIMPLE test are constructed using the spectral information of the observed adjacency matrix. In this sense, our work is related to the class of spectral clustering methods, which is one of the most scalable tools for community detection and has been popularly used in the literature. See, e.g., [von Luxburg \(2007\)](#) for a tutorial of spectral clustering methods. See also [Rohe et al. \(2011\)](#); [Lei and Rinaldo \(2015\)](#); [Jin \(2015\)](#) among many others for the specifics on the implementation of spectral methods for community detection. In addition, the optimality for the case of two communities has been established by [Abbe et al. \(2017\)](#). Our work is related to but substantially different from the link prediction problem ([Liben-Nowell and Kleinberg, 2007](#); [Wu et al., 2018](#)), which can be thought of as predicting pairs of nodes as linked or non-linked. The major difference is that in link prediction, only part of the adjacency matrix is observed and one tries to predict the latent links among the nodes which are unobserved. Moreover, link prediction methods usually do not provide statistical confidence levels.

Our work falls into the category of hypothesis testing with network data. In the literature, hypothesis testing has been used for different purposes. For example, [Arias-Castro and](#)

Verzelen (2014) and Verzelen and Arias-Castro (2015) formalized the problem of community detection in a given random graph as a hypothesis testing problem in dense and sparse random networks, respectively. Under the stochastic block model assumption, Bickel and Sarkar (2016) proposed a recursive bipartitioning algorithm to automatically estimate the number of communities using hypothesis test constructed from the largest principal eigenvalue of the suitably centered and scaled adjacency matrix. The null hypothesis of their test is that the network has only $K = 1$ community. Lei (2016) generalized their idea and proposed a test allowing for $K \geq 1$ communities in the stochastic block model under the null hypothesis. The number of communities can then be estimated by sequential testing. Wang and Bickel (2017) proposed a likelihood ratio test for selecting the correct K under the setting of SBM.

The rest of the paper is organized as follows. Section 2 introduces the model setting and technical preparation. We present the SIMPLE method and its asymptotic theory as well as the implementation details of SIMPLE in Section 3. Sections 4 and 5 provide several simulation and real data examples illustrating the finite-sample performance and utility of our newly suggested method. We discuss some implications and extensions of our work in Section 6. All the proofs and technical details are provided in the Supplementary Material.

2 Statistical inference in large networks

2.1 Model setting

Consider an undirected graph $\mathcal{N} = (V, E)$ with n nodes, where $V = \{1, \dots, n\}$ is the set of nodes and E is the set of links. Throughout the paper, we use the notation $[n] = \{1, \dots, n\}$. Let $\mathbf{X} = (x_{ij}) \in \mathbb{R}^{n \times n}$ be the symmetric adjacency matrix representing the connectivity structure of graph $\mathcal{N}$, where $x_{ij} = 1$ if there is a link connecting nodes i and j , and $x_{ij} = 0$ otherwise. We consider the general case when graph $\mathcal{N}$ may or may not admit self loops, where in the latter scenario $x_{ii} = 0$ for all $i \in [n]$. Under a probabilistic model, we will assume that x_{ij} is an independent realization from a Bernoulli random variable for all upper triangular entries of random matrix $\mathbf{X}$.

To model the connectivity pattern of graph $\mathcal{N}$, consider a symmetric binary random matrix $\mathbf{X}^*$ with the following latent structure

$$\mathbf{X}^* = \mathbf{H} + \mathbf{W}^*, \tag{1}$$

where $\mathbf{H} = (h_{ij}) \in \mathbb{R}^{n \times n}$ is the deterministic mean matrix (or probability matrix) of low rank $K \geq 1$ (see (5) later for a specification) and $\mathbf{W}^* = (w_{ij}^*) \in \mathbb{R}^{n \times n}$ is a symmetric random matrix with mean zero and independent entries on and above the diagonal. Assume that the observed adjacency matrix $\mathbf{X}$ is either $\mathbf{X}^*$ or $\mathbf{X}^* - \text{diag}(\mathbf{X}^*)$, corresponding to the cases

with or without self loops, respectively. In either case, we have the following decomposition

$$\mathbf{X} = \mathbf{H} + \mathbf{W}, \quad (2)$$

where $\mathbf{W} = \mathbf{W}^*$ in the presence of self loops and $\mathbf{W} = \mathbf{W}^* - \text{diag}(\mathbf{X}^*)$ in the absence of self loops. We can see that in either case, $\mathbf{W}$ in (2) is symmetric with independent entries on and above the diagonal. Our study will cover both cases. Hereafter to simplify the presentation, we will slightly abuse the notation by referring to $\mathbf{H}$ as the mean matrix and $\mathbf{W}$ as the noise matrix.

Assume that there is an underlying latent community structure that the network $\mathcal{N}$ can be decomposed into K latent disjoint communities

$$\mathcal{C}_1, \dots, \mathcal{C}_K,$$

where each node i is associated with the community membership probability vector $\boldsymbol{\pi}_i = (\boldsymbol{\pi}_i(1), \dots, \boldsymbol{\pi}_i(K))^T \in \mathbb{R}^K$ such that

$$P(\text{node } i \text{ belongs to community } \mathcal{C}_k) = \boldsymbol{\pi}_i(k), \quad k = 1, \dots, K. \quad (3)$$

Throughout the paper, we assume that the number of communities K is *unknown* but bounded away from infinity.

For any given pair of nodes $i, j \in V$ with $i \neq j$, our goal is to infer whether they share the same community identity or not with quantified uncertainty level from the observed adjacency matrix $\mathbf{X}$ in the general model (2). In other words, for each pair of nodes $i, j \in V$ with $i \neq j$, we are interested in testing the hypothesis

$$H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j \quad \text{versus} \quad H_a : \boldsymbol{\pi}_i \neq \boldsymbol{\pi}_j. \quad (4)$$

Throughout the paper, we consider the preselected pair (i, j) and thus nodes i and j are fixed.

To make the problem more explicit, we consider the degree-corrected mixed membership (DCMM) model. Using the same formulation as in Jin et al. (2017), the probability of a link between nodes i and j with $i \neq j$ under the DCMM model can be written as

$$P(x_{ij} = 1) = \theta_i \theta_j \sum_{k=1}^K \sum_{l=1}^K \boldsymbol{\pi}_i(k) \boldsymbol{\pi}_j(l) p_{kl}. \quad (5)$$

Here, $\theta_i > 0$, $i \in [n]$, measures the degree heterogeneity, and p_{kl} can be interpreted as the probability of a typical member ($\theta_i = 1$, say) in community $\mathcal{C}_k$ connects with a typical member ($\theta_j = 1$, say) in community $\mathcal{C}_l$, as in the stochastic block model. Writing (5) in the matrix form, we have

$$\mathbf{H} = \mathbf{\Theta} \mathbf{\Pi} \mathbf{P} \mathbf{\Pi}^T \mathbf{\Theta}, \quad (6)$$

where $\mathbf{\Theta} = \text{diag}(\theta_1, \dots, \theta_n)$ stands for the degree heterogeneity matrix, $\mathbf{\Pi} = (\boldsymbol{\pi}_1, \dots, \boldsymbol{\pi}_n)^T \in \mathbb{R}^{n \times K}$ is the matrix of community membership probability vectors, and $\mathbf{P} = (p_{kl}) \in \mathbb{R}^{K \times K}$ is a nonsingular matrix with $p_{kl} \in [0, 1]$, $1 \leq k, l \leq K$.

The family of DCMM models in (6) contains several popularly used network models for community detection as special cases. For example, when $\mathbf{\Theta} = \sqrt{\theta} \mathbf{I}_n$ and $\boldsymbol{\pi}_i \in \{\mathbf{e}_1, \dots, \mathbf{e}_K\}$ with $\mathbf{e}_k$ a unit vector whose k th component is one and all other components are zero, the model reduces to the stochastic block model with non-overlapping communities. When $\mathbf{\Theta} = \sqrt{\theta} \mathbf{I}_n$ and $\boldsymbol{\pi}_i$'s are general community membership probability vectors, the model becomes the mixed membership model. Each of these models has been studied extensively in the literature. Yet almost all these existing works have focused on the community detection perspective, which is a statistical estimation problem. In this paper, however we will concentrate on the statistical inference problem (4).

2.2 Technical preparation

When $\mathbf{\Theta} = \sqrt{\theta} \mathbf{I}_n$, we have $\mathbf{H} = \theta \mathbf{\Pi} \mathbf{P} \mathbf{\Pi}^T$. Thus the column space spanned by $\mathbf{\Pi}$ is the same as the eigenspace spanned by the top K eigenvectors of matrix $\mathbf{H}$. In other words, the membership profiles of the network are encoded in the eigen-structure of the mean matrix $\mathbf{H}$. Denote by $\mathbf{H} = \mathbf{V} \mathbf{D} \mathbf{V}^T$ the eigen-decomposition of the mean matrix, where $\mathbf{D} = \text{diag}(d_1, \dots, d_K)$ with $|d_1| \geq |d_2| \geq \dots \geq |d_K| > 0$ is the matrix of all K nonzero eigenvalues and $\mathbf{V} = (\mathbf{v}_1, \dots, \mathbf{v}_K) \in \mathbb{R}^{n \times K}$ is the corresponding orthonormal matrix of eigenvectors. In practice, one replaces the matrices $\mathbf{D}$ and $\mathbf{V}$ by those of the observed adjacency matrix $\mathbf{X}$. Denote by $\hat{d}_1, \dots, \hat{d}_n$ the eigenvalues of matrix $\mathbf{X}$ and $\hat{\mathbf{v}}_1, \dots, \hat{\mathbf{v}}_n$ the corresponding eigenvectors. Without loss of generality, assume that $|\hat{d}_1| \geq |\hat{d}_2| \geq \dots \geq |\hat{d}_n|$ and let $\hat{\mathbf{V}} = (\hat{\mathbf{v}}_1, \dots, \hat{\mathbf{v}}_K) \in \mathbb{R}^{n \times K}$. Denote by $\mathbf{W} = (w_{ij})$ and define $\alpha_n = \{\max_{1 \leq j \leq n} \sum_{i=1}^n \text{var}(w_{ij})\}^{1/2}$, which is simply the maximum standard deviation of the column sums (node degrees).

The asymptotic mean of the empirical eigenvalue $\hat{d}_k$ for $k \in [K]$ has been derived in Fan et al. (2020), which is a population quantity t_k and will be used frequently in our paper. Its definition is somewhat complicated which we now describe as follows. Let a_k and b_k be defined as

$$a_k = \begin{cases} \frac{d_k}{1+c_0/2} & \text{if } d_k > 0 \\ (1+c_0/2)d_k & \text{if } d_k < 0 \end{cases}, \quad b_k = \begin{cases} (1+c_0/2)d_k & \text{if } d_k > 0 \\ \frac{d_k}{1+c_0/2} & \text{if } d_k < 0 \end{cases},$$

where the eigen-ratio gap constant $c_0 > 0$ is given in Condition 1 in Section 3.1. For any deterministic real-valued matrices $\mathbf{M}_1$ and $\mathbf{M}_2$ of appropriate dimensions and complex

number $z \neq 0$, define

$$\mathcal{R}(\mathbf{M}_1, \mathbf{M}_2, z) = -\frac{1}{z} \mathbf{M}_1^T \mathbf{M}_2 - \sum_{l=2}^L \frac{1}{z^{l+1}} \mathbf{M}_1^T \mathbb{E} \mathbf{W}^l \mathbf{M}_2 \quad (7)$$

with L the smallest positive integer such that uniformly over $k \in [K]$,

$$\left(\frac{\alpha_n}{|z|} \right)^L \leq \min \left\{ \frac{1}{n^4}, \frac{1}{|z|^4} \right\}, \quad z \in [a_k, b_k], \quad (8)$$

where $|z|$ denotes the modulus of complex number z . We can see that as long as $\frac{|d_K|}{\alpha_n} \geq n^\epsilon$ with some positive constant ϵ , which is guaranteed by Condition 1 and Condition 2 (or 4) in Section 3.1 (or Section 3.2), the existence of the desired positive integer L can be ensured.

We are now ready to define the asymptotic mean t_k of the sample eigenvalue $\hat{d}_k$. For each $k \in [K]$, define t_k as the solution to equation

$$1 + d_k \left\{ \mathcal{R}(\mathbf{v}_k, \mathbf{v}_k, z) - \mathcal{R}(\mathbf{v}_k, \mathbf{V}_{-k}, z) [\mathbf{D}_{-k}^{-1} + \mathcal{R}(\mathbf{V}_{-k}, \mathbf{V}_{-k}, z)]^{-1} \mathcal{R}(\mathbf{V}_{-k}, \mathbf{v}_k, z) \right\} = 0 \quad (9)$$

when restricted to the interval $z \in [a_k, b_k]$, where $\mathbf{V}_{-k}$ is the submatrix of $\mathbf{V}$ formed by removing the k th column and $\mathbf{D}_{-k}$ is formed by removing the k th diagonal entry of $\mathbf{D}$. Then as shown in Fan et al. (2020), for each $k \in [K]$, t_k is the asymptotic mean of the sample eigenvalue $\hat{d}_k$ and $t_k/d_k \rightarrow 1$ as $n \rightarrow \infty$. See also Lemma 12 in Section C.6 of Supplementary Material, where the existence, uniqueness, and asymptotic property of t_k 's are stated.

To facilitate the technical presentation, we further introduce some notation that will be used throughout the paper. We use $a \ll b$ to represent $a/b \rightarrow 0$. For a matrix $\mathbf{A} = (\mathbf{A}_{ij})$, denote by $\lambda_j(\mathbf{A})$ the j th largest eigenvalue, and $\|\mathbf{A}\|_F = \sqrt{\text{tr}(\mathbf{A}\mathbf{A}^T)}$, $\|\mathbf{A}\|_2$, and $\|\mathbf{A}\|_\infty = \max_{i,j} |\mathbf{A}_{ij}|$ the Frobenius norm, the spectral norm, and the entrywise maximum norm, respectively. In addition, we use $\mathbf{A}(k)$ to denote the k th row of a matrix $\mathbf{A}$, and $\mathbf{a}(k)$ to denote the k th component of a vector $\mathbf{a}$. For a unit vector $\mathbf{x} = (x_1, \dots, x_n)^T$, let $d_{\mathbf{x}} = \max_{1 \leq i \leq n} |x_i|$. Also define $\theta_{\max} = \max_{1 \leq i \leq n} \theta_i$ and $\theta_{\min} = \min_{1 \leq i \leq n} \theta_i$ as the maximum and minimum node degrees, respectively. For each $1 \leq k \leq K$, denote by $\mathcal{N}_k = \{i : 1 \leq i \leq n, \pi_i(k) = 1\}$ the set of pure nodes in community k , where each pure node belongs to only a single community. Some additional definitions and notation are given at the beginning of Section A.

3 SIMPLE and its asymptotic theory

3.1 SIMPLE for mixed membership models

We first consider the hypothesis testing problem (4) in the mixed membership model without degree heterogeneity whose mean matrix takes the form (6) with $\Theta = \sqrt{\theta}\mathbf{I}_n$, that is,

$$\mathbb{E}\mathbf{X} = \mathbf{H} = \theta\Pi\Pi\Pi^T. \quad (10)$$

Here θ is allowed to converge to zero as $n \rightarrow \infty$. This model is a simple version of the mixed membership stochastic block (MMSB) model considered in [Airoldi et al. \(2008\)](#). As mentioned before, this model includes the stochastic block model with non-overlapping communities as a special case.

Under model (10), if $\pi_i = \pi_j$ then nodes i and j are exchangeable and it holds that $\mathbf{V}(i) = \mathbf{V}(j)$ by a simple permutation argument (see the beginning of the proof of Theorem 1 in Section A.1). Motivated by this observation, we consider the following test statistic for assessing the membership information of the i th and j th nodes

$$T_{ij} = \left[\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j) \right]^T \boldsymbol{\Sigma}_1^{-1} \left[\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j) \right], \quad (11)$$

where $\boldsymbol{\Sigma}_1$ is the asymptotic variance of $\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j)$ that is challenging to derive and estimate. Nevertheless, we will show that $\boldsymbol{\Sigma}_1 = \text{cov}[(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{WVD}^{-1}]$ whose expression is given in (28) later, and provide an estimator with required accuracy.

We need the following regularity conditions in establishing the asymptotic null and alternative distributions of test statistic T_{ij} .

Condition 1. *There exists some positive constant c_0 such that*

$$\min\left\{\frac{|d_i|}{|d_j|} : 1 \leq i < j \leq K, d_i \neq -d_j\right\} \geq 1 + c_0.$$

In addition, $\alpha_n \rightarrow \infty$ as $n \rightarrow \infty$.

Condition 2. *There exist some constants $0 < c_0 < 1$, $0 \leq c_2 < 1/2$, $0 < c_1 < 1 - 2c_2$ such that $\lambda_K(\Pi^T \Pi) \geq c_0 n$, $\lambda_K(\mathbf{P}) \geq n^{-c_2}$, and $\theta \geq n^{-c_1}$.*

Condition 3. *As $n \rightarrow \infty$, all the eigenvalues of $\theta^{-1}\mathbf{D}\boldsymbol{\Sigma}_1\mathbf{D}$ are bounded away from 0 and ∞ .*

The constant c_0 in Condition 1 can be replaced with some $o(1)$ term that vanishes as n grows at the cost of significantly more tedious calculations in our technical analysis. This condition is imposed to exclude the complicated case of multiplicity, which can lead to the singularity of $\boldsymbol{\Sigma}_1$, making our test ill-defined. A potential remedy is to use the Moore–Penrose generalized inverse of matrix $\boldsymbol{\Sigma}_1$ in defining T_{ij} , which we will leave to the future study due

to the extra technical challenge. We acknowledge that in some special models, results on community detection have been established allowing multiplicity (e.g., [Gao et al. \(2018\)](#)). Condition 2 is a standard regularity assumption imposed for the case of mixed membership models. In particular, θ measures the degree density and is allowed to converge to zero at the polynomial rate n^{-c_1} with constant c_1 arbitrarily close to one. Condition 3 is a technical condition for establishing the asymptotic properties of T_{ij} . We provide sufficient conditions for ensuring Condition 3 in Section D of Supplementary file. As shown in the proof of Theorem 1, under Conditions 1 and 2, we have $\text{var}[(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k] \sim \theta$ for all $k = 1, \dots, K$, which explains the normalization factor θ^{-1} in Condition 3. Our conditions accommodate the case where the magnitudes of spiked eigenvalues $|d_1|, \dots, |d_K|$ are of different orders.

Example 1. Consider SBM with $K = 2$ communities of equal sizes $n_1 = n_2 = n/2$ and $n^{-c_1} \leq \theta < 1$. Further assume that $\mathbf{P}$ has diagonal entries equal to a and off-diagonal entries equal to b , with a and b some positive constants satisfying $a > b$. Then we have $d_1 = n(a + b)\theta$ and $d_2 = n(a - b)\theta$. Some direct calculations show that Conditions 1–3 all hold.

The following theorem summarizes the asymptotic distribution of test statistic T_{ij} under the null and alternative hypotheses.

Theorem 1. Assume that Conditions 1–2 hold under the mixed membership model (10).

i) Under the null hypothesis $H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$, if in addition Condition 3 holds, then we have

$$T_{ij} \xrightarrow{\mathcal{D}} \chi_K^2 \quad (12)$$

as $n \rightarrow \infty$, where χ_K^2 is the chi-square distribution with K degrees of freedom.

ii) Under the contiguous alternative hypothesis $H_a : \boldsymbol{\pi}_i \neq \boldsymbol{\pi}_j$ but $n^{1/2-c_2}\sqrt{\theta}\|\boldsymbol{\pi}_i - \boldsymbol{\pi}_j\| \rightarrow \infty$, then for arbitrarily large constant $C > 0$, we have

$$P(T_{ij} > C) \rightarrow 1 \quad (13)$$

as $n \rightarrow \infty$. Moreover, if Condition 3 holds, $c_2 = 0$, $\|\boldsymbol{\pi}_i - \boldsymbol{\pi}_j\| \sim \frac{1}{\sqrt{n\theta}}$, and $[\mathbf{V}(i) - \mathbf{V}(j)]^T \boldsymbol{\Sigma}_1^{-1} [\mathbf{V}(i) - \mathbf{V}(j)] \rightarrow \mu$ with μ some constant, then it holds that

$$T_{ij} \xrightarrow{\mathcal{D}} \chi_K^2(\mu) \quad (14)$$

as $n \rightarrow \infty$, where $\chi_K^2(\mu)$ is a noncentral chi-square distribution with mean μ and K degrees of freedom.

Remark 1. Under the joint null hypotheses $H_{0,ij} : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$ for all $1 \leq i \neq j \leq n$, we have

in fact proved a uniform version of the result in (12):

$$\lim_{n \rightarrow \infty} \sup_{1 \leq i \neq j \leq n} |P(T_{ij} \leq x) - P(X \leq x)| = 0 \text{ for all } x \in \mathbb{R}, \quad (15)$$

where $X \sim \chi_K^2$. See Section E of Supplementary Material for more details.

In the special case of stochastic block model with non-overlapping communities, we can see that $\|\pi_i - \pi_j\| = 0$ under the null hypothesis H_0 , and $\|\pi_i - \pi_j\| = \sqrt{2}$ under the alternative hypothesis H_a . Thus under the null hypothesis H_0 and Conditions 1–3, the test statistic T_{ij} has asymptotic distribution (12). Under the alternative hypotheses H_a and Conditions 1–2, we have $n^{1/2-c_2} \sqrt{\theta} \|\pi_i - \pi_j\| \rightarrow \infty$ and thus the limiting result (13) holds.

The test statistic T_{ij} is, however, not directly applicable because of the unknown population parameters K and Σ_1 . We next show that for consistent estimators satisfying the following conditions

$$P(\hat{K} = K) = 1 - o(1), \quad (16)$$

$$\theta^{-1} \|\mathbf{D}(\hat{\mathbf{S}}_1 - \Sigma_1)\mathbf{D}\|_2 = o_p(1), \quad (17)$$

the asymptotic results in Theorem 1 continue to hold.

Theorem 2. Assume that estimators $\hat{K}$ and $\hat{\mathbf{S}}_1$ satisfy (16) and (17), respectively. Let $\hat{T}_{ij}$ be the test statistic constructed by replacing K and Σ_1 in (11) with $\hat{K}$ and $\hat{\mathbf{S}}_1$, respectively. Then Theorem 1 holds with T_{ij} replaced by $\hat{T}_{ij}$ under the same conditions.

Theorem 2 suggests that at significance level α , to test the null hypothesis H_0 in (4), we can construct the following rejection region

$$\{\hat{T}_{ij} > \chi_{\hat{K}, 1-\alpha}^2\}, \quad (18)$$

where $\chi_{\hat{K}, 1-\alpha}^2$ is the $100(1-\alpha)$ th percentile of the chi-square distribution with $\hat{K}$ degrees of freedom. The following corollary justifies the asymptotic size and power of our test.

Corollary 1. Assume that $\hat{K}$ and $\hat{\mathbf{S}}_1$ satisfy (16) and (17), respectively. Under the same conditions for ensuring (12), event (18) holds with asymptotic probability α . Under the same conditions for ensuring (13), event (18) holds with asymptotic probability one.

3.2 SIMPLE for degree-corrected mixed membership models

In this section, we further consider the hypothesis testing problem (4) in the more general DCMM model (6). Degree heterogeneity in network models has been explored in the statistics literature. To name a few, Jin et al. (2017) considered the estimation of node membership assuming the average degree of the nodes to be much larger than $\log n$. Jin and Ke (2017)

established a sharp lower bound for the estimated node membership allowing the average node degree to diverge with the order $\log^2 n$ or faster. Zhang et al. (2020) proposed a spectral-based detection algorithm to recover the node membership assuming that $\theta_{\max}/\theta_{\min}$ is bounded by some positive constant. Our assumption on the degree heterogeneity is similar to that in Zhang et al. (2020) and will be presented in Condition 4 below.

The test statistic T_{ij} defined in Section 3.1 is no longer applicable due to the degree heterogeneity. A simple algebra shows that degree heterogeneity can be eliminated by the ratios of eigenvectors (columnwise division). Thus, following Jin (2015), to correct the degree heterogeneity we define the following componentwise ratio

$$Y(i, k) = \frac{\hat{\mathbf{v}}_k(i)}{\hat{\mathbf{v}}_1(i)}, \quad 1 \leq i \leq n, 2 \leq k \leq K, \quad (19)$$

where $0/0$ is defined as 1 by convention. Note that the division here is to get rid of the degree heterogeneity and the equality

$$\frac{\mathbf{v}_k(i)}{\mathbf{v}_1(i)} = \frac{\mathbf{v}_k(j)}{\mathbf{v}_1(j)}, \quad 2 \leq k \leq K \quad (20)$$

holds under the null hypothesis, which is due to the exchangeability of nodes i and j under the mixed membership model; see (18) at the beginning of the proof of Theorem 3 in Section A.4. Denote by $\mathbf{Y}_i = (Y(i, 2), \dots, Y(i, K))^T$. Our new test statistic will be built upon $\mathbf{Y}_i$.

To test the null hypothesis $H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$, using (19) and (20), we propose to use the following test statistic

$$G_{ij} = (\mathbf{Y}_i - \mathbf{Y}_j)^T \boldsymbol{\Sigma}_2^{-1} (\mathbf{Y}_i - \mathbf{Y}_j) \quad (21)$$

for assessing the null hypothesis H_0 in (4), where $\boldsymbol{\Sigma}_2$ is the asymptotic variance of $\mathbf{Y}_i - \mathbf{Y}_j$. This is even much harder to derive and estimate. Nevertheless, we will show $\boldsymbol{\Sigma}_2 = \text{cov}(\mathbf{f})$ with $\mathbf{f} = (f_2, \dots, f_K)^T$ and

$$f_k = \frac{\mathbf{e}_i^T \mathbf{W} \mathbf{v}_k}{t_k \mathbf{v}_1(i)} - \frac{\mathbf{e}_j^T \mathbf{W} \mathbf{v}_k}{t_k \mathbf{v}_1(j)} - \frac{\mathbf{v}_k(i) \mathbf{e}_i^T \mathbf{W} \mathbf{v}_1}{t_1 \mathbf{v}_1^2(i)} + \frac{\mathbf{v}_k(j) \mathbf{e}_j^T \mathbf{W} \mathbf{v}_1}{t_1 \mathbf{v}_1^2(j)}. \quad (22)$$

The entries of $\boldsymbol{\Sigma}_2$ are given by (29) later that also involves the asymptotic mean of $\hat{d}_k$.

The following conditions are needed for investigating the asymptotic properties of test statistic G_{ij} .

Condition 4. *There exist some constants $c_2 \in [0, 1/2)$, $c_3 \in (0, 1 - 2c_2)$, $c_5 \in (0, 1)$ and $c_4 > 0$ such that $\lambda_K(\mathbf{P}) \geq n^{-c_2}$, $\min_{1 \leq k \leq K} |\mathcal{N}_k| \geq c_5 n$, $\theta_{\max} \leq c_4 \theta_{\min}$, and $\theta_{\min}^2 \geq n^{-c_3}$.*

Condition 5. *Matrix $\mathbf{P} = (p_{kl})$ is positive definite, irreducible, and has unit diagonal entries. Moreover $n \min_{1 \leq k \leq K, t=i,j} \text{var}(\mathbf{e}_t^T \mathbf{W} \mathbf{v}_k) \sim n \theta_{\max}^2 \rightarrow \infty$.*

Condition 6. *It holds that all the eigenvalues of $(n \theta_{\max}^2)^{-1} \mathbf{D} \text{cov}(\mathbf{f}) \mathbf{D}$ are bounded away from 0 and ∞ .*

Condition 7. Let $\boldsymbol{\eta}_1$ be the first right singular vector of $\mathbf{P}\boldsymbol{\Pi}^\top \boldsymbol{\Theta}^2 \boldsymbol{\Pi}$. It holds that

$$\min_{1 \leq k \leq K} \boldsymbol{\eta}_1(k) > 0, \quad \text{and} \quad \frac{\max_{1 \leq k \leq K} \boldsymbol{\eta}_1(k)}{\min_{1 \leq k \leq K} \boldsymbol{\eta}_1(k)} \leq C,$$

for some positive constant C , where $\boldsymbol{\eta}_1(k)$ is the k -th entry of $\boldsymbol{\eta}_1$.

Conditions 4–7 are similar to those in Jin et al. (2017). In particular, Conditions 4, 5 and 7 are special cases of (2.13), (2.14) and (2.16) therein. Same as in the previous section, the degree density is measured by $\theta_{\min}^2$ and is allowed to converge to zero at rate n^{-c_3} , and our conditions accommodate the case where $|d_1|, \dots, |d_K|$ are of different orders.

Theorem 3. Assume that Conditions 1 and 4–7 hold under the degree-corrected mixed membership model (6).

i) Under the null hypothesis $H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$, we have as $n \rightarrow \infty$,

$$G_{ij} \xrightarrow{\mathcal{D}} \chi_{K-1}^2. \quad (23)$$

ii) Under the contiguous alternative hypothesis with $\lambda_2(\boldsymbol{\pi}_i \boldsymbol{\pi}_i^\top + \boldsymbol{\pi}_j \boldsymbol{\pi}_j^\top) \gg \frac{1}{n^{1-2c_2} \theta_{\min}^2}$, we have for any arbitrarily large constant $C > 0$,

$$P(G_{ij} > C) \rightarrow 1 \text{ as } n \rightarrow \infty. \quad (24)$$

A uniform result similar to (15) has also been proved in Section E of Supplementary Material under the DCMM. The test statistic G_{ij} is not directly applicable in practice due to the presence of the unknown population parameters K and $\boldsymbol{\Sigma}_2$. Nevertheless, certain consistent estimators can be constructed and the results in Theorem 3 remain valid. In particular, for the estimator $\hat{K}$ of K , we require condition (16) and for the estimator $\hat{\mathbf{S}}_2$ of $\boldsymbol{\Sigma}_2$, we need the following property

$$(n\theta_{\max}^2)^{-1} \|\mathbf{D}(\hat{\mathbf{S}}_2 - \boldsymbol{\Sigma}_2)\mathbf{D}\|_2 = o_p(1). \quad (25)$$

Theorem 4. Assume that the estimators $\hat{K}$ and $\hat{\mathbf{S}}_2$ of parameters K and $\boldsymbol{\Sigma}_2$ satisfy (16) and (25), respectively. Let $\hat{G}_{ij}$ be the test statistic constructed by replacing K and $\boldsymbol{\Sigma}_2$ with $\hat{K}$ and $\hat{\mathbf{S}}_2$, respectively. Then Theorem 3 holds with G_{ij} replaced by $\hat{G}_{ij}$ under the same conditions.

Theorem 4 suggests that with significance level α , the rejection region can be constructed as

$$\{\hat{G}_{ij} > \chi_{\hat{K}-1, 1-\alpha}^2\}. \quad (26)$$

We have similar results to Corollary 1 regarding the type I and type II errors of the above rejection region.

Corollary 2. *Assume that $\hat{K}$ and $\hat{\mathbf{S}}_2$ satisfy (16) and (25), respectively. Under the same conditions for ensuring (23), event (26) holds with asymptotic probability α . Under the same conditions for ensuring (24), event (26) holds with asymptotic probability one.*

It is worth mentioning that since the DCMM model (6) is more general than the mixed membership model (10), the test statistic $\hat{G}_{ij}$ can be applied even under model (10). However, as will be shown in our simulation studies in Section 4, the finite-sample performance of $\hat{T}_{ij}$ can be better than that of $\hat{G}_{ij}$ in such a model setting, which is not surprising since the latter involves ratios (see (19)) in its definition and has two sources of variations from both numerators and denominators. This is also reflected in losing one degree of freedom in (26)

3.3 Estimation of unknown parameters

We now discuss some consistent estimators of K , $\mathbf{\Sigma}_1$, and $\mathbf{\Sigma}_2$ that satisfy conditions (16), (17), and (25), respectively. There are some existing works concerning the estimation of parameter K . For example, Lei (2016); Chen and Lei (2018); Daudin et al. (2008); Latouche et al. (2012); Saldana et al. (2017); Wang and Bickel (2017), among others. Most of these works consider specific network models such as the stochastic block model or degree-corrected stochastic block model.

In our paper, since we consider the general DCMM model (6) which allows for mixed memberships, the existing methods are no longer applicable. To overcome the difficulty, we suggest a simple thresholding estimator defined as

$$\hat{K} = \left| \left\{ \hat{d}_i : \hat{d}_i^2 > 2.01(\log n)\check{d}_n, i \in [n] \right\} \right|, \quad (27)$$

where $|\cdot|$ stands for the cardinality of a set, the constant 2.01 can be replaced with any other constant that is slightly larger than 2, and $\check{d}_n = \max_{1 \leq l \leq n} \sum_{j=1}^n X_{lj}$ is the maximum degree of the network. That is, we count the number of eigenvalues of matrix $\mathbf{X}$ whose magnitudes exceed a certain threshold. The following lemma justifies the consistency of $\hat{K}$ defined in (27) as an estimator of the true number of communities K .

Lemma 1. *Assume that Condition 1 holds, $|d_K| \gg \sqrt{\log(n)}\alpha_n$ and $\alpha_n \geq n^{c_5}$ for some positive constant c_5 . Then $\hat{K}$ defined in (27) is consistent, that is, it satisfies condition (16).*

Observe in Theorems 1–4 that we need the condition of $K \geq 1$ for test statistic $\hat{T}_{ij}$ and the condition of $K \geq 2$ for test statistic $\hat{G}_{ij}$. Motivated by such an observation, we propose to use $\max\{\hat{K}, 1\}$ and $\max\{\hat{K}, 2\}$ as the estimated number of communities in implementing test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$, respectively.

We next discuss the estimation of $\mathbf{\Sigma}_1$ and $\mathbf{\Sigma}_2$. The following two lemmas provide the expansions of these two matrices which serve as the foundation for our proposed estimators.

Lemma 2. The (a, b) th entry of matrix Σ_1 is given by

$$\frac{1}{d_a d_b} \left\{ \sum_{t \in \{i, j\}} \sum_{l=1}^n \sigma_{tl}^2 \mathbf{v}_a(l) \mathbf{v}_b(l) - \sigma_{ij}^2 [\mathbf{v}_a(j) \mathbf{v}_b(i) + \mathbf{v}_a(i) \mathbf{v}_b(j)] \right\}, \quad (28)$$

where $\sigma_{ab}^2 = \text{var}(w_{ab})$ for $1 \leq a, b \leq n$.

Lemma 3. The (a, b) th entry of matrix Σ_2 is given by

$$\begin{aligned} & \frac{1}{t_1^2} \left\{ \sum_{l=1, l \neq j}^n \sigma_{il}^2 \left[\frac{t_1 \mathbf{v}_{a+1}(l)}{t_{a+1} \mathbf{v}_1(i)} - \frac{\mathbf{v}_{a+1}(i) \mathbf{v}_1(l)}{\mathbf{v}_1(i)^2} \right] \left[\frac{t_1 \mathbf{v}_{b+1}(l)}{t_{b+1} \mathbf{v}_1(i)} - \frac{\mathbf{v}_{b+1}(i) \mathbf{v}_1(l)}{\mathbf{v}_1(i)^2} \right] \right. \\ & + \sum_{l=1, l \neq i}^n \sigma_{jl}^2 \left[\frac{t_1 \mathbf{v}_{a+1}(l)}{t_{a+1} \mathbf{v}_1(j)} - \frac{\mathbf{v}_{a+1}(j) \mathbf{v}_1(l)}{\mathbf{v}_1(j)^2} \right] \left[\frac{t_1 \mathbf{v}_{b+1}(l)}{t_{b+1} \mathbf{v}_1(j)} - \frac{\mathbf{v}_{b+1}(j) \mathbf{v}_1(l)}{\mathbf{v}_1(j)^2} \right] \\ & + \sigma_{ij}^2 \left[\frac{t_1 \mathbf{v}_{a+1}(j)}{t_{a+1} \mathbf{v}_1(i)} - \frac{\mathbf{v}_{a+1}(i) \mathbf{v}_1(j)}{\mathbf{v}_1(i)^2} - \frac{t_1 \mathbf{v}_{a+1}(i)}{t_{a+1} \mathbf{v}_1(j)} + \frac{\mathbf{v}_{a+1}(j) \mathbf{v}_1(i)}{\mathbf{v}_1(j)^2} \right] \\ & \times \left. \left[\frac{t_1 \mathbf{v}_{b+1}(j)}{t_{b+1} \mathbf{v}_1(i)} - \frac{\mathbf{v}_{b+1}(i) \mathbf{v}_1(j)}{\mathbf{v}_1(i)^2} - \frac{t_1 \mathbf{v}_{b+1}(i)}{t_{b+1} \mathbf{v}_1(j)} + \frac{\mathbf{v}_{b+1}(j) \mathbf{v}_1(i)}{\mathbf{v}_1(j)^2} \right] \right\}. \quad (29) \end{aligned}$$

The above expansions in Lemmas 2–3 suggest that the covariance matrices Σ_1 and Σ_2 can be estimated by plugging in the sample estimates to replace the unknown population parameters. In particular, $\mathbf{v}_a$ and d_a can be estimated by $\hat{\mathbf{v}}_a$ and $\hat{d}_a$, respectively, and the last result in Lemma 12 suggests that t_k can be estimated by $\hat{d}_k$ very well. The estimation of σ_{ab}^2 is more complicated and we will discuss it in more details below.

Recall that $\sigma_{ab}^2 = \text{var}(w_{ab})$. With estimated $\hat{K}$, a naive estimator of σ_{ab}^2 is $\hat{w}_{0,ab}^2$ with $\hat{\mathbf{W}}_0 = (\hat{w}_{0,ab}) = \mathbf{X} - \sum_{k=1}^{\hat{K}} \hat{d}_k \hat{\mathbf{v}}_k \hat{\mathbf{v}}_k^T$. The good news is that it appears in (28) and (29) in the form of the average and hence the variance will be averaged out. However, this estimator is not good enough to make (17) and (25) hold due to the well-known fact that $\hat{d}_k$ is biased up. Thus we propose the following one-step refinement procedure to estimate σ_{ab}^2 , which is motivated from the higher-order asymptotic expansion of empirical eigenvalue $\hat{d}_k$ in our theoretical analysis and shrinks $\hat{d}_k$ to make the bias at a more reasonable level.

- 1). Calculate the initial estimator $\hat{\mathbf{W}}_0 = \mathbf{X} - \sum_{k=1}^{\hat{K}} \hat{d}_k \hat{\mathbf{v}}_k \hat{\mathbf{v}}_k^T$.
- 2). With the initial estimator $\hat{\mathbf{W}}_0$, update the estimator of eigenvalue d_k as

$$\tilde{d}_k = \left[\frac{1}{\hat{d}_k} + \frac{\hat{\mathbf{v}}_k^T \text{diag}(\hat{\mathbf{W}}_0^2) \hat{\mathbf{v}}_k}{\hat{d}_k^3} \right]^{-1}.$$

- 3). Then update the estimator of $\mathbf{W}$ as $\hat{\mathbf{W}} \equiv (\hat{w}_{ij}) = \mathbf{X} - \sum_{k=1}^{\hat{K}} \tilde{d}_k \hat{\mathbf{v}}_k \hat{\mathbf{v}}_k^T$ and estimate σ_{ab}^2 as $\hat{\sigma}_{ab}^2 = \hat{w}_{ab}^2$.

To summarize, we propose to estimate matrix Σ_1 by replacing d_k , $\mathbf{v}_k$, and σ_{ab}^2 with $\hat{d}_k$, $\hat{\mathbf{v}}_k$, and $\hat{\sigma}_{ab}^2$, respectively, in (28). The covariance matrix Σ_2 can be estimated in a similar

way by replacing t_k , $\mathbf{v}_k$, and σ_{ab}^2 with $\hat{d}_k$, $\hat{\mathbf{v}}_k$, and $\hat{\sigma}_{ab}^2$, respectively, in (29). Denote by $\hat{\mathbf{S}}_1$ and $\hat{\mathbf{S}}_2$ the resulting estimators, respectively. The following lemma justifies the effectiveness of these two estimators.

Theorem 5. *Under Conditions 1–3, estimator $\hat{\mathbf{S}}_1$ satisfies condition (17). Under Conditions 1 and 4–7, estimator $\hat{\mathbf{S}}_2$ satisfies condition (25).*

We remark that through higher moments calculations, it can be proved that (17) and (25) uniformly hold for all $1 \leq i \neq j \leq n$.

4 Simulation studies

We use simulation examples to examine the finite-sample performance of our new SIMPLE test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$ with true and estimated numbers of communities K , respectively. In particular, we consider the following two model settings.

Model 1: the mixed membership model (10). We consider $K = 3$ communities, where there are n_0 pure nodes within each community. Thus for the k th community, the community membership probability vector for each pure node is $\boldsymbol{\pi} = \mathbf{e}_k \in \mathbb{R}^K$. The remaining $n - 3n_0$ nodes are divided equally into 4 groups, where within the l th group all nodes have mixed memberships with community membership probability vector $\mathbf{a}_l$, $l = 1, \dots, 4$. We set $\mathbf{a}_1 = (0.2, 0.6, 0.2)^T$, $\mathbf{a}_2 = (0.6, 0.2, 0.2)^T$, $\mathbf{a}_3 = (0.2, 0.2, 0.6)^T$, and $\mathbf{a}_4 = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})^T$. Matrix $\mathbf{P}$ has diagonal entries one and (i, j) th entry equal to $\frac{\rho}{|i-j|}$ for $i \neq j$. We experiment with two sets of parameters $(\rho, n, n_0) = (0.2, 3000, 500)$ and $(0.2, 1500, 300)$, and vary the value of θ from 0.2 to 0.9 with step size 0.1. It is clear that parameter θ has direct impact on the average degree and hence measures the signal strength.

Model 2: the DCMM model (6). Both matrices $\boldsymbol{\Pi}$ and $\mathbf{P}$ are the same as in Model 1. For the degree heterogeneity matrix $\boldsymbol{\Theta} = \text{diag}(\theta_1, \dots, \theta_n)$, we simulate $\frac{1}{\theta_i}$ as independent and identically distributed (i.i.d.) random variables from the uniform distribution on $[\frac{1}{r}, \frac{2}{r}]$ with $r \in (0, 1]$. We consider different choices of r with $r^2 \in \{0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9\}$. We can see that as parameter r^2 increases, the signal becomes stronger.

4.1 Hypothesis testing with K known

Recall that our test statistics are designed to test the membership information for each preselected pair of nodes (i, j) with $1 \leq i \neq j \leq n$. To examine the empirical size of our tests, we preselect (i, j) as two nodes with community membership probability vector $(0.2, 0.6, 0.2)^T$. To examine the empirical power of our tests, we preselect i as a node with community membership probability vector $(0.2, 0.6, 0.2)^T$ and j as a node with community membership probability vector $(0, 1, 0)^T$. The nominal significance level is set to be 0.05 when calculating the critical points and the number of repetitions is chosen as 500.

We first generate simulated data from Model 1 introduced above and examine the empirical size and power of test statistic $\hat{T}_{ij}$ with estimated $\mathbf{\Sigma}_1$, but with the true value of K . Then we consider Model 2 and examine the empirical size and power of test statistic $\hat{G}_{ij}$ with estimated $\mathbf{\Sigma}_2$ and the true value of K . The empirical size and power at different signal levels are reported in Tables 1 and 2, corresponding to sample sizes $n = 1500$ and 3000 , respectively. As shown in Tables 1 and 2, the size and power of our tests converge quickly to the nominal significance level 0.05 and the value of one, respectively, as the signal strength θ (related to effective sample size) increases. As demonstrated in Figure 1, the empirical null distributions are well described by our theoretical results. These results provide stark empirical evidence supporting our theoretical findings, albeit complicated formulas (28) and (29).

Table 1: The size and power of test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$ when the true value of K is used. The nominal level is 0.05 and sample size is $n = 1500$.

Model 1	θ	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
	Size	0.058	0.046	0.06	0.05	0.05	0.058	0.036	0.05
	Power	0.734	0.936	0.986	0.998	1	1	1	1
Model 2	r^2	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
	Size	0.076	0.062	0.072	0.062	0.074	0.046	0.044	0.056
	Power	0.426	0.562	0.696	0.77	0.89	0.93	0.952	0.976

Table 2: The size and power of test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$ when the true value of K is used. The nominal level is 0.05 and sample size is $n = 3000$.

Model 1	θ	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
	Size	0.082	0.066	0.052	0.052	0.044	0.042	0.038	0.062
	Power	0.936	0.994	1	1	1	1	1	1
Model 2	r^2	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
	Size	0.082	0.06	0.062	0.058	0.062	0.066	0.064	0.06
	Power	0.67	0.842	0.918	0.972	0.99	1	1	1

Figure 1 presents how the asymptotic null distributions change with sample size n when $\theta = 1/(2 \log n)$ and $r^2 = 1/(2 \log n)$, respectively, for Model 1 and Model 2. It is seen that the network become sparser as its size increases. The top panel shows the histogram plots when $n = 1500$ and the bottom panel corresponds to $n = 3000$. One can observe that as sample size increases, the χ^2 -distribution fits the empirical null distribution better, which is consistent with our theoretical results.

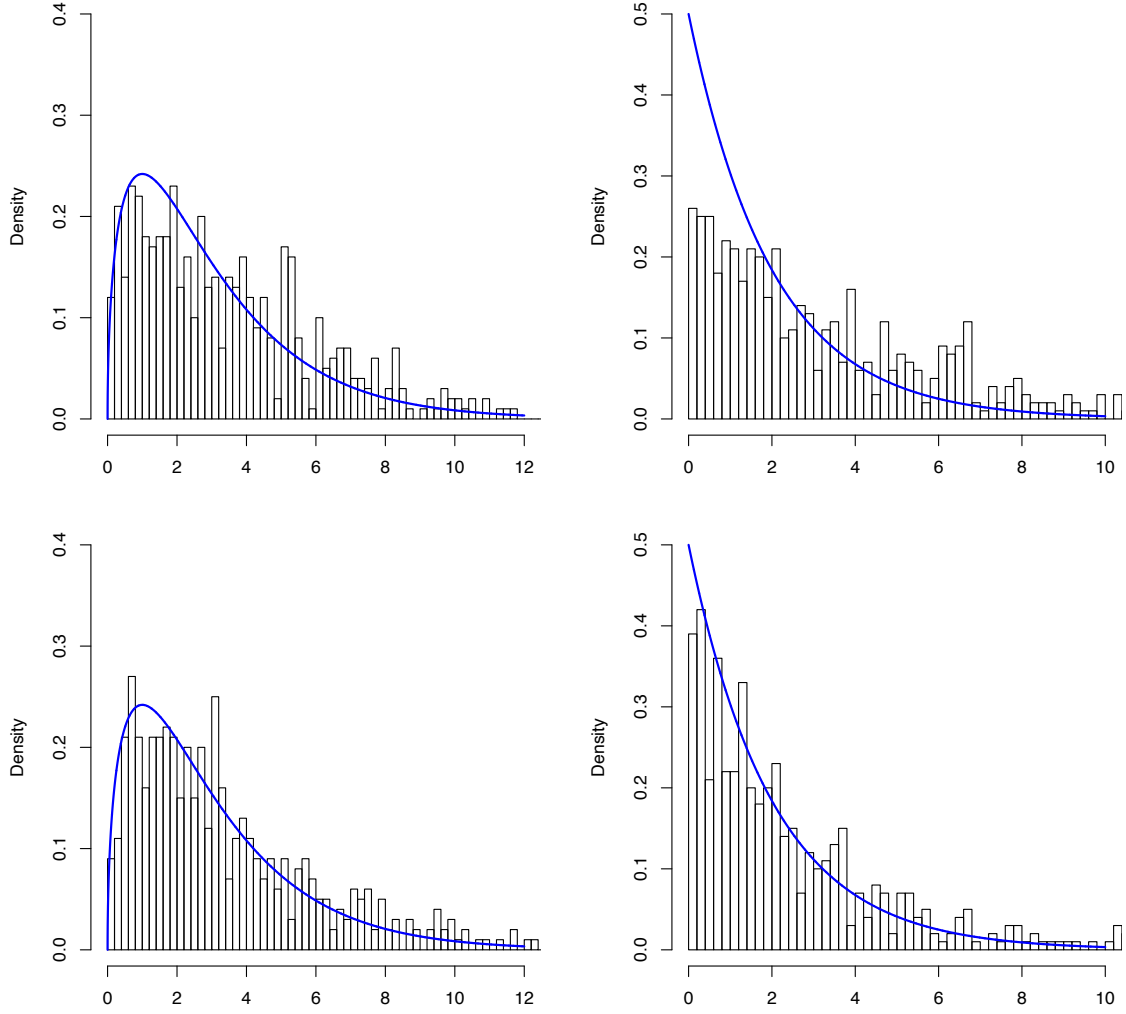


Figure 1: Left: the histogram of test statistic $\hat{T}_{ij}$ under null hypothesis with known K when $\theta = \frac{1}{2\log n}$. Blue curve is the density function of χ^2_3 . Right: the histogram of test statistic $\hat{G}_{ij}$ under null hypothesis with known K when $r^2 = \frac{1}{2\log n}$. Blue curve is the density function of χ^2_2 . Top panel is for sample size $n = 1500$ and bottom panel is for sample size $n = 3000$. Here $n_0 = \frac{n}{5}$.

4.2 Hypothesis testing with estimated K

We now examine the finite-sample performance of our test statistics $\hat{T}_{ij}$ and estimated $\hat{G}_{ij}$ with estimated K . The simulation settings are identical to those in Section 4.1 except that we explore only the setting with sample size $n = 3000$.

In Table 3, we report the proportion of correctly estimated K using the thresholding rule (27) in both simulation settings of Models 1 and 2. It is seen that as the signal becomes stronger (i.e., as θ or r^2 increases), the estimation accuracy becomes higher. We also observe that for relatively weak signals, the thresholding rule in (27) tends to underestimate K , resulting in low estimation accuracy. We can see from the same table that over all repetitions,

K is either correctly estimated or underestimated. The critical values are constructed based on these estimated values of K .

Table 3: Estimation accuracy of K using the thresholding rule (27)

	θ or r^2	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
Model 1	$P(\hat{K} = K)$	1	1	1	1	1	1	1	1
	$P(\hat{K} \leq K)$	1	1	1	1	1	1	1	1
Model 2	$P(\hat{K} = K)$	0	0	0	1	1	1	1	1
	$P(\hat{K} \leq K)$	1	1	1	1	1	1	1	1

Table 4: The size and power of test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$ when the estimated value of K is used. The nominal level is 0.05 and sample size is $n = 3000$.

	θ	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
Model 1	Size	0.082	0.066	0.052	0.052	0.044	0.042	0.038	0.062
	Power	0.936	0.994	1	1	1	1	1	1
	r^2	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9
Model 2	Size	0.054	0.058	0.062	0.058	0.062	0.066	0.064	0.06
	Power	0.074	0.042	0.918	0.972	0.99	1	1	1

Same as in Section 4.1, we also examine the empirical size and power of our tests at different levels of signal strength. The results are presented in Table 4. It is seen that the performance of $\hat{T}_{ij}$ is identical to that in Table 2, and the performance of $\hat{G}_{ij}$ is the same as in Table 2 for all $r^2 > 0.3$. This is expected because of the nearly perfect estimation of K as shown in Table 3 in these scenarios and/or the relatively strong signal strength. When $r^2 \leq 0.3$, $\hat{G}_{ij}$ has poor power because of the underestimated K (see Table 3). Nevertheless, we observe the same trend as the signal strength increases, which provides support for our theoretical results. We have also applied our tests to nodes with more distinct membership probability vectors $(0.2, 0.6, 0.2)^T$ and $(0, 0, 1)^T$, and the impact of estimated K is much smaller. These additional simulation results are available upon request.

5 Real data applications

5.1 U.S. political data

The U.S. political data set consists of 105 political books sold by an online bookseller in the year of 2004. Each book is represented by a node and links between nodes represent the frequency of co-purchasing of books by the same buyers. The network was compiled by V. Krebs (source: <http://www.orgnet.com>). The books have been assigned manually three labels (conservative, liberal, and neutral) by M. E. J. Newman based on the reviews and

descriptions of the books. Note that such labels may not be very accurate. In fact, as argued in multiple papers (e.g., [Koutsourelakis and Eliassi-Rad \(2008\)](#)), the mixed membership model may better suit this data set.

Since our SIMPLE tests $\hat{T}_{ij}$ and $\hat{G}_{ij}$ do not differentiate network models with or without mixed memberships, we will view the network as having $K = 2$ communities (conservative and liberal) and treat the neutral nodes as having mixed memberships. To connect our results with the literature, we consider the same 9 books reported in [Jin et al. \(2017\)](#). Another reason of considering the same 9 books as in [Jin et al. \(2017\)](#) is that our test statistic $\hat{G}_{ij}$ is constructed using the SCORE statistic which is closely related to [Jin et al. \(2017\)](#). The book names as well as labels (provided by Newman) are reported in Table 5. The p-values based on test statistics $\hat{T}_{ij}$ and $\hat{G}_{ij}$ for testing the pairwise membership profiles of these 9 nodes are summarized in Tables 6 and 7, respectively.

From Table 7, we see that our results based on test statistic $\hat{G}_{ij}$ are mostly consistent with the labels provided by Newman and also very consistent with those in Table 5 of [Jin \(2015\)](#). For example, books 59 and 50 are both labeled as “conservative” by Newman and our tests return large p-values between them. These two books generally have much smaller p-values with books labeled as “neutral.” Book 78, which was labeled as “conservative” by Newman, seems to be more similar to some neutral books. This phenomenon was also observed in [Jin et al. \(2017\)](#), who interpreted this as a result of having a liberal author. Among the nodes labeled by Newman as “neutral,” “All the Shah’s Men,” or book 29, has relatively larger p-values with conservative books. However, this book has even larger p-values with some other neutral books such as book 104, “The Future of Freedom,” which is consistent with the results in [Jin et al. \(2017\)](#) who reported that these two books have very close membership probability vectors. In summary, our SIMPLE method provides statistical significance for the membership probability vectors estimated in [Jin et al. \(2017\)](#).

For a summary of our testing results, we also provide the multidimensional scaling map of the nodes based on test statistics $\hat{G}_{ij}$ on the left panel of Figure 2. The graph on the right panel of Figure 2 is defined by the pairwise p-value matrix calculated from $\hat{G}_{ij}$. Specifically, we first apply the hard-thresholding to the p-value matrix by setting all entries below 0.05 to 0. Denote by $\tilde{P}$ the resulting matrix. Then we plot the graph using the entries of $\tilde{P}$ as edge weights so that zeros correspond to unconnected pairs of nodes and larger entries mean more closely connected nodes with thicker edges. The nodes in both graphs are color coded according to Newman’s labels, with red representing “conservative,” blue representing “liberal,” and orange representing “neutral.” It is seen that both graphs are mostly consistent with Newman’s labels, with a few exceptions as partially discussed before. We also would like to mention that the hard-thresholding step in p-value graph is to make the graph less dense and easier to view. In fact, a small perturbation of the threshold does not change much of the overall layout of the graph.

Table 5: Political books with labels

Title	Label (by Newman)	Node index
Empire	Neutral	105
The Future of Freedom	Neutral	104
Rise of the Vulcans	Conservative	59
All the Shah's Men	Neutral	29
Bush at War	Conservative	78
Plan of Attack	Neutral	77
Power Plays	Neutral	47
Meant To Be	Neutral	19
The Bushes	Conservative	50

Table 6: P-values based on test statistics $\hat{T}_{ij}$. The labels provided by Newman are in the parentheses.

Node No.	105(N)	104(N)	59(C)	29(N)	78(C)	77(N)	47(N)	19(N)	50(C)
105(N)	1.0000	0.6766	0.0298	0.3112	0.0248	0.0000	0.0574	0.1013	0.0449
104(N)	0.6766	1.0000	0.0261	0.2487	0.0204	0.0000	0.0643	0.1184	0.0407
59(C)	0.0298	0.0261	1.0000	0.1546	0.2129	0.0013	0.0326	0.0513	0.9249
29(N)	0.3112	0.2487	0.1546	1.0000	0.3206	0.0034	0.0236	0.0497	0.2121
78(C)	0.0248	0.0204	0.2129	0.3206	1.0000	0.0991	0.0042	0.0084	0.2574
77(N)	0.0000	0.0000	0.0013	0.0034	0.0991	1.0000	0.0000	0.0000	0.0035
47(N)	0.0574	0.0643	0.0326	0.0236	0.0042	0.0000	1.0000	0.9004	0.0834
19(N)	0.1013	0.1184	0.0513	0.0497	0.0084	0.0000	0.9004	1.0000	0.1113
50(C)	0.0449	0.0407	0.9249	0.2121	0.2574	0.0035	0.0834	0.1113	1.0000

Table 7: P-values based on test statistics $\hat{G}_{ij}$. The labels provided by Newman are in the parentheses.

Node No.	105(N)	104(N)	59(C)	29(N)	78(C)	77(N)	47(N)	19(N)	50(C)
105(N)	1.0000	0.4403	0.1730	0.4563	0.8307	0.5361	0.0000	0.0000	0.1920
104(N)	0.4403	1.0000	0.0773	0.9721	0.3665	0.6972	0.0000	0.0000	0.1144
59(C)	0.1730	0.0773	1.0000	0.0792	0.1337	0.0885	0.0000	0.0000	0.8141
29(N)	0.4563	0.9721	0.0792	1.0000	0.4256	0.7624	0.0000	0.0000	0.1153
78(C)	0.8307	0.3665	0.1337	0.4256	1.0000	0.5402	0.0000	0.0000	0.1591
77(N)	0.5361	0.6972	0.0885	0.7624	0.5402	1.0000	0.0000	0.0000	0.1294
47(N)	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	1.0000	0.9778	0.0000
19(N)	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.9778	1.0000	0.0000
50(C)	0.1920	0.1144	0.8141	0.1153	0.1591	0.1294	0.0000	0.0000	1.0000

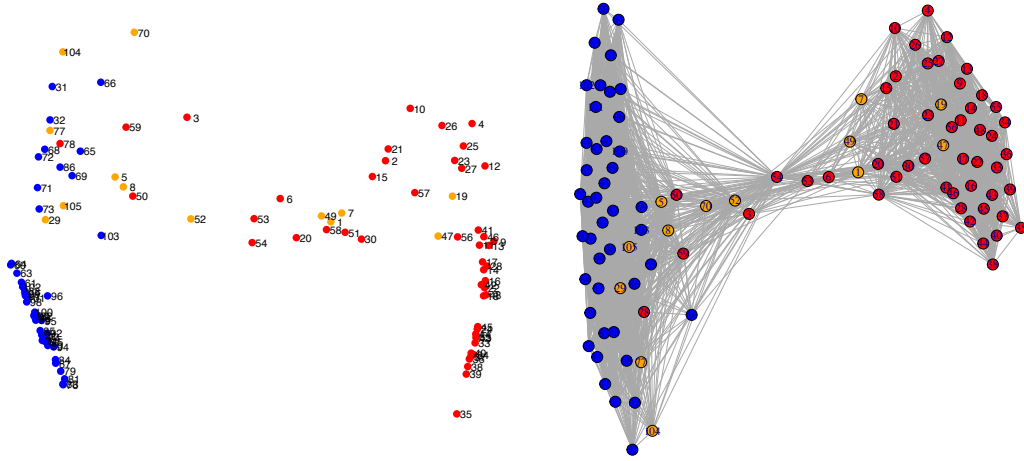


Figure 2: Left panel: the multidimensional scaling map of the nodes based on test statistics $\hat{G}_{ij}$. Right panel: the connectivity graph generated from the thresholded p-value matrix based on $\hat{G}_{ij}$. The nodes are color coded according to Newman’s labels, with red representing “conservative,” blue representing “liberal,” and orange representing “neutral.”

5.2 Stock data

We consider a larger network of stocks in this section. Specifically, daily prices of stocks in the S&P 500 from the period of January 2, 2009 to December 30, 2019 were collected and converted into log returns. After some pre-processing (e.g., removing stocks with missing values or very low node degrees), we ended up with 404 stocks. All data analyses in this section were conducted using those 404 stocks. It is well known that much variation in stock excess returns can be captured by factors such as the Fama–French three factors. We first remove these common factors by fitting a factor model, and then the adjacency matrix of stocks is constructed as the correlation matrix of idiosyncratic components from the factor model.

Since stocks are commonly believed to have heterogeneous node degrees, we only apply $\hat{G}_{ij}$ to the constructed adjacency matrix. The estimated number of communities is $\hat{K} = 3$. For each pair of stocks, we calculate its p-value using $\hat{G}_{ij}$ and the asymptotic null distribution χ^2_2 . This forms a p-value matrix, denoted as $\mathbf{A}$. To better visualize the results, we provide the multiscale plot of the distance matrix $\mathbf{1} - \mathbf{A}$ with $\mathbf{1}$ the matrix with all entries being 1, and present the results in Figure 3. It is seen that the scatter plot roughly has three legs and a central cluster. The three legs can be interpreted as the three communities with nodes having relatively more pure membership profiles, and the central cluster can be understood as for nodes with mixed membership profiles. For easier visualization, we provide zoomed plots for the three legs and the central cluster in Figure 4. The first three subplots a)–c) correspond to the three legs, and the last subplot d) corresponds to the central cluster. We observe some interesting clustering effects. Figure 4a) corresponds to the top leg in

	HD	L	AAPL	INTC	MCHP	AEE	NEE	EVRG	ADBE
TGT	0.29643	0.71361	0.00000	0.00000	0.00000	0.00000	0.00000	0.00000	0.00033
HD	1.00000	0.14934	0.00000	0.00000	0.00000	0.00000	0.00000	0.00000	0.00025
L	0.14934	1.00000	0.00000	0.00000	0.00000	0.00000	0.00000	0.00000	0.00031
AAPL	0.00000	0.00000	1.00000	0.00780	0.01933	0.00004	0.00003	0.00010	0.00395
INTC	0.00000	0.00000	0.00780	1.00000	0.00024	0.00000	0.00000	0.00000	0.00148
MCHP	0.00000	0.00000	0.01933	0.00024	1.00000	0.00000	0.00000	0.00000	0.00202
AEE	0.00000	0.00000	0.00004	0.00000	0.00000	1.00000	0.93719	0.46490	0.00467
NEE	0.00000	0.00000	0.00003	0.00000	0.00000	0.93719	1.00000	0.24407	0.00467
EVRG	0.00000	0.00000	0.00010	0.00000	0.00000	0.46490	0.24407	1.00000	0.00465
ADBE	0.00025	0.00031	0.00395	0.00148	0.00202	0.00467	0.00467	0.00465	1.00000

Table 8: The p-value matrix for selected stocks.

Figure 3. When it is far away from the central cluster (i.e., top left of this subplot), we have stocks mostly related to the retail and restaurant industry (e.g., TGT, HD, LOW, DRI), and when it moves closer to the central cluster (i.e, bottom right of this subplot), the companies are mostly in the real estate (e.g., EXR, VTR, PSA, AVB, PLD). Figure 4b) mostly consists of tech companies such as AAPL, MCHP, MU, INTC, XLNX, QCOM, ADI, among many others in similar category. Figure 4c) roughly has two subclusters. The left cluster mostly consists of companies in or related to the health industry such as DGX, VAR, GLW, MDT, CERN, TEL, UNH, PFE, BMY, and many other similar ones. The right cluster has predominately companies in the energy industry such as AEE, NEE, EVRG, PNW, DUK, LNT, LNT, ES. Figure 4d) is a zoomed plot that roughly shows the central cluster. It contains a wide range of companies including, but not limited to, risk management and investment companies (BEN, HIG, NDAQ), transportation industry (AAL, NSC, UAL), and communication industry (CTL, VRSN, CTXS).

In Table 8 below, we also present the p-value matrix for selected stocks. The first three stocks (TGT, HD, L) are all in the retail industry, the next three stocks (AAPL, INTC, MCHP) are all in the tech industry, stocks 7 to 9 (AEE, NEE, EVRG) are all in the energy industry, and the remaining one (ADBE) is taken from the central cluster. It is seen that the first three groups of stocks have high pairwise p-values within groups, but almost zero p-values with stocks from other groups. In particular, Adobe (ADBE) seems to be connected to most of these selected stocks, which is consistent with the common sense. We would also like to point out that these results were obtained after removing the three common factors from the stock returns, and the clustering structure discovered here should be interpreted as complementary to the ones already captured by the factors.

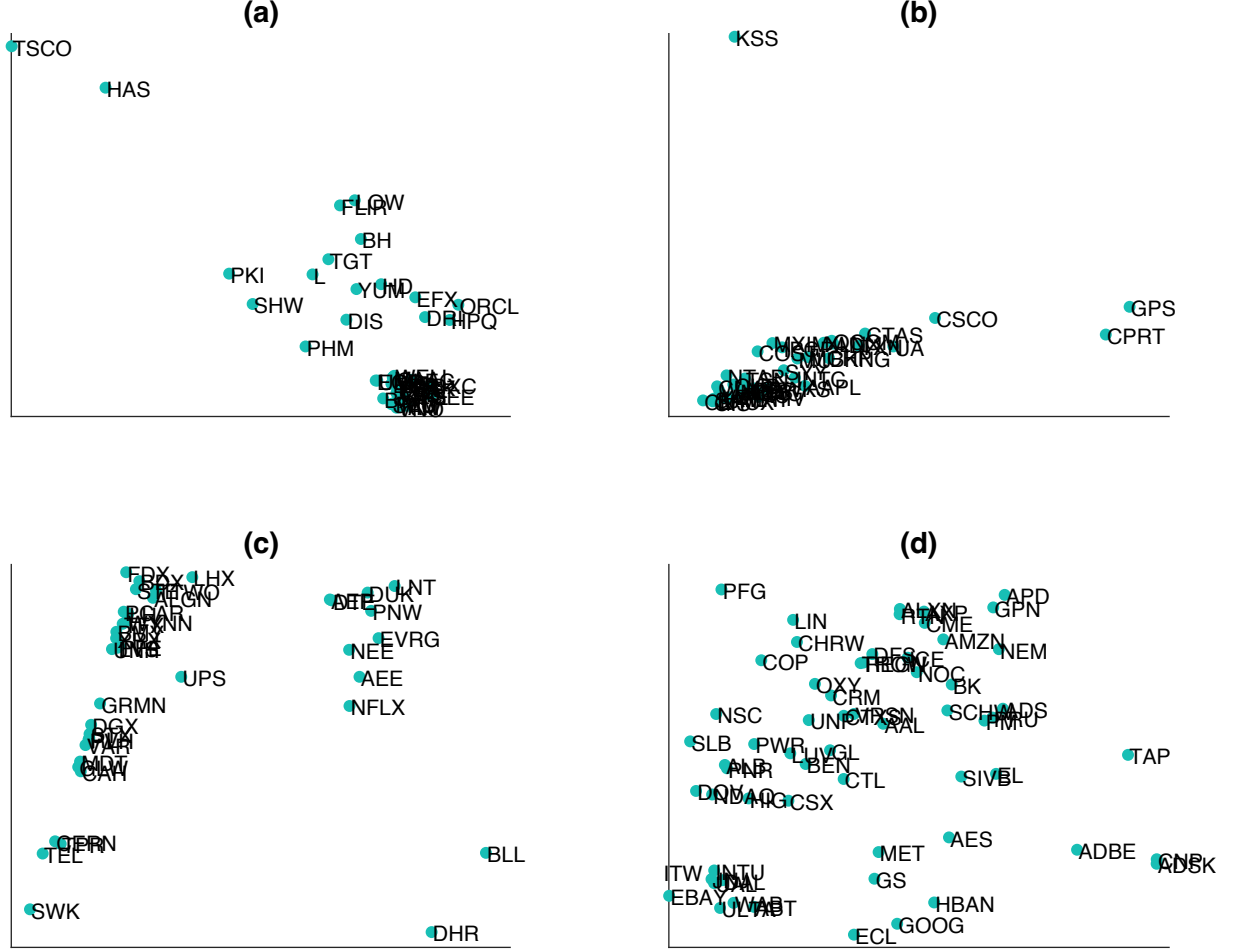


Figure 4: Zoomed multiscale plots based on the distance matrix $\mathbf{1} - \mathbf{A}$, where $\mathbf{A}$ is the p-value matrix based on $\hat{G}_{ij}$.

all the nodes within the selected set indeed share the same membership profile information. It would also be interesting to quantify and control the statistical inference error rates when one is interested in performing a set of hypothesis tests simultaneously for network data. Moreover, it would be interesting to investigate the hypothesis testing problem for more general network models as well as for statistical models beyond network data such as for large collections of text documents.

In addition, it would be interesting to connect the growing literature on sparse covariance matrices and sparse precision matrices with that on network models. Such connections can be made via modeling the graph Laplacian through a precision matrix or covariance matrix (Brownlees et al., 2019). A natural question is then how well the network profiles can be inferred from a panel of time series data. The same question also arises if the panel of time series data admits a factor structure (Fan et al., 2008, 2013). These problems and extensions are beyond the scope of the current paper and will be interesting topics for future research.

References

- Abbe, E. (2017). Community detection and stochastic block models: recent developments. *Journal of Machine Learning Research* 18, 177:1–177:86.
- Abbe, E., J. Fan, K. Wang, and Y. Zhong (2017). Entrywise eigenvector analysis of random matrices with low expected rank. *arXiv preprint arXiv:1709.09565*.
- Airoldi, E. M., D. M. Blei, S. E. Fienberg, and E. P. Xing (2008). Mixed membership stochastic blockmodels. *Journal of Machine Learning Research* 9, 1981–2014.
- Arias-Castro, E. and N. Verzelen (2014, 06). Community detection in dense random networks. *Ann. Statist.* 42(3), 940–969.
- Bickel, P. J. and P. Sarkar (2016). Hypothesis testing for automated community detection in networks. *Journal of the Royal Statistical Society Series B* 78, 253–273.
- Brownlees, C., G. G. Stefan, and G. Lugosi (2019). Community detection in partial correlation network models. *Manuscript*.
- Chen, K. and J. Lei (2018). Network cross-validation for determining the number of communities in network data. *Journal of the American Statistical Association* 113, 241–251.
- Daudin, J.-J., F. Picard, and S. Robin (2008). A mixture model for random graphs. *Statistics and Computing* 18, 173–183.
- Fan, J., Y. Fan, X. Han, and J. Lv (2020). Asymptotic theory of eigenvectors for random matrices with diverging spikes. *Journal of the American Statistical Association*, to appear.
- Fan, J., Y. Fan, and J. Lv (2008). High dimensional covariance matrix estimation using a factor model. *Journal of Econometrics* 147, 186–197.
- Fan, J., Y. Liao, and M. Mincheva (2013). Large covariance estimation by thresholding principal orthogonal complements (with discussion). *Journal of the Royal Statistical Society Series B* 75, 603–680.
- Gao, C., Z. Ma, A. Y. Zhang, and H. H. Zhou (2018, 10). Community detection in degree-corrected block models. *Ann. Statist.* 46(5), 2153–2185.
- Goldenberg, A., A. X. Zheng, S. E. Fienberg, and E. M. Airoldi (2010). A survey of statistical network models. *Found. Trends Mach. Learn.* 2, 129–233.
- Holland, P. W., K. B. Laskey, and S. Leinhardt (1983). Stochastic blockmodels: First steps. *Social Networks* 5, 109–137.
- Jin, J. (2015). Fast community detection by SCORE. *Ann. Statist.* 43, 57–89.

- Jin, J. and Z. T. Ke (2017). A sharp lower bound for mixed-membership estimation. *arXiv preprint arXiv:1709.05603*.
- Jin, J., Z. T. Ke, and S. Luo (2017). Estimating network memberships by simplex vertex hunting. *arXiv preprint arXiv:1708.07852*.
- Karrer, B. and M. E. J. Newman (2011). Stochastic blockmodels and community structure in networks. *Phys. Rev. E* *83*, 016107.
- Koutsourelakis, P.-S. and T. Eliassi-Rad (2008). Finding mixed-memberships in social networks. *AAAI Spring Symposium: Social Information Processing*, 48–53.
- Latouche, P., E. Birmelé, and C. Ambroise (2012). Variational bayesian inference and complexity control for stochastic block models. *Statistical Modelling* *12*, 93–115.
- Lei, J. (2016). A goodness-of-fit test for stochastic block models. *Ann. Statist.* *44*, 401–424.
- Lei, J. and A. Rinaldo (2015). Consistency of spectral clustering in stochastic block models. *Ann. Statist.* *43*, 215–237.
- Liben-Nowell, D. and J. Kleinberg (2007). The link-prediction problem for social networks. *J. Am. Soc. Inf. Sci. Technol.* *58*, 1019–1031.
- Newman, M. E. J. (2013a). Community detection and graph partitioning. *EPL (Europhysics Letters)* *103*, 28003.
- Newman, M. E. J. (2013b). Spectral methods for community detection and graph partitioning. *Phys. Rev. E* *88*, 042822.
- Newman, M. E. J. and T. P. Peixoto (2015). Generalized communities in networks. *Phys. Rev. Lett.* *115*, 088701.
- Rohe, K., S. Chatterjee, and B. Yu (2011). Spectral clustering and the high-dimensional stochastic blockmodel. *The Annals of Statistics* *39*, 1878–1915.
- Saldana, D., Y. Yu, and Y. Feng (2017). How many communities are there? *Journal of Computational and Graphical Statistics* *26*, 171–181.
- Verzelen, N. and E. Arias-Castro (2015, 12). Community detection in sparse random networks. *Ann. Appl. Probab.* *25*(6), 3465–3510.
- von Luxburg, U. (2007). A tutorial on spectral clustering. *Statistics and Computing* *17*, 395–416.
- Wang, Y. J. and G. Y. Wong (1987). Stochastic blockmodels for directed graphs. *Journal of the American Statistical Association* *82*, 8–19.

- Wang, Y. X. R. and P. J. Bickel (2017). Likelihood-based model selection for stochastic block models. *Ann. Statist.* *45*, 500–528.
- Wu, Y.-J., E. Levina, and J. Zhu (2018). Link prediction for egocentrically sampled networks. *arXiv preprint arXiv:1803.04084*.
- Zhang, P. and C. Moore (2014). Scalable detection of statistically significant communities and hierarchies, using message passing for modularity. *Proceedings of the National Academy of Sciences* *111*, 18144–18149.
- Zhang, Y., E. Levina, and J. Zhu (2020). Detecting overlapping communities in networks using spectral methods. *SIAM Journal on Mathematics of Data Science* *2*(2), 265–283.
- Zhao, Y., E. Levina, and J. Zhu (2012). Consistency of community detection in networks under degree-corrected stochastic block models. *Ann. Statist.* *40*, 2266–2292.

Supplementary Material to “SIMPLE: Statistical Inference on Membership Profiles in Large Networks”

Jianqing Fan¹, Yingying Fan², Xiao Han³ and Jinchi Lv²

Princeton University¹, University of Southern California²
and University of Science and Technology of China³

September 1, 2021

This Supplementary Material contains all the proofs and technical details.

A Proofs of main results

To facilitate the technical presentation, we list two definitions below, where n represents the network size and dimensionality of eigenvectors.

Definition 1. Let ζ and ξ be a pair of random variables that may depend on n . We say that they satisfy $\xi = O_{\prec}(\zeta)$ if for any pair of positive constants (a, b) , there exists some positive integer $n_0(a, b)$ depending only on a and b such that $\mathbb{P}(|\xi| > n^a |\zeta|) \leq n^{-b}$ for all $n \geq n_0(a, b)$.

Definition 2. We say that an event $\mathfrak{A}_n$ holds with high probability if for any positive constant a , there exists some positive integer $n_0(a)$ depending only on a such that $\mathbb{P}(\mathfrak{A}_n) \geq 1 - n^{-a}$ for all $n \geq n_0(a)$.

From Definitions 1 and 2 above, we can see that if $\xi = O_{\prec}(\zeta)$, then it holds that $\xi = O(n^a |\zeta|)$ with high probability for any positive constant a . The strong probabilistic bounds in the statements of Definitions 1 and 2 are in fact consequences of analyzing large binary random matrices given by networks.

Let us introduce some additional notation. Since the eigenvectors are always up to a sign change, for simplicity we fix the orientation of the empirical eigenvector $\widehat{\mathbf{v}}_k$ such that $\widehat{\mathbf{v}}_k^T \mathbf{v}_k \geq 0$ for each $1 \leq k \leq K$, where $\mathbf{v}_k$ is the k th population eigenvector of the low-rank mean matrix $\mathbf{H}$ in our general network model (2). It is worth mentioning that all the variables are real-valued throughout the paper except that variable z can be complex-valued. For any nonzero complex number z , deterministic matrices $\mathbf{M}_1$ and $\mathbf{M}_2$ of appropriate dimensions,

$1 \leq k \leq K$, and n -dimensional unit vector $\mathbf{u}$, we define

$$\mathcal{P}(\mathbf{M}_1, \mathbf{M}_2, z) = z\mathcal{R}(\mathbf{M}_1, \mathbf{M}_2, z), \quad \tilde{\mathcal{P}}_{k,z} = [z^2(A_{\mathbf{v}_k,k,z}/z)']^{-1}, \quad (1)$$

$$\mathbf{b}_{\mathbf{u},k,z} = \mathbf{u} - \mathbf{V}_{-k} [(\mathbf{D}_{-k})^{-1} + \mathcal{R}(\mathbf{V}_{-k}, \mathbf{V}_{-k}, z)]^{-1} \mathcal{R}^T(\mathbf{u}, \mathbf{V}_{-k}, z), \quad (2)$$

where $\mathcal{R}(\mathbf{M}_1, \mathbf{M}_2, z)$ is defined in (7),

$$A_{\mathbf{u},k,z} = \mathcal{P}(\mathbf{u}, \mathbf{v}_k, z) - \mathcal{P}(\mathbf{u}, \mathbf{V}_{-k}, z) [z(\mathbf{D}_{-k})^{-1} + \mathcal{P}(\mathbf{V}_{-k}, \mathbf{V}_{-k}, z)]^{-1} \mathcal{P}(\mathbf{V}_{-k}, \mathbf{v}_k, z), \quad (3)$$

$(A_{\mathbf{v}_k,k,z}/z)'$ denotes the derivative of $A_{\mathbf{v}_k,k,z}/z$ with respect to complex variable z , $\mathbf{V}_{-k}$ represents a submatrix of $\mathbf{V} = (\mathbf{v}_1, \dots, \mathbf{v}_K)$ by removing the k th column, and $\mathbf{D}_{-k}$ stands for a principal submatrix of $\mathbf{D} = \text{diag}(d_1, \dots, d_K)$ by removing the k th diagonal entry.

A.1 Proof of Theorem 1

We first prove the conclusion in the first part of Theorem 1 under the null hypothesis $H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$, where (i, j) with $1 \leq i < j \leq n$ represents a given pair of nodes in the network. In particular, Lemma 6 in Section B.8 of Supplementary Material plays a key role in the technical analysis. For the given pair (i, j) , let us define a new random matrix $\tilde{\mathbf{X}} = (\tilde{\mathbf{x}}_{lm})_{1 \leq l, m \leq n}$ based on the original random matrix $\mathbf{X} = (\mathbf{x}_{lm})_{1 \leq l, m \leq n}$ by swapping the roles of nodes i and j , namely by setting

$$\tilde{\mathbf{x}}_{lm} = \begin{cases} \mathbf{x}_{lm}, & l, m \in \{i, j\}^c \\ \mathbf{x}_{im}, & l = j, m \in \{i, j\}^c \\ \mathbf{x}_{jm}, & l = i, m \in \{i, j\}^c \\ \mathbf{x}_{li}, & m = j, l \in \{i, j\}^c \\ \mathbf{x}_{lj}, & m = i, l \in \{i, j\}^c \end{cases} \quad \text{and} \quad \tilde{\mathbf{x}}_{lm} = \begin{cases} \mathbf{x}_{ij}, & (l, m) = (i, j) \text{ or } (j, i) \\ \mathbf{x}_{ii}, & l = m = j \\ \mathbf{x}_{jj}, & l = m = i \end{cases}, \quad (4)$$

where $\{i, j\}^c$ stands for the complement of set $\{i, j\}$ in the node set $\{1, \dots, n\}$. It is easy to see that the new symmetric random matrix $\tilde{\mathbf{X}}$ defined in (4) is simply the adjacency matrix of a network given by the mixed membership model (10) by swapping the i th and j th rows, $\boldsymbol{\pi}_i$ and $\boldsymbol{\pi}_j$, of the community membership probability matrix $\boldsymbol{\Pi} = (\boldsymbol{\pi}_1, \dots, \boldsymbol{\pi}_n)^T$.

By the above definition of $\tilde{\mathbf{X}}$, we can see that under the null hypothesis $H_0 : \boldsymbol{\pi}_i = \boldsymbol{\pi}_j$, it holds that

$$\tilde{\mathbf{X}} \stackrel{d}{=} \mathbf{X}, \quad (5)$$

where $\stackrel{d}{=}$ denotes being equal in distribution. The representation in (5) entails that for each $1 \leq k \leq K$, the i th and j th components of the k th population eigenvector $\mathbf{v}_k$ are identical;

that is,

$$\mathbf{v}_k(i) = \mathbf{v}_k(j).$$

This identity along with the asymptotic expansion of the empirical eigenvector $\hat{\mathbf{v}}_k$ in (B.25) given in Lemma 6 results in

$$\hat{\mathbf{v}}_k(i) - \hat{\mathbf{v}}_k(j) = \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k}{t_k} + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n}|d_k|^2} + \frac{1}{\sqrt{n}|d_k|} \right). \quad (6)$$

Note that although the expectation of $\mathbf{e}_i^T \mathbf{W} \mathbf{v}_k$ can be nonzero, the difference of expectations $\mathbb{E}(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k = 0$ under the null hypothesis by (5). It follows from Lemma 4 in Section B.6 and Lemma 12 in Section C.6 of Supplementary Material that

$$n^{1-c_2\theta} \lesssim d_k \sim t_k \lesssim n\theta \quad \text{and} \quad \alpha_n = O(\sqrt{n\theta}),$$

where $\sim$ denotes the same asymptotic order. Condition 3 ensures that there exists some positive constant ϵ such that

$$\text{SD}((\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k) \sim \sqrt{\theta} \gg n^\epsilon n^{c_2-1/2} \gtrsim n^\epsilon \left(\frac{\alpha_n^2}{\sqrt{n}|d_k|} + \frac{1}{\sqrt{n}} \right), \quad (7)$$

which guarantees that $O_{\prec}(\frac{\alpha_n^2}{\sqrt{n}|d_k|^2} + \frac{1}{\sqrt{n}|d_k|})$ in (6) is negligible compared to the first term on the right hand side. Here SD represents the standard deviation of a random variable. Moreover, by Lemma 3 in Section B.5 of Supplementary Material we have $\|\mathbf{V}\|_\infty = O(\frac{1}{\sqrt{n}}) \ll \min_{1 \leq k \leq K} \text{SD}((\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k) \sim \sqrt{\theta}$, and hence $((\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_k)_{k=1}^K$ satisfies the conditions of Lemma 1 in Section B.3 of Supplementary Material with $h_n = \theta$. Then it holds that

$$\begin{aligned} & \Sigma_1^{-1/2}(\hat{\mathbf{V}}(i) - \hat{\mathbf{V}}(j)) \\ &= \Sigma_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T + o_p(1) \xrightarrow{\mathcal{D}} N(\mathbf{0}, \mathbf{I}), \end{aligned} \quad (8)$$

which proves (12).

We next establish (13) under the condition of $\sqrt{n^{1-2c_2\theta}} \|\boldsymbol{\pi}_i - \boldsymbol{\pi}_j\| \rightarrow \infty$. By (B.25) in Lemma 6, we have

$$\begin{aligned} & \mathbf{D}(\hat{\mathbf{V}}(i) - \hat{\mathbf{V}}(j)) \\ &= \mathbf{D}(\mathbf{V}(i) - \mathbf{V}(j)) + \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n}|d_K|} \right). \end{aligned} \quad (9)$$

In view of (7), it holds that

$$\left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right) = O_p(\sqrt{\theta}).$$

Thus it suffices to show that

$$\|\mathbf{D}(\mathbf{V}(i) - \mathbf{V}(j))\| \gg \sqrt{\theta}.$$

In fact, it follows from (B.17) that

$$\mathbf{D}(\mathbf{V}(i) - \mathbf{V}(j)) = \mathbf{D}\mathbf{B}(\boldsymbol{\pi}_i - \boldsymbol{\pi}_j).$$

This along with (B.18) and Condition 2 leads to

$$\begin{aligned} \|\mathbf{D}(\mathbf{V}(i) - \mathbf{V}(j))\| &= \|\mathbf{D}(\boldsymbol{\pi}_i - \boldsymbol{\pi}_j)^T \mathbf{B}\| \geq |d_K| \sqrt{(\boldsymbol{\pi}_i - \boldsymbol{\pi}_j)^T (\boldsymbol{\Pi}^T \boldsymbol{\Pi})^{-1} (\boldsymbol{\pi}_i - \boldsymbol{\pi}_j)} \\ &\gtrsim \|\pi_1 - \pi_2\| n^{1/2 - c_2} \theta \gg \sqrt{\theta}, \end{aligned}$$

which concludes the proof of (13).

Finally, we prove (14). The conclusion follows immediately from (9) and $(\mathbf{V}(i) - \mathbf{V}(j))^T \boldsymbol{\Sigma}_1^{-1} (\mathbf{V}(i) - \mathbf{V}(j)) \rightarrow \mu$ as $n \rightarrow \infty$. This completes the proof of Theorem 1.

A.2 Proof of Theorem 2

As guaranteed by Slutsky's lemma, the asymptotic distributions of test statistics after replacing $\boldsymbol{\Sigma}_1$ with $\hat{\mathbf{S}}_1$ stay the same. Thus we need only to prove that the asymptotic distributions are the same after replacing K with its estimate $\hat{K}$ in the test statistics.

To ease the presentation, we write $T_{ij} = T_{ij}(K)$ and $\hat{T}_{ij} = T_{ij}(\hat{K})$ to emphasize their dependency on K and $\hat{K}$, respectively. By (12) of Theorem 1, we have for any $t > 0$,

$$\lim_{n \rightarrow \infty} \mathbb{P}(T_{ij}(K) < t) = \mathbb{P}(\chi_K^2 < t). \quad (10)$$

By the condition on $\hat{K}$, it holds that

$$\mathbb{P}(\hat{K} = K) = 1 - o(1). \quad (11)$$

Then by the properties of conditional probability, we deduce

$$\begin{aligned} \mathbb{P}(T_{ij}(\hat{K}) < t) &= \mathbb{P}(T_{ij}(\hat{K}) < t | \hat{K} = K) \mathbb{P}(\hat{K} = K) + \mathbb{P}(T_{ij}(\hat{K}) < t | \hat{K} \neq K) \mathbb{P}(\hat{K} \neq K) \\ &= \mathbb{P}(T_{ij}(K) < t | \hat{K} = K) \mathbb{P}(\hat{K} = K) + o(1) \\ &= \mathbb{P}(T_{ij}(K) < t | \hat{K} = K) \mathbb{P}(\hat{K} = K) + \mathbb{P}(T_{ij}(K) < t | \hat{K} \neq K) \mathbb{P}(\hat{K} \neq K) + o(1) \\ &= \mathbb{P}(T_{ij}(K) < t) + o(1). \end{aligned} \quad (12)$$

Observe that the $o(1)$ term comes from (11) and thus it holds uniformly for any t . Combining

(12) with (10), we can show that

$$\lim_{n \rightarrow \infty} \mathbb{P}(T_{ij}(\hat{K}) < t) = \mathbb{P}(\chi_K^2 < t). \quad (13)$$

Therefore, the same conclusion as in (12) of Theorem 1 is proved. Results in (13) and (14) can be shown using similar arguments and are omitted here for simplicity. This concludes the proof of Theorem 2.

A.3 Proof of Corollary 2

Recall that in the proof of Theorem 2, we denote by $T_{ij} = T_{ij}(K)$ and $\hat{T}_{ij} = T_{ij}(\hat{K})$ to emphasize their dependency on K and $\hat{K}$. It suffices to prove that the impact of the use of $\hat{K}$ in place of K is asymptotically negligible. In fact, we can deduce that

$$\begin{aligned} \mathbb{P}(T_{ij}(\hat{K}) > \chi_{\hat{K}, 1-\alpha}^2) &= \mathbb{P}(T_{ij}(\hat{K}) > \chi_{\hat{K}, 1-\alpha}^2 | \hat{K} = K) \mathbb{P}(\hat{K} = K) \\ &\quad + \mathbb{P}(T_{ij}(\hat{K}) > \chi_{\hat{K}, 1-\alpha}^2 | \hat{K} \neq K) \mathbb{P}(\hat{K} \neq K) \\ &= \mathbb{P}(T_{ij}(K) > \chi_{K, 1-\alpha}^2 | \hat{K} = K) \mathbb{P}(\hat{K} = K) + o(1) \\ &= \mathbb{P}(T_{ij}(K) > \chi_{K, 1-\alpha}^2 | \hat{K} = K) \mathbb{P}(\hat{K} = K) \\ &\quad + \mathbb{P}(T_{ij}(K) > \chi_{K, 1-\alpha}^2 | \hat{K} \neq K) \mathbb{P}(\hat{K} \neq K) + o(1) \\ &= \mathbb{P}(T_{ij}(K) > \chi_{K, 1-\alpha}^2) + o(1). \end{aligned} \quad (14)$$

By (14), under the null hypothesis we have

$$\lim_{n \rightarrow \infty} \mathbb{P}(\hat{T}_{ij} > \chi_{\hat{K}, 1-\alpha}^2) = \lim_{n \rightarrow \infty} \mathbb{P}(T_{ij} > \chi_{K, 1-\alpha}^2) = \alpha \quad (15)$$

for any constant $\alpha \in (0, 1)$. Moreover, by (12), under the alternative hypothesis, for any arbitrarily large constant $C > 0$ it holds that

$$\lim_{n \rightarrow \infty} \mathbb{P}(\hat{T}_{ij} > C) = \lim_{n \rightarrow \infty} \mathbb{P}(T_{ij} > C) = 1. \quad (16)$$

Therefore, combining (15) and (16) completes the proof of Corollary 2.

A.4 Proof of Theorem 3

We begin with listing some basic properties of $\mathbf{v}_k$ and d_k :

- 1). We can choose a direction such that all components of $\mathbf{v}_1$ are nonnegative. Moreover, $\min_{1 \leq l \leq n} \{\mathbf{v}_1(l)\} \sim \frac{1}{\sqrt{n}}$.
- 2). $\max_{1 \leq k \leq K} \|\mathbf{v}_k\|_\infty \leq \frac{C}{\sqrt{n}}$ for some positive constant C .
- 3). $\alpha_n \leq \sqrt{n} \theta_{\max}$.

4). $|d_K| \geq cn^{1-2c_2}\theta_{\min}^2$ and $|d_1| \leq c^{-1}n\theta_{\max}^2$ for some positive constant c .

Here the second statement is ensured by Lemma 3. The third and fourth statements are guaranteed by Lemma 4, and the remaining properties are entailed by Lemma B.2 of ?. One should notice that the proof of Lemma B.2 of ? does not require $\{d_k\}_{k=1}^K$ have the same order.

By Condition 5 and Statement 4 above, we have

$$\frac{1}{n^{1/2-c_2}|t_k|} \ll \frac{\min_{1 \leq k \leq K, t=i,j} \sqrt{\text{var}(\mathbf{e}_t^T \mathbf{W} \mathbf{v}_k)}}{|t_k|}.$$

By (B.19), there exists some $K \times K$ matrix $\mathbf{B}$ such that

$$\mathbf{V} = \mathbf{\Theta} \mathbf{\Pi} \mathbf{B}. \quad (17)$$

Recall that $\mathbf{\Theta}$ is a diagonal matrix. Then it follows from (17) that under the null hypothesis, we have

$$\frac{\mathbf{v}_k(i)}{\mathbf{v}_1(i)} = \frac{\mathbf{v}_k(j)}{\mathbf{v}_1(j)}, \quad k = 1, \dots, K. \quad (18)$$

Here we use the exchangeability between rows i and j of matrix $\mathbf{\Pi} \mathbf{B}$ under the null hypothesis as argued under the mixed membership model (see the beginning of the proof of Theorem 1). In light of the asymptotic expansion in Lemma 6, we deduce

$$\widehat{\mathbf{v}}_k(i) = \mathbf{v}_k(i) + \frac{\mathbf{e}_i^T \mathbf{W} \mathbf{v}_k}{t_k} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}|t_k|}\right). \quad (19)$$

Moreover, it follows from Corollary 1 in Section C.2 of Supplementary Material, Condition 4, and the statements at the beginning of this proof that

$$\frac{\mathbf{e}_s^T \mathbf{W} \mathbf{v}_k}{t_k} = O_{\prec}\left(\frac{\theta_{\max}}{|t_k|}\right), \quad s = i, j, \quad k = 1, \dots, K. \quad (20)$$

Thus, by (18)–(20) and Statement 1 above we have under the null hypothesis that

$$\begin{aligned} \mathbf{Y}(i, k) - \mathbf{Y}(j, k) &= \frac{\widehat{\mathbf{v}}_k(i)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{\mathbf{v}}_k(j)}{\widehat{\mathbf{v}}_1(j)} \\ &= \frac{\mathbf{v}_k(i) + \frac{\mathbf{e}_i^T \mathbf{W} \mathbf{v}_k}{t_k} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}|t_k|}\right)}{\mathbf{v}_1(i) + \frac{\mathbf{e}_i^T \mathbf{W} \mathbf{v}_1}{t_1} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}|t_1|}\right)} - \frac{\mathbf{v}_k(j) + \frac{\mathbf{e}_j^T \mathbf{W} \mathbf{v}_k}{t_k} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}|t_k|}\right)}{\mathbf{v}_1(j) + \frac{\mathbf{e}_j^T \mathbf{W} \mathbf{v}_1}{t_1} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}|t_1|}\right)} \\ &= \frac{\mathbf{e}_i^T \mathbf{W} \mathbf{v}_k}{t_k \mathbf{v}_1(i)} - \frac{\mathbf{e}_j^T \mathbf{W} \mathbf{v}_k}{t_k \mathbf{v}_1(j)} - \frac{\mathbf{v}_k(i) \mathbf{e}_i^T \mathbf{W} \mathbf{v}_1}{t_1 \mathbf{v}_1^2(i)} + \frac{\mathbf{v}_k(j) \mathbf{e}_j^T \mathbf{W} \mathbf{v}_1}{t_1 \mathbf{v}_1^2(j)} + O_{\prec}\left(\frac{n^{c_2}}{|t_k|}\right) \\ &= \frac{\mathbf{e}_i^T \mathbf{W} [\mathbf{v}_k - \frac{t_k \mathbf{v}_k(i)}{t_1 \mathbf{v}_1(i)} \mathbf{v}_1]}{t_k \mathbf{v}_1(i)} - \frac{\mathbf{e}_j^T \mathbf{W} [\mathbf{v}_k - \frac{t_k \mathbf{v}_k(j)}{t_1 \mathbf{v}_1(j)} \mathbf{v}_1]}{t_k \mathbf{v}_1(j)} + O_{\prec}\left(\frac{n^{c_2}}{|t_k|}\right). \end{aligned} \quad (21)$$

Denote by $\mathbf{y}_k = \frac{\mathbf{v}_k - \frac{t_k \mathbf{v}_k(i)}{t_1 \mathbf{v}_1(i)} \mathbf{v}_1}{t_k \mathbf{v}_1(i)}$ and $\mathbf{z}_k = \frac{\mathbf{v}_k - \frac{t_k \mathbf{v}_k(j)}{t_1 \mathbf{v}_1(j)} \mathbf{v}_1}{t_k \mathbf{v}_1(j)}$. Then we have $f_k = \mathbf{e}_i^T \mathbf{W} \mathbf{y}_k - \mathbf{e}_j^T \mathbf{W} \mathbf{z}_k$ with f_k defined in (22), and

$$\mathbf{Y}(i, k) - \mathbf{Y}(j, k) = f_k + O_{\prec}\left(\frac{n^{c_2}}{|t_k|}\right). \quad (22)$$

To establish the central limit theorem, we need to compare the order of the variance of f_k with that of the residual term $O_{\prec}\left(\frac{n^{2c_2}}{t_k^2}\right)$. The variance of f_k is

$$\text{var}(f_k) = \sum_{l=1}^n \text{var}(w_{il}) \mathbf{y}_k^2(l) + \sum_{l=1}^n \text{var}(w_{jl}) \mathbf{z}_k^2(l) - \text{var}(w_{ij}) [\mathbf{y}_k(i) \mathbf{z}_k(j) + \mathbf{y}_k(j) \mathbf{z}_k(i)]. \quad (23)$$

By Statements 1 and 2 at the beginning of this proof and (18), we can conclude that $\max_{1 \leq l \leq n} \{|\mathbf{y}_k(l)|, |\mathbf{z}_k(l)|\} = O\left(\frac{1}{|t_k|}\right)$ and $\mathbf{y}_k(l) \sim \mathbf{z}_k(l)$, $l = 1, \dots, n$. Consequently, we obtain

$$\text{var}(w_{ij}) [\mathbf{y}_k(i) \mathbf{z}_k(j) + \mathbf{y}_k(j) \mathbf{z}_k(i)] = O\left(\frac{1}{t_k^2}\right). \quad (24)$$

By Condition 6, it holds that $(n\theta_{\max}^2)^{-1} d_k^2 \text{var}(f_k) = (n\theta_{\max}^2)^{-1} d_k^2 \text{var}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k - \mathbf{e}_j^T \mathbf{W} \mathbf{z}_k) \sim 1$. Combining the previous two results and by Statement 4, the last term on the left hand side of (23) is asymptotically negligible compared to the right hand side.

Note that under the null hypothesis $\boldsymbol{\pi}_i = \boldsymbol{\pi}_j$ and model (6), we have $\frac{\mathbf{H}_{il}}{\theta_i} = \frac{\mathbf{H}_{jl}}{\theta_j}$. Since $\mathbf{X} = \mathbf{H} + \mathbf{W}$ with $\mathbf{W}$ a generalized Wigner matrix, it follows from the properties of Bernoulli random variables that $\text{var}(w_{il}) \sim \text{var}(w_{jl})$. Thus the first two terms on the left hand side of (23) are comparable and satisfy that

$$\begin{aligned} (n\theta_{\max}^2)^{-1} d_k^2 \text{var}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k) &= (n\theta_{\max}^2)^{-1} d_k^2 \sum_{l=1}^n 2\text{var}(w_{il}) \mathbf{y}_k^2(l) \\ &\sim (n\theta_{\max}^2)^{-1} d_k^2 \sum_{l=1}^n 2\text{var}(w_{jl}) \mathbf{z}_k^2(l) = (n\theta_{\max}^2)^{-1} d_k^2 \text{var}(\mathbf{e}_j^T \mathbf{W} \mathbf{z}_k) \\ &\sim (n\theta_{\max}^2)^{-1} d_k^2 \text{var}(f_k) \sim 1. \end{aligned} \quad (25)$$

Consequently, $\text{var}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k) \sim \text{var}(\mathbf{e}_j^T \mathbf{W} \mathbf{z}_k) \sim \text{var}(f_k)$.

Now we are ready to check the conditions of Lemma 1. By $\max_l \{|\mathbf{y}_k(l)|, |\mathbf{z}_k(l)|\} = O\left(\frac{1}{|t_k|}\right)$ (see (24) above) and noticing that the expectations of the off-diagonal entries of $\mathbf{W}$ are zero, we have $|\mathbb{E}(f_k)| = |\mathbb{E}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k - \mathbf{e}_j^T \mathbf{W} \mathbf{z}_k)| = |\mathbb{E}(w_{ii} \mathbf{y}_k(i) - w_{jj} \mathbf{z}_k(j))| \leq |\mathbf{y}_k(i)| + |\mathbf{z}_k(j)| = O\left(\frac{1}{|t_k|}\right)$, which means that the expectation of $\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k - \mathbf{e}_j^T \mathbf{W} \mathbf{z}_k$ is asymptotically negligible compared to its standard deviation. Moreover, by (25) it holds that $\max_l \{|\mathbf{y}_k(l)|, |\mathbf{z}_k(l)|\} \ll \min_{1 \leq k \leq K} \min\{\text{SD}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_k), \text{SD}(\mathbf{e}_j^T \mathbf{W} \mathbf{z}_k)\}$ and hence they satisfy the conditions of Lemma

1 with $h_n = n\theta_{\max}^2$. Thus we arrive at

$$\text{cov}(\mathbf{e}_i^T \mathbf{W} \mathbf{y}_2, \mathbf{e}_j^T \mathbf{W} \mathbf{z}_2, \dots, \mathbf{e}_j^T \mathbf{W} \mathbf{z}_K)^{-1/2} (\mathbf{e}_i^T \mathbf{W} \mathbf{y}_2, \mathbf{e}_j^T \mathbf{W} \mathbf{z}_2, \dots, \mathbf{e}_j^T \mathbf{W} \mathbf{z}_K)^T \xrightarrow{\mathcal{D}} N(\mathbf{0}, \mathbf{I}). \quad (26)$$

Using the compact notation, (26) can be rewritten as

$$\Sigma_2^{-1/2} (f_2, \dots, f_K)^T \xrightarrow{\mathcal{D}} N(\mathbf{0}, \mathbf{I}). \quad (27)$$

Furthermore, there exists some positive constant ϵ such that $\text{SD}(f_k) \sim \frac{\sqrt{n}\theta_{\max}}{|t_k|} \gg n^\epsilon \frac{n^{c_2}}{|t_k|}$ by Condition 4. Hence $O_{\prec}(\frac{n^{c_2}}{|t_k|})$ involved in (22) is negligible compared to f_k . Finally, we can obtain from (22) and (27) that

$$\Sigma_2^{-1/2} (\mathbf{Y}_i - \mathbf{Y}_j) \xrightarrow{\mathcal{D}} N(\mathbf{0}, \mathbf{I}),$$

which completes the proof for part i) of Theorem 3.

It remains to prove part ii) of Theorem 3. Under the alternative hypothesis that $\pi_i \neq \pi_j$, we have the generalized asymptotic expansion

$$\mathbf{Y}(i, k) - \mathbf{Y}(j, k) = \frac{\mathbf{v}_k(i)}{\mathbf{v}_1(i)} - \frac{\mathbf{v}_k(j)}{\mathbf{v}_1(j)} + \mathbf{e}_i^T \mathbf{W} \mathbf{y}_k - \mathbf{e}_j^T \mathbf{W} \mathbf{z}_k + O_{\prec}(\frac{n^{c_2}}{|t_k|}). \quad (28)$$

In view of (26), to complete the proof it suffices to show that

$$\left\| \frac{\mathbf{V}(i)}{\mathbf{v}_1(i)} - \frac{\mathbf{V}(j)}{\mathbf{v}_1(j)} \right\| \gg \frac{1}{n^{1/2-c_2}\theta_{\min}}. \quad (29)$$

Denote by $\mathbf{B}(i)$ the i th column of matrix $\mathbf{B}$ in (17). It follows from (17) that

$$\frac{\mathbf{V}(i)}{\mathbf{v}_1(i)} = \frac{\pi_i^T \mathbf{B}}{\pi_i^T \mathbf{B}(1)} \quad \text{and} \quad \frac{\mathbf{V}(j)}{\mathbf{v}_1(j)} = \frac{\pi_j^T \mathbf{B}}{\pi_j^T \mathbf{B}(1)}.$$

Let $a_i = \pi_i^T \mathbf{B}(1)$ and $a_j = \pi_j^T \mathbf{B}(1)$. Note that by Statements 1 and 2 at the beginning of this proof, we have $\mathbf{v}_1(i) \sim \mathbf{v}_1(j) \sim \frac{1}{\sqrt{n}}$. In light of (17), it holds that $\mathbf{v}_1(i) = \theta_i a_i$ and $\mathbf{v}_1(j) = \theta_j a_j$. Combining these two results yields

$$a_i \sim a_j \sim \frac{1}{\sqrt{n}\theta_{\min}}.$$

Moreover, it holds that

$$\frac{\pi_i^T \mathbf{B}}{\pi_i^T \mathbf{B}(1)} - \frac{\pi_j^T \mathbf{B}}{\pi_j^T \mathbf{B}(1)} = (a_i^{-1}, -a_j^{-1})(\pi_i, \pi_j)^T \mathbf{B},$$

which entails that

$$\left\| \frac{\mathbf{V}(i)}{\mathbf{v}_1(i)} - \frac{\mathbf{V}(j)}{\mathbf{v}_1(j)} \right\|^2 \geq \|(a_i^{-1}, -a_j^{-1})\|^2 \lambda_{\min}((\boldsymbol{\pi}_i, \boldsymbol{\pi}_j)^T (\boldsymbol{\pi}_i, \boldsymbol{\pi}_j)) \lambda_{\min}(\mathbf{B}\mathbf{B}^T).$$

Here $\lambda_{\min}(\cdot)$ stands for the smallest eigenvalue. By (17), similar to (B.18) we can show that

$$\mathbf{B}\mathbf{B}^T = (\boldsymbol{\Pi}^T \boldsymbol{\Theta}^2 \boldsymbol{\Pi})^{-1}.$$

Thus $\lambda_{\min}(\mathbf{B}\mathbf{B}^T) \sim \frac{1}{n\theta_{\min}^2}$. By the condition that $\lambda_2(\boldsymbol{\pi}_i \boldsymbol{\pi}_i^T + \boldsymbol{\pi}_j \boldsymbol{\pi}_j^T) \gg \frac{1}{n^{1-2c_2}\theta_{\min}^2}$ in Theorem 3, it holds that

$$\lambda_{\min}((\boldsymbol{\pi}_i, \boldsymbol{\pi}_j)^T (\boldsymbol{\pi}_i, \boldsymbol{\pi}_j)) = \lambda_2(\boldsymbol{\pi}_i \boldsymbol{\pi}_i^T + \boldsymbol{\pi}_j \boldsymbol{\pi}_j^T) \gg \frac{1}{n^{1-2c_2}\theta_{\min}^2}.$$

Therefore, combining the above arguments results in

$$\left\| \frac{\mathbf{V}(i)}{\mathbf{v}_1(i)} - \frac{\mathbf{V}(j)}{\mathbf{v}_1(j)} \right\|^2 \gg \frac{1}{n^{1-2c_2}\theta_{\min}^2},$$

which concludes the proof of Theorem 3.

A.5 Proof of Theorem 4

The arguments for the proof of Theorem 4 are similar to those for the proof of Theorem 2 in Section A.2.

A.6 Proof of Theorem 5

By Lemma 1, (16) holds. Since $\hat{K}$ is bounded with probability tending to 1, it suffices to show the entrywise convergence of $\hat{\boldsymbol{\Sigma}}_1 = \theta^{-1} \mathbf{D} \hat{\mathbf{S}}_1 \mathbf{D}$ and $\hat{\boldsymbol{\Sigma}}_2 = (n\theta_{\max})^{-1} \mathbf{D} \hat{\mathbf{S}}_2 \mathbf{D}$. As will be made clear later, the proof relies heavily on the asymptotic expansions of $(\hat{\boldsymbol{\Sigma}}_1)_{11}$, $(\hat{\boldsymbol{\Sigma}}_1)_{12}$, $(\hat{\boldsymbol{\Sigma}}_2)_{11}$, and $(\hat{\boldsymbol{\Sigma}}_2)_{12}$. We will provide only the full details on the convergence of $(\hat{\boldsymbol{\Sigma}}_1)_{11}$. For the other cases, the asymptotic expansions will be provided and the technical details will be mostly omitted since the arguments of the proof are similar. Throughout the proof, we will use repeatedly the results in Lemma 6, and the node indices i and j are fixed.

We start with considering $(\hat{\boldsymbol{\Sigma}}_1)_{11}$. First, by definitions of $\hat{\mathbf{W}}$ we have the following expansions

$$(\theta^{-1} \mathbf{D} \boldsymbol{\Sigma}_1 \mathbf{D})_{11} = \theta^{-1} \sum_{t=i,j, 1 \leq l \leq n} \left[\sigma_{tl}^2 \mathbf{v}_1^2(l) - 2\sigma_{ij}^2 \mathbf{v}_1(j) \mathbf{v}_1(i) \right] \quad (30)$$

and

$$(\widehat{\boldsymbol{\Sigma}}_1)_{11} = (\theta^{-1} \mathbf{D} \widehat{\mathbf{S}}_1 \mathbf{D})_{11} = \theta^{-1} \sum_{t=i,j, 1 \leq l \leq n} \left[\widehat{w}_{il}^2 \mathbf{v}_1^2(l) - 2\widehat{w}_{ij}^2 \widehat{\mathbf{v}}_1(l) \widehat{\mathbf{v}}_1(i) \right]. \quad (31)$$

It follows from Lemma 7 in Section B.9 of Supplementary Material that $\widehat{w}_{ij}^2 \widehat{\mathbf{v}}_1(j) \widehat{\mathbf{v}}_1(i) = O_{\prec}(\frac{1}{n})$. In addition, by Lemmas 3 and 4 it holds that

$$\text{var} \left[\sum_{1 \leq l \leq n} (w_{il}^2 - \sigma_{il}^2) \mathbf{v}_1^2(l) \right] \leq \sum_{1 \leq l \leq n} \mathbf{v}_1^4(l) \mathbb{E} w_{il}^2 = O\left(\frac{1}{n^2}\right)(\alpha_n^2 + 1) = O\left(\frac{\theta}{n}\right). \quad (32)$$

The same inequality also holds for $\text{var}[\sum_{1 \leq l \leq n} (w_{jl}^2 - \sigma_{jl}^2) \mathbf{v}_1^2(l)]$. Thus we have

$$\sum_{t=i,j, 1 \leq l \leq n} (w_{il}^2 - \sigma_{il}^2) \mathbf{v}_1^2(l) = O_p\left(\frac{\sqrt{\theta}}{\sqrt{n}}\right), \quad (33)$$

which implies that

$$\sum_{t=i,j, 1 \leq l \leq n} w_{il}^2 \mathbf{v}_1^2(l) = \sum_{t=i,j, 1 \leq l \leq n} \sigma_{il}^2 \mathbf{v}_1^2(l) + O_p\left(\frac{\sqrt{\theta}}{\sqrt{n}}\right). \quad (34)$$

By Lemmas 4 and 6, we have

$$\widehat{\mathbf{v}}_k(j) = \mathbf{v}_k(j) + \frac{\mathbf{e}_j^T \mathbf{W} \mathbf{v}_k}{t_k} + O_{\prec}\left(\frac{1}{n^{3/2-2c_2}\theta}\right).$$

It follows from Corollary 1 in Section C.2 and Lemma 10 in Section C.4 of Supplementary Material that

$$\begin{aligned} \sum_{t=i,j, 1 \leq l \leq n} w_{il}^2 [\mathbf{v}_1^2(l) - \widehat{\mathbf{v}}_1^2(l)] &= 2 \sum_{t=i,j, 1 \leq l \leq n} w_{il}^2 \mathbf{v}_1(j) [\mathbf{v}_1(l) - \widehat{\mathbf{v}}_1(l)] + O_{\prec}(n^{2c_2-1}) \\ &= -\frac{2}{t_1} \sum_{t=i,j, 1 \leq l \leq n} w_{il}^2 \mathbf{v}_1(l) \mathbf{e}_l^T \mathbf{W} \mathbf{v}_1 + O_{\prec}\left(\frac{1}{n^{1-2c_2}}\right) \\ &= O_{\prec}\left(\frac{\sqrt{\theta}}{n^{1/2-c_2}}\right). \end{aligned} \quad (35)$$

Similarly, by Lemma 7 we have

$$\sum_{t=i,j, 1 \leq l \leq n} w_{il}^2 \widehat{\mathbf{v}}_1^2(l) = \sum_{t=i,j, 1 \leq l \leq n} \widehat{w}_{il}^2 \widehat{\mathbf{v}}_k^2(l) + O_{\prec}\left(\frac{\sqrt{\theta}}{n^{1/2-c_2}}\right). \quad (36)$$

Combining the equalities (30)–(36) yields

$$(\widehat{\boldsymbol{\Sigma}}_1)_{11} = \theta^{-1} (\mathbf{D} \boldsymbol{\Sigma}_1 \mathbf{D})_{11} + O_{\prec}\left(\frac{1}{n^{1/2-c_2}\sqrt{\theta}}\right) + O_p\left(\frac{1}{\sqrt{n\theta}}\right) = \theta^{-1} (\boldsymbol{\Sigma}_1)_{11} + o_p(1), \quad (37)$$

where we have used $O_{\prec}(\frac{1}{n^{1/2-c_2}\sqrt{\theta}}) = o_p(1)$ by Condition 2. This has proved the convergence of $(\widehat{\Sigma}_1)_{11}$ to $(\Sigma_1)_{11}$.

We next consider $(\widehat{\Sigma}_1)_{12}$. By definitions, we have the following expansions

$$(\theta^{-1}\mathbf{D}\Sigma_1\mathbf{D})_{12} = \theta^{-1}\left\{\sum_{t=i,j}\sigma_{ti}^2\mathbf{v}_1(l)\mathbf{v}_2(l) - \sigma_{ij}^2[\mathbf{v}_1(j)\mathbf{v}_2(i) + \mathbf{v}_1(i)\mathbf{v}_2(j)]\right\} \quad (38)$$

and

$$(\widehat{\Sigma}_1)_{12} = \theta^{-1}\left\{\sum_{t=i,j}\widehat{w}_{ti}^2\widehat{\mathbf{v}}_1(l)\widehat{\mathbf{v}}_2(l) - \widehat{w}_{ij}^2[\widehat{\mathbf{v}}_1(j)\widehat{\mathbf{v}}_2(i) + \widehat{\mathbf{v}}_1(i)\widehat{\mathbf{v}}_2(j)]\right\}. \quad (39)$$

Based on the above two expansions, using similar arguments to those for proving (37) we can show that

$$(\widehat{\Sigma}_1)_{12} = \theta^{-1}(\mathbf{D}\Sigma_1\mathbf{D})_{12} + o_p(1). \quad (40)$$

Now let us consider $\widehat{\Sigma}_2$. Similar as above, we will provide only the asymptotic expansions for $(\widehat{\Sigma}_2)_{11}$ and $(\widehat{\Sigma}_2)_{12}$, and the remaining arguments are similar. By definitions, we can deduce that

$$\begin{aligned} ((n\theta_{\max}^2)^{-1}\mathbf{D}\Sigma_2\mathbf{D})_{11} &= (n\theta_{\max}^2)^{-1}d_2^2\text{var}(f_2) \\ &= \frac{d_2^2}{t_2^2n\theta_{\max}^2}\left\{\sum_{l \neq j}\sigma_{il}^2\left[\frac{\mathbf{v}_2(l)}{\mathbf{v}_1(i)} - \frac{t_2\mathbf{v}_2(i)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(i)^2}\right]^2 + \sum_{l \neq i}\sigma_{jl}^2\left[\frac{\mathbf{v}_2(l)}{\mathbf{v}_1(j)} - \frac{t_2\mathbf{v}_2(j)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(j)^2}\right]^2\right. \\ &\quad \left.+ \sigma_{ij}^2\left[\frac{\mathbf{v}_2(j)}{\mathbf{v}_1(i)} - \frac{t_2\mathbf{v}_2(i)\mathbf{v}_1(j)}{t_1\mathbf{v}_1(i)^2} - \frac{\mathbf{v}_2(i)}{\mathbf{v}_1(j)} + \frac{t_2\mathbf{v}_2(j)\mathbf{v}_1(i)}{t_1\mathbf{v}_1(j)^2}\right]^2\right\} \end{aligned}$$

and

$$\begin{aligned} (\widehat{\Sigma}_2)_{11} &= \frac{d_2^2}{\widehat{d}_2^2n\theta_{\max}^2}\left\{\sum_{l \neq j}\widehat{w}_{il}^2\left[\frac{\widehat{\mathbf{v}}_2(l)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(i)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2}\right]^2 + \sum_{l \neq i}\widehat{w}_{jl}^2\left[\frac{\widehat{\mathbf{v}}_2(l)}{\widehat{\mathbf{v}}_1(j)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(j)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2}\right]^2\right. \\ &\quad \left.+ \widehat{w}_{ij}^2\left[\frac{\widehat{\mathbf{v}}_2(j)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(i)\widehat{\mathbf{v}}_1(j)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2} - \frac{\widehat{\mathbf{v}}_2(i)}{\widehat{\mathbf{v}}_1(j)} + \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(j)\widehat{\mathbf{v}}_1(i)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2}\right]^2\right\}. \end{aligned}$$

Note that the expression of $(n\theta_{\max}^2\mathbf{D}\Sigma_2\mathbf{D})_{11}$ is essentially the same as (30) up to a normalization factor involving $\mathbf{v}_1(i)$ and $\mathbf{v}_1(j)$. Thus applying the similar arguments to those for proving (17), we can establish the desired result.

Finally, the consistency of $(\widehat{\Sigma}_2)_{12}$ can also be shown similarly using the following expan-

sions

$$\begin{aligned}
& ((n\theta_{\max}^2)^{-1}\mathbf{D}\Sigma_2\mathbf{D})_{12} \\
&= \frac{d_2d_3}{t_2t_3n\theta_{\max}^2} \left\{ \sum_{l \neq j} \sigma_{il}^2 \left[\frac{\mathbf{v}_2(l)}{\mathbf{v}_1(i)} - \frac{t_2\mathbf{v}_2(i)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(i)^2} \right] \left[\frac{\mathbf{v}_3(l)}{\mathbf{v}_1(i)} - \frac{t_3\mathbf{v}_3(i)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(i)^2} \right] \right. \\
&\quad + \sum_{l \neq i} \sigma_{jl}^2 \left[\frac{\mathbf{v}_2(l)}{\mathbf{v}_1(j)} - \frac{t_2\mathbf{v}_2(j)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(j)^2} \right] \left[\frac{\mathbf{v}_2(l)}{\mathbf{v}_1(j)} - \frac{t_3\mathbf{v}_3(j)\mathbf{v}_1(l)}{t_1\mathbf{v}_1(j)^2} \right] \\
&\quad + \sigma_{ij}^2 \left[\frac{\mathbf{v}_2(j)}{\mathbf{v}_1(i)} - \frac{t_2\mathbf{v}_2(i)\mathbf{v}_1(j)}{t_1\mathbf{v}_1(i)^2} - \frac{\mathbf{v}_2(i)}{\mathbf{v}_1(j)} + \frac{t_2\mathbf{v}_2(j)\mathbf{v}_1(i)}{t_1\mathbf{v}_1(j)^2} \right] \\
&\quad \times \left. \left[\frac{\mathbf{v}_3(j)}{\mathbf{v}_1(i)} - \frac{t_3\mathbf{v}_3(i)\mathbf{v}_1(j)}{t_1\mathbf{v}_1(i)^2} - \frac{\mathbf{v}_3(i)}{\mathbf{v}_1(j)} + \frac{t_3\mathbf{v}_3(j)\mathbf{v}_1(i)}{t_1\mathbf{v}_1(j)^2} \right] \right\}
\end{aligned}$$

and

$$\begin{aligned}
(\widehat{\Sigma}_2)_{12} &= \frac{d_2d_3}{\widehat{d}_2\widehat{d}_3n\theta_{\max}^2} \left\{ \sum_{l \neq j} \widehat{w}_{il}^2 \left[\frac{\widehat{\mathbf{v}}_2(l)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(i)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2} \right] \left[\frac{\widehat{\mathbf{v}}_3(l)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_3\widehat{\mathbf{v}}_3(i)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2} \right] \right. \\
&\quad + \sum_{l \neq i} \widehat{w}_{jl}^2 \left[\frac{\widehat{\mathbf{v}}_2(l)}{\widehat{\mathbf{v}}_1(j)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(j)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2} \right] \left[\frac{\widehat{\mathbf{v}}_3(l)}{\widehat{\mathbf{v}}_1(j)} - \frac{\widehat{d}_3\widehat{\mathbf{v}}_3(j)\widehat{\mathbf{v}}_1(l)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2} \right] \\
&\quad + \widehat{w}_{ij}^2 \left[\frac{\widehat{\mathbf{v}}_2(j)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(i)\widehat{\mathbf{v}}_1(j)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2} - \frac{\widehat{\mathbf{v}}_2(i)}{\widehat{\mathbf{v}}_1(j)} + \frac{\widehat{d}_2\widehat{\mathbf{v}}_2(j)\widehat{\mathbf{v}}_1(i)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2} \right] \\
&\quad \times \left. \left[\frac{\widehat{\mathbf{v}}_3(j)}{\widehat{\mathbf{v}}_1(i)} - \frac{\widehat{d}_3\widehat{\mathbf{v}}_3(i)\widehat{\mathbf{v}}_1(j)}{\widehat{d}_1\widehat{\mathbf{v}}_1(i)^2} - \frac{\widehat{\mathbf{v}}_3(i)}{\widehat{\mathbf{v}}_1(j)} + \frac{\widehat{d}_3\widehat{\mathbf{v}}_3(j)\widehat{\mathbf{v}}_1(i)}{\widehat{d}_1\widehat{\mathbf{v}}_1(j)^2} \right] \right\}.
\end{aligned}$$

This completes the proof of Theorem 5.

B Some key lemmas and their proofs

B.1 Proof of Lemma 1

For each pair (i, j) with $i \neq j$, let us define a matrix $\mathbf{W}(i, j) = w_{ij}(\mathbf{e}_i\mathbf{e}_j^T + \mathbf{e}_j\mathbf{e}_i^T)$. For $i = j$, we define a matrix $\mathbf{W}(i, j) = (w_{ii} - \mathbb{E}w_{ii})\mathbf{e}_i\mathbf{e}_i^T$. Then it is easy to see that

$$\left\| \sum_{1 \leq i \leq j \leq n} \mathbf{W}(i, j) - \mathbf{W} \right\| = \|\text{diag}(\mathbf{W} - \mathbb{E}\mathbf{W})\| \leq 1. \quad (\text{B.1})$$

It is straightforward to show that

$$\left\| \sum_{1 \leq i \leq j \leq n} \mathbb{E}\mathbf{W}(i, j)^2 \right\| = \alpha_n^2.$$

By Theorem 6.2 of ?, for any constant $c > \sqrt{2}$ we have

$$\mathbb{P}(\|\sum_{1 \leq i \leq j \leq n} \mathbf{W}(i, j)\| \geq c\sqrt{\log n \alpha_n} - 1) \leq n \exp \left[\frac{-(c\sqrt{\log n \alpha_n} - 1)^2}{2\alpha_n^2 + 2(c\sqrt{\log n \alpha_n} - 1)} \right] = o(1). \quad (\text{B.2})$$

This together with (B.1) entails that

$$\mathbb{P}(\|\mathbf{W}\| \leq c\sqrt{\log n \alpha_n}) \geq 1 - o(1). \quad (\text{B.3})$$

Note that this result is weaker than Lemma 11 in Section C.5.

By (B.3) and $|\hat{d}_K - d_K| \leq \|\mathbf{W}\|$, and using the assumption of $|d_K| \gg \sqrt{\log n \alpha_n}$, it holds that

$$|\hat{d}_K| \gg \sqrt{\log n \alpha_n} \quad (\text{B.4})$$

with probability tending to one. Finally, by Weyl's inequality we have

$$\lambda_n(\mathbf{W}) = \lambda_n(\mathbf{W}) - \lambda_{K+1}(\mathbf{H}) \leq \lambda_{K+1}(\mathbf{X}) = \lambda_{K+1}(\mathbf{H} + \mathbf{W}) \leq \lambda_1(\mathbf{W}) + \lambda_{K+1}(\mathbf{H}) = \lambda_1(\mathbf{W}),$$

which leads to

$$|\hat{d}_{K+1}| = |\lambda_{K+1}(\mathbf{X})| \leq \|\mathbf{W}\|. \quad (\text{B.5})$$

Let us choose $c = \sqrt{2.01}$ and define

$$\tilde{K} = \# \left\{ |\hat{d}_i| > \sqrt{2.01 \log n \alpha_n}, i = 1, \dots, n \right\}. \quad (\text{B.6})$$

Then by (B.4)–(B.5), we can show that

$$\mathbb{P}(\tilde{K} = K) = 1 - o(1). \quad (\text{B.7})$$

Recall that X_{ij} follows the Bernoulli distribution. Thus it holds that

$$\sum_{j=1} \mathbb{E} w_{ij}^2 \leq \sum_{j=1} \mathbb{E} X_{ij}.$$

By Lemma 8 in Section C.1, choosing $l = 1$, $\mathbf{x} = \mathbf{e}_i$, and $\mathbf{y} = \frac{1}{\sqrt{n}} \mathbf{1}$ yields

$$\sum_{j=1} \mathbb{E} X_{ij} = \sum_{j=1} X_{ij} + O_{\prec}(\alpha_n),$$

where we have used $X_{ij} - \mathbb{E} X_{ij} = w_{ij}$. Thus it holds that

$$\max_i \sum_{j=1} X_{ij} \geq \max_i \sum_{j=1} \mathbb{E} w_{ij}^2 + O_{\prec}(\alpha_n) = \alpha_n^2 + O_{\prec}(\alpha_n).$$

This together with (B.6) and (B.7) results in

$$\mathbb{P}(\widehat{K} = K) = 1 - o(1), \quad (\text{B.8})$$

which completes the proof of Lemma 1.

B.2 Proofs of Lemmas 2 and 3

The proofs of Lemmas 2 and 3 involve standard calculations and thus are omitted for brevity.

B.3 Lemma 1 and its proof

Lemma 1. *Let m be a fixed positive integer, $\mathbf{x}_i$ and $\mathbf{y}_i$ be n -dimensional unit vectors for $1 \leq i \leq m$, and $\Sigma = (\Sigma_{ij})$ the covariance matrix with $\Sigma_{ij} = \text{cov}(\mathbf{x}_i^T \mathbf{W} \mathbf{y}_i, \mathbf{x}_j^T \mathbf{W} \mathbf{y}_j)$. Assume that there exists some positive sequence (h_n) such that $\|\Sigma^{-1}\| \sim \|\Sigma\| \sim h_n$ and $\max_k \{\|\mathbf{x}_k\|_\infty \|\mathbf{y}_k\|_\infty\} \ll \|\Sigma^{1/2}\|$. Then it holds that*

$$\Sigma^{-1/2} (\mathbf{x}_1^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_1, \dots, \mathbf{x}_m^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_m)^T \xrightarrow{\mathcal{D}} N(\mathbf{0}, \mathbf{I}). \quad (\text{B.9})$$

Proof. Note that it suffices to show that for any unit vector $\mathbf{c} = (c_1, \dots, c_m)^T$, we have

$$\mathbf{c}^T \Sigma^{-1/2} (\mathbf{x}_1^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_1, \dots, \mathbf{x}_m^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_m)^T \xrightarrow{\mathcal{D}} N(0, 1). \quad (\text{B.10})$$

Let $\mathbf{x}_i = (x_{1i}, \dots, x_{ni})^T$ and $\mathbf{y}_i = (y_{1i}, \dots, y_{ni})^T$, $i = 1, \dots, m$. Since $\mathbf{W}$ is a symmetric random matrix of independent entries on and above the diagonal, we can deduce

$$\mathbf{x}_i^T \mathbf{W} \mathbf{y}_i - \mathbf{x}_i^T \mathbb{E} \mathbf{W} \mathbf{y}_i = \sum_{1 \leq s, t \leq n, s < t} w_{st} (x_{si} y_{ti} + x_{ti} y_{si}) + \sum_{1 \leq s \leq n} (w_{ss} - \mathbb{E} w_{ss}) x_{si} y_{si} \quad (\text{B.11})$$

and

$$\begin{aligned} s_n^2 &:= \text{var} \left[\mathbf{c}^T \Sigma^{-1/2} (\mathbf{x}_1^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_1, \dots, \mathbf{x}_m^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_m)^T \right] \\ &= \mathbf{c}^T \Sigma^{-1/2} \text{cov} [(\mathbf{x}_1^T \mathbf{W} \mathbf{y}_1, \dots, \mathbf{x}_m^T \mathbf{W} \mathbf{y}_m)^T] \Sigma^{-1/2} \mathbf{c} = \mathbf{c}^T \mathbf{c} = 1. \end{aligned} \quad (\text{B.12})$$

Denote by $\tilde{\mathbf{c}} = \Sigma^{-1/2} \mathbf{c} = (\tilde{c}_1, \dots, \tilde{c}_m)^T$. Then it holds that

$$\mathbf{c}^T \Sigma^{-1/2} (\mathbf{x}_1^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_1, \dots, \mathbf{x}_m^T (\mathbf{W} - \mathbb{E} \mathbf{W}) \mathbf{y}_m)^T = \text{tr} \left[(\mathbf{W} - \mathbb{E} \mathbf{W}) \sum_{s=1}^m \tilde{c}_s \mathbf{y}_s \mathbf{x}_s^T \right].$$

Let $\mathbf{M} = (M_{ij}) = \sum_{s=1}^m \tilde{c}_s \mathbf{y}_s \mathbf{x}_s^T$. By assumption, we have $\max_k \|\mathbf{x}_k \mathbf{y}_k^T\|_\infty \ll \|\Sigma^{1/2}\| \sim \|\Sigma^{-1/2}\|$, which entails that

$$\|\mathbf{M}\|_\infty \ll 1. \quad (\text{B.13})$$

Then it follows from the assumption of $\max_{1 \leq i, j \leq n} |w_{ij}| \leq 1$ and (B.13) that

$$\begin{aligned}
& \frac{1}{|s_n|^3} \left(\sum_{1 \leq i, j \leq n, i < j} \mathbb{E}|w_{ij}|^3 |M_{ij} + M_{ji}|^3 + \sum_{1 \leq i \leq n} \mathbb{E}|w_{ii} - \mathbb{E}w_{ii}|^3 |M_{ii}|^3 \right) \\
& \leq \frac{2}{|s_n|^3} \left(\sum_{1 \leq i, j \leq n, i < j} \mathbb{E}|w_{ij}|^2 |M_{ij} + M_{ji}|^3 + \sum_{1 \leq i \leq n} \mathbb{E}|w_{ii} - \mathbb{E}w_{ii}|^2 |M_{ii}|^3 \right) \\
& \ll \frac{2}{|s_n|^3} \left(\sum_{1 \leq i, j \leq n, i < j} \mathbb{E}|w_{ij}|^2 |M_{ij} + M_{ji}|^2 + \sum_{1 \leq i \leq n} \mathbb{E}|w_{ii} - \mathbb{E}w_{ii}|^2 |M_{ii}|^2 \right) \\
& \leq 2.
\end{aligned} \tag{B.14}$$

Since w_{ij} with $1 \leq i < j \leq n$ and $w_{ii} - \mathbb{E}w_{ii}$ with $1 \leq i \leq n$ are independent random variables with zero mean, by the Lyapunov condition (see, for example, Theorem 27.3 of ?) we can conclude that (B.10) holds. This concludes the proof of Lemma 1.

B.4 Lemma 2 and its proof

Lemma 2. *Under either model (10) and Conditions 1–2, or model (6) and Conditions 1 and 4, it holds that*

$$\|(\mathbf{D}_{-k})^{-1} + \mathcal{R}(\mathbf{V}_{-k}, \mathbf{V}_{-k}, z)\| = O(|z|) \text{ for any } z \in [a_k, b_k], \tag{B.15}$$

where a_k and b_k are defined in (8).

Proof. The conclusion of Lemma 2 has been proved in (A.16) of ?.

B.5 Lemma 3 and its proof

Lemma 3. *Under model (10) and Conditions 1–2, we have*

$$\max_{1 \leq k \leq K} \|\mathbf{v}_k\|_\infty = O\left(\frac{1}{\sqrt{n}}\right). \tag{B.16}$$

The same conclusion also holds under model (6) and Conditions 1 and 4.

Proof. We first consider model (10) and prove (B.16) under Conditions 1 and 2. In light of $\theta \mathbf{\Pi} \mathbf{\Pi}^T = \mathbf{V} \mathbf{D} \mathbf{V}^T$, we have $\theta \mathbf{\Pi} (\mathbf{\Pi}^T \mathbf{V} \mathbf{D}^{-1}) = \mathbf{V}$. This shows that $\mathbf{V}$ belongs to the space expanded by $\mathbf{\Pi}$. Thus there exists some $K \times K$ matrix $\mathbf{B}$ such that

$$\mathbf{V} = \mathbf{\Pi} \mathbf{B}. \tag{B.17}$$

Since $\mathbf{V}^T \mathbf{V} = \mathbf{I}$, it holds that $\mathbf{B}^T \mathbf{\Pi}^T \mathbf{\Pi} \mathbf{B} = \mathbf{I}$, which entails that $\mathbf{B} \mathbf{B}^T \mathbf{\Pi}^T \mathbf{\Pi} \mathbf{B} \mathbf{B}^T = \mathbf{B} \mathbf{B}^T$ and

$$\mathbf{B} \mathbf{B}^T = (\mathbf{\Pi}^T \mathbf{\Pi})^{-1}. \tag{B.18}$$

By Condition 2, we can conclude that $\|(\mathbf{\Pi}^T \mathbf{\Pi})^{-1}\| = O(n^{-1})$ and thus each entry of matrix $\mathbf{B}$ is of order $O(\frac{1}{\sqrt{n}})$. Hence in view of (B.17), the desired result can be established.

Now let us consider model (6) under Conditions 1 and 4. For this model, we also have $\mathbf{\Theta} \mathbf{\Pi} \mathbf{\Pi}^T \mathbf{\Theta} = \mathbf{V} \mathbf{D} \mathbf{V}^T$ and thus

$$\mathbf{\Theta} \mathbf{\Pi} (\mathbf{P} \mathbf{\Pi}^T \mathbf{\Theta} \mathbf{V} \mathbf{D}^{-1}) = \mathbf{V}. \quad (\text{B.19})$$

Since $\mathbf{\Theta}$ is a diagonal matrix, we can see that $\mathbf{V}$ belongs to the space expanded by $\mathbf{\Pi}$. Let $\tilde{\mathbf{\Pi}} = (\tilde{\boldsymbol{\pi}}_1, \dots, \tilde{\boldsymbol{\pi}}_n)^T$ be the submatrix of $\mathbf{\Pi}$ such that

$$\tilde{\boldsymbol{\pi}}_i = \begin{cases} \boldsymbol{\pi}_i & \text{if there exists some } 1 \leq k \leq K \text{ such that } \boldsymbol{\pi}_i(k) = 1, \\ 0 & \text{otherwise.} \end{cases}$$

By Condition 4, it holds that $c_2 n^{c_2} \mathbf{I} \leq \tilde{\mathbf{\Pi}}^T \tilde{\mathbf{\Pi}} = \sum_{i=1}^n \tilde{\boldsymbol{\pi}}_i \tilde{\boldsymbol{\pi}}_i^T \leq \sum_{i=1}^n \boldsymbol{\pi}_i \boldsymbol{\pi}_i^T = \mathbf{\Pi}^T \mathbf{\Pi}$, which leads to $\|(\mathbf{\Pi}^T \mathbf{\Pi})^{-1}\| = O(n^{-1})$. Therefore, an application of similar arguments to those for (B.17)–(B.18) concludes the proof of Lemma 3.

B.6 Lemma 4 and its proof

Lemma 4. *Under model (10) and Condition 2, it holds that*

$$\alpha_n^2 \leq n\theta, \quad d_k \gtrsim n^{1-c_2}\theta, \quad d_1 = O(n\theta), \quad k = 1, \dots, K. \quad (\text{B.20})$$

Under model (6) and Condition 4, similarly we have

$$\alpha_n^2 \leq n\theta_{\max}^2, \quad d_k \gtrsim n^{1-c_2}\theta_{\min}^2, \quad d_1 = O(n\theta_{\max}^2), \quad k = 1, \dots, K. \quad (\text{B.21})$$

Proof. We show (B.20) first. It follows from $\sum_{k=1}^K \boldsymbol{\pi}_i(k) = 1$ that $\|\mathbf{\Pi}\|_F^2 = \sum_{i=1}^n \sum_{k=1}^K \boldsymbol{\pi}_i^2(k) \leq n$ and $\lambda_1(\mathbf{\Pi}^T \mathbf{\Pi}) = O(n)$. By Condition 2, we have

$$d_K = \theta \lambda_K(\mathbf{P} \mathbf{\Pi}^T \mathbf{\Pi}) \geq \theta \lambda_K(\mathbf{\Pi}^T \mathbf{\Pi}) \lambda_K(\mathbf{P}) \geq c_0^2 \theta n^{1-c_2}$$

and

$$d_1 \leq \theta \lambda_1(\mathbf{\Pi}^T \mathbf{\Pi}) \lambda_1(\mathbf{P}) = O(\theta n).$$

Thus the second result in (B.20) is proved. Next by model (10), the (i, j) th entry h_{ij} of matrix $\mathbf{H}$ satisfies that

$$h_{ij} = \theta \sum_{s,t=1}^K \boldsymbol{\pi}_i(s) \boldsymbol{\pi}_j(t) p_{st} \leq \theta. \quad (\text{B.22})$$

Since the entries of $\mathbf{X}$ follow the Bernoulli distributions, it follows from (B.22) that $\text{var}(w_{ij}) \leq$

θ . Therefore, in view of the definition of α_n , we have

$$\alpha_n^2 = \max_j \sum_{i=1}^n \text{var}(w_{ij}) \leq n\theta.$$

The results in (B.21) can also be proved using similar arguments. This completes the proof of Lemma 4.

B.7 Lemma 5

The following 3 Lemmas follow from Lemma 9 and exactly the same proof as ?

Lemma 5. *Under either model (10) and Conditions 1–2, or model (6) and Conditions 1 and 4, for $\mathbf{u} = \mathbf{e}_i$ or $\mathbf{v}_k$ we have the following asymptotic expansions*

$$\begin{aligned} \mathbf{u}^T \hat{\mathbf{v}}_k \hat{\mathbf{v}}_k^T \mathbf{v}_k &= \left[\tilde{\mathcal{P}}_{k,t_k} - 2t_k^{-1} \tilde{\mathcal{P}}_{k,t_k}^2 \mathbf{v}_k^T \mathbf{W} \mathbf{v}_k + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n} t_k^2} \right) \right] \left[A_{\mathbf{u},k,t_k} - t_k^{-1} \mathbf{b}_{\mathbf{u},k,t_k}^T \mathbf{W} \mathbf{v}_k + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n} t_k^2} \right) \right] \\ &\quad \times \left[A_{\mathbf{v}_k,k,t_k} - t_k^{-1} \mathbf{b}_{\mathbf{v}_k,k,t_k}^T \mathbf{W} \mathbf{v}_k + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n} t_k^2} \right) \right], \end{aligned} \quad (\text{B.23})$$

$$\hat{d}_k = t_k + \mathbf{v}_k^T \mathbf{W} \mathbf{v}_k + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n} |d_k|} \right). \quad (\text{B.24})$$

B.8 Lemma 6

Lemma 6. *Under model (10) and Conditions 1–2, we have*

$$t_k [\mathbf{e}_i^T \hat{\mathbf{v}}_k - \mathbf{v}_k(i)] = \mathbf{e}_i^T \mathbf{W} \mathbf{v}_k + O_{\prec} \left(\frac{\alpha_n^2}{\sqrt{n} |t_k|} + \frac{1}{\sqrt{n}} \right). \quad (\text{B.25})$$

The same conclusion also holds under model (6) and Conditions 1 and 4.

B.9 Lemma 7

Lemma 7. *Assume that $\hat{K} = K$. Then under the mixed membership model (10) and Conditions 1–2, it holds uniformly over all i, j that*

$$\hat{w}_{ij} = w_{ij} + O_{\prec} \left(\frac{\sqrt{\theta}}{\sqrt{n}} \right). \quad (\text{B.26})$$

Under the degree-corrected mixed membership model (6), if Conditions 1 and 4–5 are satisfied, then it holds uniformly over all i, j that

$$\hat{w}_{ij} = w_{ij} + O_{\prec} \left(\frac{\theta_{\max}}{\sqrt{n}} \right). \quad (\text{B.27})$$

C Additional technical details

C.1 Lemma 8 and its proof

Lemma 8. *For any n -dimensional unit vectors $\mathbf{x}, \mathbf{y}$ and any positive integer r , we have*

$$\mathbb{E} \left[\mathbf{x}^T (\mathbf{W}^l - \mathbb{E} \mathbf{W}^l) \mathbf{y} \right]^{2r} \leq C_r (\min\{\alpha_n^{l-1}, d_{\mathbf{x}} \alpha_n^l, d_{\mathbf{y}} \alpha_n^l\})^{2r}, \quad (\text{C.1})$$

where l is any positive integer and C_r is some positive constant determined only by r .

Proof. The main idea of the proof is similar to that for Lemma 4 in ?, which is to count the number of nonzero terms in the expansion of $\mathbb{E}[\mathbf{x}^T (\mathbf{W}^l - \mathbb{E} \mathbf{W}^l) \mathbf{y}]^{2r}$. It will be made clear that the nonzero terms in the expansion consist of terms such as w_{ij}^s with $s \geq 2$. In counting the nonzero terms, we will fix one index, say i , and vary the other index j which ranges from 1 to n . Note that for any $i = 1, \dots, n$ and $s \geq 2$, we have $\sum_{j=1}^n \mathbb{E}|w_{ij}|^s \leq \alpha_n^2$ since $|w_{ij}| \leq 1$. Thus roughly speaking, counting the maximal moment of α_n is the crucial step in our proof.

Let $\mathbf{x} = (x_1, \dots, x_n)^T$, $\mathbf{y} = (y_1, \dots, y_n)^T$, and C_r be a positive constant depending only on r and whose value may change from line to line. Recall that $l, r \geq 1$ are two integers. We can expand $\mathbb{E}(\mathbf{x}^T \mathbf{W}^l \mathbf{y} - \mathbb{E} \mathbf{x}^T \mathbf{W}^l \mathbf{y})^{2r}$ to obtain the following expression

$$\begin{aligned} & \mathbb{E}(\mathbf{x}^T \mathbf{W}^l \mathbf{y} - \mathbb{E} \mathbf{x}^T \mathbf{W}^l \mathbf{y})^{2r} \\ &= \sum_{\substack{1 \leq i_1, \dots, i_{l+1}, i_{l+2}, \dots, i_{2l+2}, \dots, \\ i_{(2r-1)(l+1)+1}, \dots, i_{2r(l+1)} \leq n}} \mathbb{E} \left[(x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}} - \mathbb{E} x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}}) \times \cdots \right. \\ & \quad \times (x_{i_{(2r-1)(l+1)+1}} w_{i_{(2r-1)(l+1)+1} i_{(2r-1)(l+1)+2}} w_{i_{(2r-1)(l+1)+2} i_{(2r-1)(l+1)+3}} \cdots w_{i_{2r(l+1)-1} i_{2r(l+1)}} y_{i_{2r(l+1)}} \\ & \quad \left. - \mathbb{E} x_{i_{(2r-1)(l+1)+1}} w_{i_{(2r-1)(l+1)+1} i_{(2r-1)(l+1)+2}} w_{i_{(2r-1)(l+1)+2} i_{(2r-1)(l+1)+3}} \cdots w_{i_{2r(l+1)-1} i_{2r(l+1)}} y_{i_{2r(l+1)}}) \right]. \end{aligned} \quad (\text{C.2})$$

Let $\mathbf{i}^{(j)} = (i_{(j-1)(l+1)+1}, \dots, i_{j(l+1)})$, $j = 1, \dots, 2r$, be $2r$ vectors taking values in $\{1, \dots, n\}^{l+1}$. Then for each $\mathbf{i}^{(j)}$, we define a graph $\mathcal{G}^{(j)}$ whose vertices represent distinct values of the components of $\mathbf{i}^{(j)}$. Each adjacent component of $\mathbf{i}^{(j)}$ is connected by an undirected edge in $\mathcal{G}^{(j)}$. It can be seen that for each j , $\mathcal{G}^{(j)}$ is a connected graph, which means that there exists some path connecting any two nodes in $\mathcal{G}^{(j)}$. For each fixed $i_1, \dots, i_{l+1}, \dots, i_{(2r-1)(l+1)+1}, \dots, i_{2r(l+1)}$, consider the following term

$$\begin{aligned} & \mathbb{E} \left[(x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}} - \mathbb{E} x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}}) \times \cdots \right. \\ & \quad \times (x_{i_{(2r-1)(l+1)+1}} w_{i_{(2r-1)(l+1)+1} i_{(2r-1)(l+1)+2}} w_{i_{(2r-1)(l+1)+2} i_{(2r-1)(l+1)+3}} \cdots w_{i_{2r(l+1)-1} i_{2r(l+1)}} y_{i_{2r(l+1)}} \\ & \quad \left. - \mathbb{E} x_{i_{(2r-1)(l+1)+1}} w_{i_{(2r-1)(l+1)+1} i_{(2r-1)(l+1)+2}} w_{i_{(2r-1)(l+1)+2} i_{(2r-1)(l+1)+3}} \cdots w_{i_{2r(l+1)-1} i_{2r(l+1)}} y_{i_{2r(l+1)}}) \right], \end{aligned} \quad (\text{C.3})$$

which corresponds to graph $\mathcal{G}^{(1)} \cup \dots \cup \mathcal{G}^{(2r)}$. If there exists one graph $\mathcal{G}^{(s)}$ that is unconnected to the remaining graphs $\mathcal{G}^{(j)}$, $j \neq s$, then the corresponding expectation in (C.3) is equal to

zero. This shows that for any graph $\mathcal{G}^{(s)}$, there exists at least one connected $\mathcal{G}^{(s')}$ to ensure the nonzero expectation in (C.3). To analyze each nonzero (C.3), we next calculate how many distinct vertices are contained in the graph $\mathcal{G}^{(1)} \cup \dots \cup \mathcal{G}^{(2r)}$.

Denote by $\mathfrak{S}(2r)$ the set of partitions of the integers $\{1, 2, \dots, 2r\}$ and $\mathfrak{S}_{\geq 2}(2r)$ the subset of $\mathfrak{S}(2r)$ whose block sizes are at least two. To simplify the notation, define

$$\mathfrak{h}_j = x_{i_{(j-1)(l+1)+1}} w_{i_{(j-1)(l+1)+1} i_{(j-1)(l+1)+2}} w_{i_{(j-1)(l+1)+2} i_{(j-1)(l+1)+3}} \cdots w_{i_{j(l+1)-1} i_{j(l+1)}} y_{i_{j(l+1)}}.$$

Let $\mathcal{A} \in \mathfrak{S}_{\geq 2}(2r)$ be a partition of $\{1, 2, \dots, 2r\}$ and $|\mathcal{A}|$ the number of groups in $\mathcal{A}$. We can further define $A_j \in \mathcal{A}$ as the j th group in $\mathcal{A}$ and $|A_j|$ as the number of integers in A_j . For example, let us consider $\mathcal{A} = \{\{1, 2, 3\}, \{4, 5, \dots, 2r\}\}$. Then we have $|\mathcal{A}| = 2$, set $A_1 = \{1, 2, 3\} \in \mathcal{A}$, and $|A_1| = 3$. It is easy to see that there is a one-to-one correspondence between the partitions of $\{1, 2, \dots, 2r\}$ and the graphs $\mathcal{G}^{(1)}, \dots, \mathcal{G}^{(2r)}$ such that $\mathcal{G}^{(s)}$ and $\mathcal{G}^{(s')}$ are connected if and only if s and s' belong to one group in the partition. For any $A_j \in \mathcal{A} \in \mathfrak{S}_{\geq 2}(2r)$, there are $|A_j|l$ edges in the graph $\bigcup_{w \in A_j} \mathcal{G}^{(j)}$ since for each integer $w \in A_j$, there is a chain containing l edges by $\mathfrak{h}_w$. Since $\mathbb{E} w_{ss'} = 0$ for $s \neq s'$, in order to obtain a nonzero value of (C.3) each edge in $\bigcup_{w \in A_j} \mathcal{G}^{(j)}$ should have at least one additional copy. Thus for each nonzero (C.3), we have $\lceil \frac{|A_j|l}{2} \rceil$ distinct edges without self loops in $\bigcup_{w \in A_j} \mathcal{G}^{(j)}$. Since the graph $\bigcup_{w \in A_j} \mathcal{G}^{(j)}$ is connected, we can conclude that there are at most $\lceil \frac{|A_j|l}{2} \rceil + 1$ distinct vertices in $\bigcup_{w \in A_j} \mathcal{G}^{(j)}$. Let $\mathcal{S}(\mathcal{A})$ be the collection of all choices of $\bigcup_{s=1}^{2r} \mathbf{i}^{(s)}$ such that

1). $\bigcup_{s=1}^{2r} \mathcal{G}^{(s)}$ has the same partition as $\mathcal{A}$ such that they are connected within the same group and unconnected between groups;

2). Within each group A_j , there are at most $\lceil \frac{|A_j|l}{2} \rceil$ distinct edges without self loops and $\lceil \frac{|A_j|l}{2} \rceil + 1$ distinct vertices.

Similarly we can define $\mathcal{S}(A_j)$ since A_j can be regarded as a special partition of A_j with only one group. Summarizing the arguments above, (C.2) can be rewritten as

$$(C.2) = \sum_{\mathcal{A} \in \mathfrak{S}_{\geq 2}(2r) \cup_{s=1}^{2r}} \sum_{\mathbf{i}^{(s)} \in \mathcal{S}(\mathcal{A})} \prod_{j=1}^{|\mathcal{A}|} \left[\mathbb{E} \prod_{\gamma \in A_j} (\mathfrak{h}_\gamma - \mathbb{E} \mathfrak{h}_\gamma) \right]. \quad (C.4)$$

Let us further simplify $\mathbb{E} \prod_{\gamma \in A_j} (\mathfrak{h}_\gamma - \mathbb{E} \mathfrak{h}_\gamma)$. Let $\mathcal{B}_j$ be the set of partitions of A_j such that each partition contains exactly two groups. Without loss of generality, let $\mathcal{B}_j = \{b_{j_1}, b_{j_2}\}$, where for any $w \in A_j$, we have $w \in b_{j_1}$ or $w \in b_{j_2}$. Then it holds that

$$\left| \mathbb{E} \prod_{\gamma \in A_j} (\mathfrak{h}_\gamma - \mathbb{E} \mathfrak{h}_\gamma) \right| \leq \sum_{\gamma \in \mathcal{B}_j} \mathbb{E} \left| \prod_{\gamma \in b_{j_1}} \mathfrak{h}_\gamma \right| \prod_{\gamma \in b_{j_2}} \left| \mathbb{E} \mathfrak{h}_\gamma \right|. \quad (C.5)$$

Observe that by definition, $\mathfrak{h}_\gamma$ is the product of some independent random variables, and $\mathfrak{h}_{\gamma_1}$ and $\mathfrak{h}_{\gamma_2}$ may share some dependency through factors $w_{ab}^{m_1}$ and $w_{ab}^{m_2}$, respectively, for some

w_{ab} and nonnegative integers m_1 and m_2 . Thus in light of the inequality

$$\mathbb{E}|w_{ab}|^{m_1} \mathbb{E}|w_{ab}|^{m_2} \leq \mathbb{E}|w_{ab}|^{m_1+m_2},$$

(C.5) can be bounded as

$$(C.5) \leq 2^{|A_j|} \mathbb{E} \left| \prod_{\gamma \in A_j} \mathfrak{h}_\gamma \right|. \quad (C.6)$$

By (C.6), we can deduce

$$\begin{aligned} (C.4) &\leq 2^{2r} \sum_{\mathcal{A} \in \mathfrak{G}_{\geq 2}(2r)} \sum_{\bigcup_{s=1}^{2r} \mathbf{i}^{(s)} \in \mathcal{S}(\mathcal{A})} \prod_{j=1}^{|\mathcal{A}|} \mathbb{E} \left| \prod_{\gamma \in A_j} \mathfrak{h}_\gamma \right| \\ &\leq 2^{2r} \sum_{\mathcal{A} \in \mathfrak{G}_{\geq 2}(2r)} \prod_{j=1}^{|\mathcal{A}|} \left(\sum_{\mathbf{i}^{(s)} \in \mathcal{S}(A_j)} \mathbb{E} \left| \prod_{\gamma \in A_j} \mathfrak{h}_\gamma \right| \right). \end{aligned} \quad (C.7)$$

Thus it suffices to show that

$$\sum_{\mathbf{i}^{(s)} \in \mathcal{S}(A_j)} \mathbb{E} \left| \prod_{\gamma \in A_j} \mathfrak{h}_\gamma \right| = C_{|A_j|} (\min\{\alpha_n^{l-1}, d_{\mathbf{x}} \alpha_n^l, d_{\mathbf{y}} \alpha_n^l\})^{|A_j|},$$

using the fact that $\sum_{j=1}^{|\mathcal{A}|} |A_j| = 2r$. Without loss of generality, we prove the most difficult case of $|\mathcal{A}| = 1$, that is, there is only one connected chain which is $A = \{1, 2, \dots, 2r\}$. It has the most components in the chain $\prod_{\gamma \in A} \mathfrak{h}_\gamma$. Other cases with smaller $|A|$ can be shown in the same way. Using the same arguments as those for (C.4), we have the basic property for this chain that there are at most $\lfloor \frac{|A|l}{2} \rfloor + 1 = rl + 1$ distinct vertices and rl distinct edges without self loops.

To facilitate our technical presentation, let us introduce some additional notation. Denote by $\psi(r, l)$ the set of partitions of the edges $\{(i_s, i_{s+1}), 1 \leq s \leq 2rl, i_s \neq i_{s+1}\}$ and $\psi_{\geq 2}(r, l)$ the subset of $\psi(r, l)$ whose blocks have size at least two. Let $\tilde{\mathbf{i}} = \bigcup_{s=1}^{2r} \mathbf{i}^{(s)}$ and $P(\tilde{\mathbf{i}}) \in \psi_{\geq 2}(2l+2)$ be the partition of $\{(i_s, i_{s+1}), 1 \leq s \leq 2rl, i_s \neq i_{s+1}\}$ that is associated with the equivalence relation $(i_{s_1}, i_{s_1+1}) \sim (i_{s_2}, i_{s_2+1})$, which is defined as if and only if $(i_{s_1}, i_{s_1+1}) = (i_{s_2}, i_{s_2+1})$ or $(i_{s_1}, i_{s_1+1}) = (i_{s_2+1}, i_{s_2})$. Denote by $|P(\tilde{\mathbf{i}})| = m$ the number of groups in the partition $P(\tilde{\mathbf{i}})$ such that the edges are equivalent within each group. We further denote the distinct edges in the partition $P(\tilde{\mathbf{i}})$ as $(s_1, s_2), (s_3, s_4), \dots, (s_{2m-1}, s_{2m})$ and the corresponding counts in each group as $r_1, \dots, r_m$, and define $\tilde{\mathbf{s}} = (s_1, s_2, \dots, s_{2m})$. For the vertices, let $\phi(2m)$ be the set of partitions of $\{1, 2, \dots, 2m\}$ and $Q(\tilde{\mathbf{s}}) \in \phi(2m)$ the partition that is associated with the equivalence relation $a \sim b$, which is defined as if and only if $s_a = s_b$. Note that $s_{2j-1} \neq s_{2j}$

by the definition of the partition. By $|w_{aa}| \leq 1$, we can deduce

$$\begin{aligned}
& \sum_{\mathbf{i}^{(s)} \in \mathcal{S}(A)} \mathbb{E} \left| \prod_{\gamma \in A} \mathfrak{h}_\gamma \right| = \sum_{\mathbf{i}^{(s)} \in \mathcal{S}(A)} \mathbb{E} \left| \prod_{\gamma=1}^{2r} \mathfrak{h}_\gamma \right| \\
& \leq \sum_{\substack{1 \leq |P(\tilde{\mathbf{i}})|=m \leq rl \\ P(\tilde{\mathbf{i}}) \in \psi_{\geq 2}(2l+2)}} \sum_{\substack{\tilde{\mathbf{i}} \text{ with partition } P(\tilde{\mathbf{i}}) \\ r_1, \dots, r_m \geq 2}} \sum_{Q(\tilde{\mathbf{s}}) \in \phi(2m)} \sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) \\
& \times \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j}. \tag{C.8}
\end{aligned}$$

Denote by $\mathcal{F}_{\tilde{\mathbf{s}}}$ the graph constructed by the edges of $\tilde{\mathbf{s}}$. Since the edges in $\tilde{\mathbf{s}}$ are the same as those of the edges in $\bigcup_{s=1}^{2r} \mathcal{G}^{(s)}$ with the structure $\mathcal{S}(A)$, we can see that $\mathcal{F}_{\tilde{\mathbf{s}}}$ is also a connected graph. In view of (C.8), putting term $|x_{i_1} y_{i_{l+1}} x_{i_{l+2}} y_{i_{2l+2}}|$ aside we need to analyze the summation

$$\sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j}.$$

If index s_{2k-1} satisfies that $s_{2k-1} \neq s$ for all $s \in \{s_1, \dots, s_{2m}\} \setminus \{s_{2k-1}\}$, that is, index s_{2k-1} appears only in one $w_{s_{2j-1}s_{2j}}$, we call s_{2k-1} a single index (or single vertex). If there exists some single index s_{2k-1} , then it holds that

$$\begin{aligned}
& \sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j} \\
& \leq \sum_{\substack{\tilde{\mathbf{s}} \setminus \{s_{2k-1}\} \text{ with partition } Q(\tilde{\mathbf{s}} \setminus \{s_{2k-1}\}) \\ 1 \leq s_1, \dots, s_{2k-2}, s_{2k+2}, s_{2m} \leq n \\ s_{2k} = s_j \text{ for some } 1 \leq j \leq 2m}} \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j} \sum_{s_{2k-1}=1}^n \mathbb{E} |w_{s_{2k-1}s_{2k}}|^{r_k}. \tag{C.9}
\end{aligned}$$

Note that since graph $\mathcal{F}_{\tilde{\mathbf{s}}}$ is connected and index s_{2k-1} is single, there exists some j such that $s_j = s_{2k}$, which means that in the summation $\sum_{s_{2k-1}=1}^n \mathbb{E} |w_{s_{2k-1}s_{2k}}|^{r_k}$, index s_{2k} is fixed. Then it follows from the definition of α_n , $|w_{ij}| \leq 1$, and $r_k \geq 2$ that

$$\sum_{s_{2k-1}=1}^n \mathbb{E} |w_{s_{2k-1}s_{2k}}|^{r_k} \leq \alpha_n^2.$$

After taking the summation over index s_{2k-1} , we can see that there is one less edge in $\mathcal{F}(\tilde{\mathbf{s}})$. That is, by taking the summation above we will have one additional α_n^2 in the upper bound while removing one edge from graph $\mathcal{F}(\tilde{\mathbf{s}})$. For the single index s_{2k} , we also have the same bound. If s_{2k_1-1} is not a single index, without loss of generality we assume that $s_{2k_1-1} = s_{2k-1}$. Then this vertex s_{2k-1} needs some delicate analysis. By the assumption of

$|w_{ij}| \leq 1$, we have

$$\mathbb{E}|w_{2k-1,2k}|^{r_k} |w_{2k_1-1,2k_1}|^{r_{k_1}} \leq \frac{\mathbb{E}|w_{2k-1,2k}|^{r_k} + \mathbb{E}|w_{2k_1-1,2k_1}|^{r_{k_1}}}{2}.$$

Then it holds that

$$\begin{aligned} & \sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^m \mathbb{E}|w_{s_{2j-1}s_{2j}}|^{r_j} \\ & \leq \frac{1}{2} \sum_{\substack{\tilde{\mathbf{s}} \setminus (s_{2k-1}, s_{2k_1-1}) \text{ with partition } Q(\tilde{\mathbf{s}} \setminus (s_{2k-1}, s_{2k_1-1})) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1, j \neq k}^m \mathbb{E}|w_{s_{2j-1}s_{2j}}|^{r_j} \\ & + \frac{1}{2} \sum_{\substack{\tilde{\mathbf{s}} \setminus (s_{2k-1}, s_{2k_1-1}) \text{ with partition } Q(\tilde{\mathbf{s}} \setminus (s_{2k-1}, s_{2k_1-1})) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1, j \neq k_1}^m \mathbb{E}|w_{s_{2j-1}s_{2j}}|^{r_j}. \end{aligned} \quad (\text{C.10})$$

Note that since $\mathcal{F}_{\tilde{\mathbf{s}}}$ is a connected graph, if we delete either edge (s_{2k-1}, s_{2k}) or edge (s_{2k_1-1}, s_{2k_1}) from graph $\mathcal{F}_{\tilde{\mathbf{s}}}$, the resulting graph is also connected. Then the two summations on the right hand side of (C.10) can be reduced to the case in (C.9) for the graph with edge (s_{2k-1}, s_{2k}) or (s_{2k_1-1}, s_{2k_1}) removed, since s_{2k-1} or s_{2k_1-1} is a single index in the subgraph. Similar to (C.9), after taking the summation over index s_{2k-1} or s_{2k_1-1} there are two less edges in graph $\mathcal{F}_{\tilde{\mathbf{s}}}$ and thus we now obtain $2\alpha_n^2$ in the upper bound.

For the general case when there are m_1 vertices belonging to the same group, without loss of generality we denote them as $w_{ab_1}, \dots, w_{ab_{m_1}}$. If for any k graph $\mathcal{F}_{\tilde{\mathbf{s}}}$ is still connected after deleting edges $(a, b_1), \dots, (a, b_{k-1}), (a, b_{k+1}), \dots, (a, b_{m_1})$, then we repeat the process in (C.10) to obtain a new connected graph by deleting $k-1$ edges in $w_{ab_1}, \dots, w_{ab_{m_1}}$ and thus obtain $k\alpha_n^2$ in the upper bound. Motivated by the key observations above, we carry out an iterative process in calculating the upper bound as follows.

- (1) If there exists some single index in $\tilde{\mathbf{s}}$, using (C.9) we can calculate the summation over such an index and then delete the edge associated with this vertex in $\mathcal{F}_{\tilde{\mathbf{s}}}$. The corresponding vertices associated with this edge are also deleted. For simplicity, we also denote the new graph as $\mathcal{F}_{\tilde{\mathbf{s}}}$. In this step, we obtain α_n^2 in the upper bound.
- (2) Repeat (1) until there is no single index in graph $\mathcal{F}_{\tilde{\mathbf{s}}}$.
- (3) Suppose there exists some index associated with k edges such that graph $\mathcal{F}_{\tilde{\mathbf{s}}}$ is still connected after deleting any $k-1$ edges. Without loss of generality, let us consider the case of $k=2$. Then we can apply (C.9) to obtain α_n^2 in the upper bound. Moreover, we delete k edges associated with this vertex in $\mathcal{F}_{\tilde{\mathbf{s}}}$.
- (4) Repeat (3) until there is no such index.
- (5) If there still exists some single index, go back to (1). Otherwise stop the iteration.

Completing the graph modification process mentioned above, we can obtain a final graph $\mathbf{Q}$ that enjoys the following properties:

- i) Each edge does not contain any single index;
- ii) Deleting any vertex makes the graph disconnected.

Let $\mathbf{S}_{\mathbf{Q}}$ be the spanning tree of graph $\mathbf{Q}$, which is defined as the subgraph of $\mathbf{Q}$ with the minimum possible number of edges. Since $\mathbf{S}_{\mathbf{Q}}$ is a subgraph of $\mathbf{Q}$, it also satisfies property ii) above. Assume that $\mathbf{S}_{\mathbf{Q}}$ contains p edges. Then the number of vertices in $\mathbf{S}_{\mathbf{Q}}$ is $p + 1$. Denote by $q_1, \dots, q_{p+1}$ the vertices of $\mathbf{S}_{\mathbf{Q}}$ and $\deg(q_i)$ the degree of vertex q_i . Then by the degree sum formula, we have $\sum_{i=1}^{p+1} \deg(q_i) = 2p$. As a result, the spanning tree has at least two vertices with degree one and thus there exists a subgraph of $\mathbf{S}_{\mathbf{Q}}$ without either of the vertices that is connected. This will result in a contradiction with property ii) above unless the number of vertices in graph $\mathbf{Q}$ is exactly one. Since l is a bounded constant, the numbers of partitions $P(\tilde{\mathbf{i}})$ and $Q(\tilde{\mathbf{s}})$ are also bounded. It follows that

$$(C.8) \leq C_r d_{\mathbf{x}}^{2r} d_{\mathbf{y}}^{2r} \sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1} s_{2j}}|^{r_j}, \quad (C.11)$$

where $d_{\mathbf{x}} = \|\mathbf{x}\|_{\infty}$, $d_{\mathbf{y}} = \|\mathbf{y}\|_{\infty}$, and C_r is some positive constant determined by l . Combining these arguments above and noticing that there are at most l distinct edges in graph $\mathcal{F}_{\tilde{\mathbf{s}}}$, we can obtain

$$\begin{aligned} (C.11) &\leq C_r d_{\mathbf{x}}^{2r} d_{\mathbf{y}}^{2r} \alpha_n^{2rl-2} \sum_{1 \leq s_{2k_0-1}, s_{2k_0} \leq n, (s_{2k_0-1}, s_{2k_0}) = \mathbf{Q}} \mathbb{E} |w_{s_{2k_0-1} s_{2k_0}}|^{r_{k_0}} \\ &\leq C_r d_{\mathbf{x}}^{2r} d_{\mathbf{y}}^{2r} \alpha_n^{2rl} n. \end{aligned} \quad (C.12)$$

Therefore, we have established a simple upper bound of $C_r d_{\mathbf{x}}^{2r} d_{\mathbf{y}}^{2r} \alpha_n^{2rl} n$.

In fact, we can improve the aforementioned upper bound to $C_r \alpha_n^{r(l-1)}$. Note that the process mentioned above did not utilize the condition that both $\mathbf{x}$ and $\mathbf{y}$ are unit vectors, that is, $\|\mathbf{x}\| = \|\mathbf{y}\| = 1$. Since term $\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)$ is involved in (C.8), we can analyze them together with random variables w_{ij} . First, we need to deal with some distinct lower indices with low moments in $\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)$. If there are two distinct lower indices, without loss of generality denoted them as i_s and $i_{s'}$ and then the corresponding entries are x_{i_s} (or y_{i_s}) and $y_{i_{s'}}$ (or $x_{i_{s'}}$). Moreover, there are only one x_{i_s} and $y_{i_{s'}}$ involved in $\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)$. Without loss of generality, let us assume that

$s = 1$ and $s' = l + 1$. Then it holds that

$$\begin{aligned} \prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) &= |x_{i_1}| |y_{i_{l+1}}| \prod_{j=2}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) \\ &\leq \frac{x_{i_1}^2}{2} \prod_{j=2}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) + \frac{y_{i_{l+1}}^2}{2} \prod_{j=2}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|). \end{aligned} \quad (\text{C.13})$$

That is, if we have two lower indices and each index appears only once in the product above, we can use (C.13) to increase the moment of x_{i_s} (or $y_{i_{s'}}$) and delete the other one. For (C.13), it is equivalent for us to consider the case when the lower index $i_1 = i_{l+1}$. Repeating the procedure (C.13), finally we can obtain a product $\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)$ with the following properties:

1). Except for one vertex i_{s_0} , for each i_s with $s \neq s_0$ there exists some $i_{s'}$ such that $i_s = i_{s'}$ with $s \neq s'$.

2). Except for one vertex i_{s_0} , for each i_s with $s \neq s_0$ the term $x_{i_s}^{m_1} y_{i_s}^{m_2}$ involved in $\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)$ satisfies the condition that $m_1 + m_2 \geq 2$. Moreover, at least one of m_1 and m_2 is larger than one.

By the properties above, let us denote by $\Upsilon(2r)$ the set of partitions of the vertices $\{i_{(j-1)(l+1)+1}, i_{j(l+1)}, j = 1, \dots, 2r\}$ such that except for one group, the remaining groups in Υ with $\Upsilon \in \Upsilon(2r)$ have blocks with size at least two. There are three different cases to consider.

Case 1). All the groups in Υ have block size two. Then it follows that

$$\left| \prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) \right| = \prod_{k=1}^{|\Upsilon|} |x_{i_s}^{m_{1k}} y_{i_k}^{m_{2k}}|, \quad (\text{C.14})$$

where $m_{1k} + m_{2k} = 2$. In fact, by the second property of Υ above, $m_{1k} = 0$ or $m_{2k} = 0$. Without loss of generality, we assume that $m_{2k} = 0$. Then we need only to consider the equation

$$\left| \prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|) \right| = \prod_{k=1}^{|\Upsilon|} |x_{i_k}|^2.$$

Then by (C.8), it remains to bound

$$\sum_{\substack{\tilde{\mathbf{s}} \text{ with partition } Q(\tilde{\mathbf{s}}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{k=1}^{|\Upsilon|} |x_{i_k}|^2 \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1} s_{2j}}|^{r_j}. \quad (\text{C.15})$$

To simplify the presentation, assume without loss of generality that $i_k = s_k$, $k =$

$1, \dots, |\Upsilon|$. Then the summation in (C.15) becomes

$$\sum_{\substack{\mathbf{s} \text{ with partition } Q(\mathbf{s}) \\ 1 \leq s_1, \dots, s_{2m} \leq n}} \prod_{j=1}^{|\Upsilon|} |x|_{s_j}^2 \prod_{j=1}^m \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j}.$$

By repeating the iterative process (1)–(5) mentioned before, we can bound the summation for fixed $s_2, \dots, s_{|\Upsilon|}$ and obtain an alternative upper bound

$$\sum_{s_1=1}^n x_{s_1}^2 \mathbb{E} |w_{s_{2j-1}s_{2j}}|^{r_j} \leq \sum_{s_1=1}^n x_{s_1}^2 = 1$$

since $\mathbf{x}$ is a unit vector. Thus for this step of the iteration, we obtain term one instead of α_n^2 in the upper bound. Repeat this step until there is only $x_{s_{|\Upsilon|}}^2$ left. Since the graph is always connected during the iteration process, there exists another vertex b such that $w_{s_{|\Upsilon|}b}$ is involved in (C.15). For index $s_{|\Upsilon|}$, we do not delete the edges containing $s_{|\Upsilon|}$ in the graph during the iterative process (1)–(5). Then after the iteration stops, the final graph $\mathbf{Q}$ satisfies properties i) and ii) defined earlier except for vertex $s_{|\Upsilon|}$. Since there are at least two vertices with degree one in $\mathbf{S}_{\mathbf{Q}}$, we will also reach a contradiction unless the number of vertices in graph $\mathbf{Q}$ is exactly one. By (C.14), it holds that $2|\Upsilon| = 4r$. As a result, we can obtain the upper bound

$$(C.8) \leq C_r \alpha_n^{2rl-2|\Upsilon|} \sum_{1 \leq s_{|\Upsilon|}, b \leq n, (s_{|\Upsilon|}, b) = \mathbf{Q}} \mathbb{E} x_{s_{|\Upsilon|}}^2 |w_{s_{|\Upsilon|}b}|^r \leq C_r \alpha_n^{2rl-2r} \quad (C.16)$$

with C_r some positive constant. Therefore, the improved bound $C_r \alpha_n^{2r(l-1)}$ is shown for this case.

Case 2). All the groups in Υ have block size at least two and there is at least one block with size larger than two. Then it follows that

$$|\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)| = \prod_{k=1}^{|\Upsilon|} |x|_{i_s}^{m_{1k}} |y|_{i_k}^{m_{2k}}.$$

Since $m_{1k} + m_{2k} \geq 2$ by the second property of Υ above, define the nonnegative integer $r_1 = \sum_{k=1}^{|\Upsilon|} (m_{1k} + m_{2k} - 2)$. There are at most $\lceil \frac{2rl+2-r_1}{2} \rceil$ distinct vertices in the graph $\mathcal{F}_{\mathbf{s}}$ and at most $\lceil \frac{2rl+2-r_1}{2} \rceil - 1$ distinct edges. Similar to Case 1 with less distinct edges, we have

$$(C.8) \leq C \alpha_n^{2\lceil \frac{2rl+2-r_1}{2} \rceil - 2|\Upsilon| - 2} \sum_{1 \leq s_1, b \leq n, (s_1, b) = \mathbf{Q}} \mathbb{E} x_{s_1}^2 |w_{s_1 b}|^r \leq C \alpha_n^{2\lceil \frac{2rl+2-r_1}{2} \rceil - 2|\Upsilon|}. \quad (C.17)$$

By the definition of r_1 and $\sum_{k=1}^{|\Upsilon|} (m_{1k} + m_{2k}) = 4r$, it holds that

$$r_1 + 2|\Upsilon| = 4r.$$

Thus r_1 is an even number and $2\lceil \frac{2rl+2-r_1}{2} \rceil - 2|\Upsilon| = 2rl - r_1 - 2|\Upsilon| + 2 \leq 2rl - 2r$. The improved bound $C_r \alpha_n^{2r(l-1)}$ is also shown for this case.

Case 3). Except for one index i_{k_0} , the other groups in Υ have block size at least two. Let us define $r'_1 = \sum_{k=1, k \neq k_0}^{|\Upsilon|} (m_{1k} + m_{2k} - 2)$. There are at most $\lceil \frac{2rl+2-r'_1}{2} \rceil$ distinct vertices and at most $\lceil \frac{2rl+2-r'_1}{2} \rceil - 1$ distinct edges. For the parameter $|x_{i_{k_0}}|$ (or $|y_{i_{k_0}}|$), we can bound it by one since $\mathbf{x}$ and $\mathbf{y}$ are unit vectors. Then similar to Case 2, we can deduce

$$(C.8) \leq C \alpha_n^{2\lceil \frac{2rl+2-r'_1}{2} \rceil - 2|\Upsilon|} \sum_{1 \leq s_1, b \leq n, (s_1, b) = \mathbf{Q}} \mathbb{E} x_{s_1}^2 |w_{s_1 b}|^r \leq C \alpha_n^{2\lceil \frac{2rl+2-r'_1}{2} \rceil - 2|\Upsilon| + 2}. \quad (C.18)$$

By the definition of r'_1 in this case, it holds that

$$r'_1 + 2|\Upsilon| = 4r + 1.$$

Then r'_1 is an odd number and thus

$$2\lceil \frac{2rl+2-r'_1}{2} \rceil - 2|\Upsilon| + 2 = 2rl - r_1 - 2|\Upsilon| + 3 \leq 2rl - 2r.$$

Summarizing the arguments above, for this case we can also obtain the desired bound $C_r \alpha_n^{2r(l-1)}$.

In addition, we can also improve the upper bound to $C_r(\min\{d_{\mathbf{x}}^{2r} \alpha_n^{2rl}, d_{\mathbf{y}}^{2r} \alpha_n^{2rl}\})$. The technical arguments for this refinement are similar to those for the improvement to order $C_r \alpha_n^{2r(l-1)}$ above. As an example, we can bound the components of $\mathbf{y}$ by $d_{\mathbf{y}} = \|\mathbf{y}\|_{\infty}$, which leads to $|\prod_{j=1}^{2r} (|x_{i_{(j-1)(l+1)+1}}| |y_{i_{j(l+1)}}|)| \leq d_{\mathbf{y}}^{2r} |\prod_{j=1}^{2r} |x_{i_{(j-1)(l+1)+1}}||$. Then the analysis becomes similar to the three cases above. The only difference is that $\sum_{k=1}^{|\Upsilon|} m_{1k} = 2r$ instead of $\sum_{k=1}^{|\Upsilon|} (m_{1k} + m_{2k}) = 4r$. For this case, we have

$$(C.8) \leq C d_{\mathbf{y}}^{2r} \alpha_n^{2rl-2|\Upsilon|} \sum_{1 \leq s_2, b \leq n, (s_2, b) = \mathbf{Q}} \mathbb{E} x_{s_1}^2 |w_{s_1 b}|^r \leq C_r d_{\mathbf{y}}^{2r} \alpha_n^{2rl}. \quad (C.19)$$

Thus we can obtain the claimed upper bound $C_r(\min\{d_{\mathbf{x}}^{2r} \alpha_n^{2rl}, d_{\mathbf{y}}^{2r} \alpha_n^{2rl}\})$. Therefore, combining the two aforementioned improved bounds yields the desired upper bound of

$$C_r(\min\{\alpha_n^{2r(l-1)}, d_{\mathbf{x}}^{2r} \alpha_n^{2rl}, d_{\mathbf{y}}^{2r} \alpha_n^{2rl}\}),$$

which completes the proof of Lemma 8.

C.2 Corollary 1 and its proof

Lemma 8 ensures the following corollary immediately.

Corollary 1. *Under the conditions of Lemma 8, it holds that for any positive constants a and b , there exists some $n_0(a, b) > 0$ such that*

$$\sup_{\|\mathbf{x}\|=\|\mathbf{y}\|=1} \mathbb{P}\left(\mathbf{x}^T(\mathbf{W}^l - \mathbb{E}\mathbf{W}^l)\mathbf{y} \geq n^a \min\{\alpha_n^{l-1}, d_{\mathbf{x}}\alpha_n^l, d_{\mathbf{y}}\alpha_n^l\}\right) \leq n^{-b} \quad (\text{C.20})$$

for any $n \geq n_0(a, b)$ and $l \geq 1$. Moreover, we have

$$\mathbf{x}^T(\mathbf{W}^l - \mathbb{E}\mathbf{W}^l)\mathbf{y} = O_{\prec}(\min\{\alpha_n^{l-1}, d_{\mathbf{x}}\alpha_n^l, d_{\mathbf{y}}\alpha_n^l\}). \quad (\text{C.21})$$

Proof. It suffices to show (C.20) because then (C.21) follows from the definition. For any positive constants a and b , there exists some integer r such that $2ar \geq b + 1$. By the Chebyshev inequality, it holds that

$$\begin{aligned} & \sup_{\|\mathbf{x}\|=\|\mathbf{y}\|=1} \mathbb{P}(|\mathbf{x}^T(\mathbf{W}^l - \mathbb{E}\mathbf{W}^l)\mathbf{y}| \geq n^a \min\{\alpha_n^{l-1}, d_{\mathbf{x}}\alpha_n^l, d_{\mathbf{y}}\alpha_n^l\}) \\ & \leq \sup_{\|\mathbf{x}\|=\|\mathbf{y}\|=1} \frac{\mathbb{E}(\mathbf{x}^T(\mathbf{W}^l - \mathbb{E}\mathbf{W}^l)\mathbf{y})^{2r}}{n^{2ar}(\min\{\alpha_n^{l-1}, d_{\mathbf{x}}\alpha_n^l, d_{\mathbf{y}}\alpha_n^l\})^{2r}} \leq \frac{C_r}{n^{b+1}}, \end{aligned}$$

which can be further bounded by n^{-b} as long as $n \geq C_r$. It is seen that C_r is determined completely by a and b . This concludes the proof of Corollary 1.

C.3 Lemma 9 and its proof

Lemma 9. *For any n -dimensional unit vectors $\mathbf{x}$ and $\mathbf{y}$, we have*

$$\mathbb{E}\mathbf{x}^T\mathbf{W}^l\mathbf{y} = O(\alpha_n^l), \quad (\text{C.22})$$

where $l \geq 2$ is a positive integer. Furthermore, if the number of nonzero components of $\mathbf{x}$ is bounded, then it holds that

$$\mathbb{E}\mathbf{x}^T\mathbf{W}^l\mathbf{y} = O(\alpha_n^l d_{\mathbf{y}}), \quad (\text{C.23})$$

where $d_{\mathbf{y}} = \|\mathbf{y}\|_{\infty}$.

Proof. The result in (C.22) follows directly from Lemma 5 of ?. Thus it remains to show (C.23). The main idea of the proof is similar to that for the proof of Lemma 8. Denote by $\mathfrak{C}$ the set of positions of the nonzero components of $\mathbf{x}$. Then we have

$$\mathbb{E}\mathbf{x}^T\mathbf{W}^l\mathbf{y} = \sum_{\substack{i_1 \in \mathfrak{C}, 1 \leq i_2, \dots, i_{l+1} \leq n \\ i_s \neq i_{s+1}}} \mathbb{E}(x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}}). \quad (\text{C.24})$$

Note that the cardinality of set $\mathfrak{C}$ is bounded. Thus it suffices to show that for fixed i_1 , we have

$$\sum_{\substack{1 \leq i_2, \dots, i_{l+1} \leq n \\ i_s \neq i_{s+1}}} \mathbb{E} (x_{i_1} w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}} y_{i_{l+1}}) = O(d_{\mathbf{y}} \alpha_n^l). \quad (\text{C.25})$$

By the definition of graph $\mathcal{G}^{(1)}$ in the proof of Lemma 8, we can also get a similar expression as (C.6) that

$$\begin{aligned} & |(\text{C.24})| \\ & \leq d_{\mathbf{y}} \sum_{\substack{\mathcal{G}^{(1)} \text{ with at most } \lfloor l/2 \rfloor \text{ distinct edges without self loops and } \lfloor l/2 \rfloor + 1 \text{ distinct vertices, } i_1 \text{ is fixed}}} \mathbb{E} |w_{i_1 i_2} w_{i_2 i_3} \cdots w_{i_l i_{l+1}}|. \end{aligned} \quad (\text{C.26})$$

Using similar arguments for bounding the order of the summation through the iterative process as those for (C.11)–(C.12) in the proof of Lemma 8, we can obtain a similar bound

$$\mathbb{E} \mathbf{x}^T \mathbf{W}^l \mathbf{y} \leq C d_{\mathbf{y}} \alpha_n^{l-2} \sum_{i_{k_0}=1}^n \mathbb{E} |w_{i_1 i_{k_0}}|^{r_0} \leq C d_{\mathbf{y}} \alpha_n^l \quad (\text{C.27})$$

with $r_0 \geq 2$. Here we do not remove the lower index i_1 during the iteration procedure. The additional factor n on the right hand side of (C.12) can be eliminated since i_1 is fixed. This completes the proof of Lemma 9.

C.4 Lemma 10 and its proof

Lemma 10. Assume that $\xi_1 = O_{\prec}(\zeta), \dots, \xi_m = O_{\prec}(\zeta)$ with $m = \lfloor n^c \rfloor$ and c some positive constant. If

$$\mathbb{P} [|\xi_i| > n^a |\zeta|] \leq n^{-b} \quad (\text{C.28})$$

uniformly for $\xi_i, i = 1, \dots, m$, and any positive constants a, b with $n \geq n_0(a, b)$, then for any positive random variables $X_1, \dots, X_m$, we have

$$\sum_{i=1}^m X_i \xi_i = O_{\prec} \left(\sum_{i=1}^m X_i \zeta \right).$$

Proof. For any positive constants a and b , let $b_1 = c + b$. By (C.28), it holds that

$$\mathbb{P} [|\xi_i| > n^a |\zeta|] \leq n^{-b_1}$$

for all $n \geq n_0(a, b_1)$, where $n_0(a, b_1)$ is determined completely by a and b_1 . Then we have

$$\mathbb{P} \left[\left| \sum_{i=1}^m X_i \xi_i \right| > n^a |\zeta| \sum_{i=1}^m X_i \right] \leq \sum_{i=1}^m \mathbb{P} [|\xi_i| > n^a |\zeta|] \leq n^{-b}$$

for large enough $n \geq n_0(a, b_1)$. Since $b_1 = c + b$ and c is fixed, the constant $n_0(a, b_1)$ is determined essentially by a and b . This concludes the proof of Lemma 10.

C.5 Lemma 11 and its proof

Lemma 11. *For any positive constant $\mathfrak{L}$, it holds that*

$$\mathbb{P}(\|\mathbf{W}\| \geq \alpha_n \log n) \leq n^{-\mathfrak{L}}$$

for all sufficiently large n .

Proof. The conclusion of Lemma 11 follows directly from Theorem 6.2 of ?. We can also prove it by (B.1) and the inequality with $c\sqrt{\log n \alpha_n} - 1$ replaced by $\alpha_n \log n$ in (B.2).

C.6 Lemma 12

Lemma 12 (?). *There exists a unique solution $z = t_k$ to equation (9) on the interval $[a_k, b_k]$, and thus t_k 's are well defined. In addition, for each $k = 1, \dots, K$, we have $t_k/d_k \rightarrow 1$ as $n \rightarrow \infty$.*

D Sufficient conditions for Condition 3

D.1 Lemma 13 and its proof

Lemma 13. *Under Conditions 1–2, if $\theta < 1$ and $\min_{1 \leq i, j \leq K} \mathbf{P}_{ij} \geq c$ for some positive constant c , then Condition 3 holds.*

Proof. The key step of the proof is to calculate $\text{cov}[(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{WV}]$. Without loss of generality, let us assume that $(i, j) = (1, 2)$. Note that the main difference between the null and alternative hypotheses is that the mean value of $(\mathbf{e}_1 - \mathbf{e}_2)^T \mathbb{E}\mathbf{W}$ is 0 under the former and is $(\mathbb{E}w_{1,1}, -\mathbb{E}w_{2,2}, 0, \dots, 0)^T$, which may be nonzero, under the latter. However, since the main idea of the proof applies to both cases, we will provide only the technical details under the null hypothesis.

First, some direct calculations show that

$$\begin{aligned} \theta^{-1} \mathbf{D} \Sigma_1 \mathbf{D} &= \theta^{-1} \text{cov}[(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{WV}] \\ &= \theta^{-1} \mathbf{V}^T \mathbb{E}(\mathbf{W}(\mathbf{e}_i - \mathbf{e}_j)(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W}) \mathbf{V} \\ &= \theta^{-1} \mathbf{V}^T \mathbf{QV}, \end{aligned} \tag{D.1}$$

where $\mathbf{Q} = \text{diag}(\mathbb{E}(\mathbf{w}_{i1} - \mathbf{w}_{j1})^2, \dots, \mathbb{E}(\mathbf{w}_{in} - \mathbf{w}_{jn})^2) + \mathbb{E}\mathbf{w}_{ij}^2 \mathbf{e}_i \mathbf{e}_j^T + \mathbb{E}\mathbf{w}_{ij}^2 \mathbf{e}_j \mathbf{e}_i^T$. By the assumptions that $\theta < 1$ and $\min_{1 \leq i, j \leq K} \mathbf{P}_{ij} \geq c$, we see that the entries of the mean matrix $\mathbf{H} = (h_{ij})$ are bounded from below by $c\theta$ and from above by θ . Since $\mathbb{E}w_{ij}^2 \sim h_{ij}$ and w_{ik} and w_{jk} are independent for $i \neq j$, it holds that

$$\theta \mathbf{I} \lesssim \text{diag}(\mathbb{E}(w_{i1} - w_{j1})^2, \dots, \mathbb{E}(w_{in} - w_{jn})^2) \lesssim \theta \mathbf{I}. \quad (\text{D.2})$$

Then it follows from (D.2) that

$$\mathbf{I} \lesssim \theta^{-1} \mathbf{V}^T \text{diag}(\mathbb{E}(w_{i1} - w_{j1})^2, \dots, \mathbb{E}(w_{in} - w_{jn})^2) \mathbf{V} \lesssim \mathbf{I}. \quad (\text{D.3})$$

Since $\Sigma_1 \in \mathbb{R}^{K \times K}$ with K a finite integer, we can deduce that

$$\|\theta^{-1} \mathbf{V}^T (\mathbb{E}w_{ij}^2 \mathbf{e}_i \mathbf{e}_j^T + \mathbb{E}w_{ij}^2 \mathbf{e}_j \mathbf{e}_i^T) \mathbf{V}\| \lesssim \frac{1}{n}. \quad (\text{D.4})$$

Therefore, combining (D.1)–(D.4), we can obtain the desired conclusion under the null hypothesis. This completes the proof of Lemma 13.

E Uniform convergence

Theorem 1. *Assume that the null hypotheses $H_{0,ij} : \pi_i = \pi_j$ hold for all $1 \leq i \neq j \leq n$. Then*

1) *Under Conditions 1–3 and the mixed membership model (10), we have for any $x \in \mathbb{R}$,*

$$\lim_{n \rightarrow \infty} \sup_{1 \leq i \neq j \leq n} |\mathbb{P}(T_{ij} \leq x) - \mathbb{P}(\chi_K^2 \leq x)| = 0. \quad (\text{E.1})$$

2) *Under Conditions 1 and 4–7 and DCMM (6), we have for any $x \in \mathbb{R}$,*

$$\lim_{n \rightarrow \infty} \sup_{1 \leq i \neq j \leq n} |\mathbb{P}(G_{ij} \leq x) - \mathbb{P}(\chi_{K-1}^2 \leq x)| = 0. \quad (\text{E.2})$$

Proof. We provide the detailed proof only for (E.1) since the proof of (E.2) is almost identical. Recall that $T_{ij} = \|\Sigma_1^{-1/2}(\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j))\|^2$. Let us investigate the asymptotic behavior of random vector $\Sigma_1^{-1/2}(\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j))$. Checking the proof of Theorem 1 in Section A.1 carefully, we can see that there exists some positive constant ϵ such that

$$\begin{aligned} & \Sigma_1^{-1/2}(\widehat{\mathbf{V}}(i) - \widehat{\mathbf{V}}(j)) \\ &= \Sigma_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T + O_{\prec}(n^{-\epsilon}), \end{aligned} \quad (\text{E.3})$$

where the $o_p(1)$ term in (8) is replaced by $O_{\prec}(n^{-\epsilon})$. By (E.3) and the continuity of the

standard multivariate Gaussian distribution, it suffices to show that for any convex set $\mathbf{S} \subset \mathbb{R}^K$, we have

$$\lim_{n \rightarrow \infty} \sup_{i \neq j} \left| \mathbb{P} \left(\boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T \in \mathbf{S} \right) - \mathbb{P}(\mathbf{x}_K \in \mathbf{S}) \right| = 0, \quad (\text{E.4})$$

where $\mathbf{x}_K \sim N(\mathbf{0}, \mathbf{I}_K)$.

For an application of Theorem 1.1 in ?, we need to rewrite

$$\boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T$$

as the sum of independent random vectors. Indeed, some direct calculations yield

$$\begin{aligned} & \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_1}{t_1/d_1}, \dots, \frac{(\mathbf{e}_i - \mathbf{e}_j)^T \mathbf{W} \mathbf{v}_K}{t_K/d_K} \right)^T \\ &= \sum_{l=1}^n \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} \left(\frac{(w_{il} - w_{jl}) \mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{(w_{il} - w_{jl}) \mathbf{v}_{Kl}}{t_K/d_K} \right)^T \\ &= \sum_{l \neq i, j} \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T \\ & \quad + \sum_{l \in \{i, j\}} \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T, \end{aligned} \quad (\text{E.5})$$

where the first term in the last step is the sum of independent random vectors. Then it follows from Lemma 3 and Condition 3 that

$$\sum_{l \in \{i, j\}} \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T = O\left(\frac{1}{\sqrt{n\theta}}\right).$$

Combining this with (E.4) and (E.5), we see that it remains to show that

$$\lim_{n \rightarrow \infty} \sup_{i \neq j} \left| \mathbb{P} \left(\sum_{l \neq i, j} \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T \in \mathbf{S} \right) - \mathbb{P}(\mathbf{x}_K \in \mathbf{S}) \right| = 0. \quad (\text{E.6})$$

From Theorem 1.1 in ?, Condition 3, and Lemma 3, we can deduce that for any fixed

i, j , there exists some positive constant C (independent of i, j) such that

$$\begin{aligned}
& \left| \mathbb{P} \left(\sum_{l \neq i, j} \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T \in \mathbf{S} \right) - \mathbb{P}(\mathbf{x}_K \in \mathbf{S}) \right| \\
& \leq C \sum_{l \neq i, j} \mathbb{E} \left\| \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} (w_{il} - w_{jl}) \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T \right\|_2^3 \\
& = C \sum_{l \neq i, j} \left(\left\| \boldsymbol{\Sigma}_1^{-1/2} \mathbf{D}^{-1} \left(\frac{\mathbf{v}_{1l}}{t_1/d_1}, \dots, \frac{\mathbf{v}_{Kl}}{t_K/d_K} \right)^T \right\|_2^3 \times \mathbb{E} |w_{il} - w_{jl}|^3 \right) \\
& \leq \frac{C^2}{\sqrt{n}} \max_{l \neq i, j} \mathbb{E} \left| \frac{w_{il} - w_{jl}}{\sqrt{\theta}} \right|^3 \\
& = O\left(\frac{1}{\sqrt{n\theta}}\right),
\end{aligned}$$

which entails (E.6). Therefore, the desired conclusions of the theorem follow immediately, which concludes the proof of Theorem E.