

Supplementary Materials for “A Doubly-Enhanced EM Algorithm for Model-Based Tensor Clustering”

Qing Mai, Xin Zhang, Yuqing Pan and Kai Deng
Florida State University

A An example to demonstrate the effect of correlations

The following example demonstrates the effect of correlations as discussed in Section 2.3. We use the shorthand notation that, for any matrix Ω and a scalar ρ , we say that $\Omega = CS(\rho)$ if $\omega_{ii} = 1$ for all i and $\omega_{ij} = \rho$ for all $i \neq j$. The following Example S.1 reveals that ignoring the correlation structure could severely reduce the clustering and variable selection accuracy. Therefore, it is important to develop a tensor clustering method that takes into account of the correlation structure.

Example S.1. Consider a two-way tensor $\mathbf{X} \in \mathbb{R}^{5 \times 5}$ and the latent clustering label $Y \in \{1, 2\}$ generated from TNMM in (2.3), with $\pi_1^* = \pi_2^* = 0.5$ and $\mu_1^* = 0$. We consider the following two models.

(a) $\mu_{2,[1:2,1:2]}^* = 1$ and $\mu_{2,ij}^* = 0$ for all other i, j , $\Sigma_1^* = \mathbf{I}_5$, $\Sigma_2^* = CS(0.5)$. Hence, only the

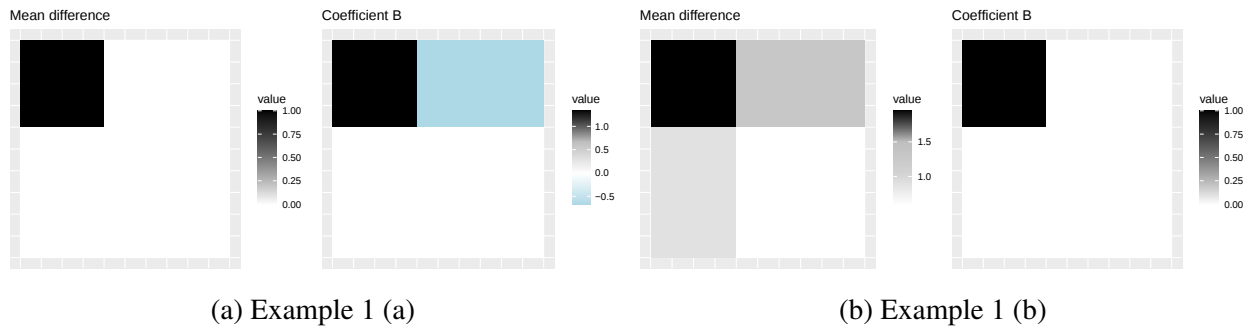


Figure S.1: Heatmaps of $\mu_2^* - \mu_1^*$ and $\mathbf{B}^*$. Each figure represents a 5×5 matrix and white spaces correspond to zero elements. Example 1 (a) shows that elements with no mean differences across clusters could still help with clustering, while Example 1 (b) shows that elements with mean differences may not help with clustering.

upper left 2×2 corner belongs to $\mathcal{A}$, but

$$\mathbf{B}^* = \begin{bmatrix} 4/3 & 4/3 & -2/3 & -2/3 & -2/3 \\ 4/3 & 4/3 & -2/3 & -2/3 & -2/3 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix}$$

has nonzero elements in the first two rows. See Figure S.1a for heatmaps of $\boldsymbol{\mu}_2^* - \boldsymbol{\mu}_1^*$ and $\mathbf{B}^*$. Hence, the $[1 : 2, 3 : 5]$ block helps clustering even though they have constant means across clusters. Moreover, on a dataset with 10000 samples, the optimal rule based on $\mathbf{B}^*$ has an error rate of 12.98%, while the independence rule based on the set $\mathcal{A}$ has an error rate of 20.95%.

(b) Set $\boldsymbol{\Sigma}_1^* = CS(0.5)$, $\boldsymbol{\Sigma}_2^* = CS(0.3)$ and

$$\boldsymbol{\mu}_2^* = \begin{bmatrix} 1.95 & 1.95 & 1.3 & 1.3 & 1.3 \\ 1.95 & 1.95 & 1.3 & 1.3 & 1.3 \\ 0.9 & 0.9 & 0.6 & 0.6 & 0.6 \\ 0.9 & 0.9 & 0.6 & 0.6 & 0.6 \\ 0.9 & 0.9 & 0.6 & 0.6 & 0.6 \end{bmatrix}.$$

Then the upper-left 2×2 corner of $\mathbf{B}^*$ has nonzero elements all equal to 1, while all other elements in $\mathbf{B}^*$ are 0. See Figure S.1b for heatmaps of $\boldsymbol{\mu}_2^* - \boldsymbol{\mu}_1^*$ and $\mathbf{B}^*$. Even though all the elements are differentially distributed across clusters, most of them do not play a role in achieving the optimal clustering error. In this case, the clustering error goes up from 6.42% to 10.85% if we ignore the correlations.

B A blockwise coordinate descent algorithm to solve (3.13)

To solve (3.13), we need the following result.

Proposition S.1. For $m = 1, \dots, M$, define $\widehat{\boldsymbol{\Sigma}}_{m,j}^{(t)}$ as the j -th column of $\widehat{\boldsymbol{\Sigma}}_m^{(t)}$. The solution of $\mathbf{B}_{\cdot\mathcal{J}} = (b_{2\mathcal{J}}, \dots, b_{K\mathcal{J}})$ to (3.13) given $b_{k\mathcal{J}'}$ for $k = 2, \dots, K$ and $\mathcal{J}' \neq \mathcal{J}$ is

$$(\widehat{b}_{2\mathcal{J}}, \dots, \widehat{b}_{K\mathcal{J}}) = \check{\mathbf{B}}_{\cdot\mathcal{J}} \left(1 - \frac{\lambda^{(t)}}{\|\check{\mathbf{B}}_{\cdot\mathcal{J}}\|} \right)_+, \quad (\text{S.1})$$

where $\check{b}_{k\mathcal{J}} = \frac{(\widehat{\mu}_{k\mathcal{J}} - \widehat{\mu}_{1\mathcal{J}}) - \llbracket \mathbf{B}_k^{\mathcal{J}}; \{\widehat{\boldsymbol{\Sigma}}_{m,j_1}^{(t)}\}^\top, \dots, \{\widehat{\boldsymbol{\Sigma}}_{m,j_M}^{(t)}\}^\top \rrbracket}{\prod_{m=1}^M \widehat{\sigma}_{m,j_m j_m}^{(t)}}$ with the elements in $\mathbf{B}_k^{\mathcal{J}}$ equals $\mathbf{B}_{k\mathcal{J}'}$ for all $\mathcal{J}' \neq \mathcal{J}$ and 0 otherwise.

Proof of Proposition S.1. The desired conclusion can be proved in the same way as Lemma 5 in Pan et al. (2019) and is thus omitted. $\square$

By Proposition S.1, we can solve (3.13) by iteratively updating $b_{k\mathcal{J}}$ while keeping all other elements fixed. Such an algorithm is summarized in Algorithm S.2.

Algorithm S.2 Algorithm to solve (3.13)

1. Input $\widehat{\Sigma}_m^{(t)}, m = 1, \dots, M$ and $\widehat{\mu}_k^{(t)}, k = 1, \dots, K$. Initialize $\widehat{\mathbf{B}}_k^{(t,0)} = 0$ for all $k = 2, \dots, K$.
 2. For $w = 1, 2, \dots$, do the following until convergence:
 For each $\mathcal{J} = (j_1, \dots, j_M)$,
 - (a) Update $\check{b}_{k,\mathcal{J}}^{(w)} \leftarrow \frac{(\widehat{\mu}_{k,\mathcal{J}}^{(t)} - \widehat{\mu}_{1,\mathcal{J}}^{(t)}) - \llbracket \mathbf{B}_k^j; \{\widehat{\Sigma}^{(t)}\}_{1,\cdot,j_1}^\top, \dots, \{\widehat{\Sigma}^{(t)}\}_{M,\cdot,j_M}^\top \rrbracket}{\prod_{m=1}^M \widehat{\sigma}_{m,j_m,j_m}};$
 - (b) Compute $\widehat{b}_{k,\mathcal{J}}^{(t,w)} \leftarrow \check{b}_{k,\mathcal{J}}^{(w)} \left(1 - \frac{\lambda^{(t)}}{\sqrt{\sum_{k=2}^K (\check{b}_{k,\mathcal{J}}^{(w)})^2}} \right)_+$ for $k = 2, \dots, K$.
 3. Output $\widehat{\mathbf{B}}_k^{(t)} = \widehat{\mathbf{B}}_k^{(t,\infty)}$ at convergence.
-

C Estimation of the number of clusters

We present a working method for the estimation of K under TNMM in high dimensions, although a thorough study along this line is out of the scope of this manuscript. Recall that, in addition to K , we also have a tuning parameter λ that controls the amount of penalty. The tuning parameter λ is selected by a BIC-type measure, defined in (3.23). We slightly modify this criterion to estimate K along with λ . Recall that $\widehat{\mathcal{D}}^\lambda = \{\mathcal{J} : \widehat{b}_{k,\mathcal{J}}^\lambda \text{ is nonzero for at least one } k\}$. We let

$$\text{BIC}(\lambda, K) = -2 \sum_{i=1}^n \log \left(\sum_{k=1}^K \widehat{\pi}_k^\lambda f_k(\mathbf{X}_i; \widehat{\boldsymbol{\theta}}_k^\lambda) \right) + \log(n) \cdot |\widehat{\mathcal{D}}^\lambda|. \quad (\text{S.2})$$

Apparently, $\text{BIC}(\lambda, K)$ is equal to $BIC(\lambda)$ in (3.23) for any fixed K . Define $\mathbb{Z}^+$ as the set of positive integers. To estimate K , we let

$$(\widehat{K}, \widehat{\lambda}) = \underset{K \in \mathbb{K}, \lambda > 0}{\operatorname{argmin}} \text{BIC}(\lambda, K), \quad (\text{S.3})$$

Model	M1	M2	M3	M4	M5	M6
Correct Selection Rate	62%	80%	77%	58%	68%	22%

Table S.1: Estimation of number of clusters. Reported are the proportions of replicates in which we correctly select K by the restricted BIC.

where $\mathcal{K} = \{K \in \mathbb{Z}^+ : \text{there are at least } 5\% \times n \text{ observations in each cluster}\}$. The set $\mathcal{K}$ contains all the K 's with reasonably balanced clusters; if a K results in very small clusters (i.e, fewer than $5\% \times n$ observations in a cluster), it is excluded. We minimize $\text{BIC}(\lambda, K)$ with the constraint that $K \in \mathcal{K}$ because otherwise $\text{BIC}(\lambda, K)$ tends to overestimate K , with many small clusters. By restricting our attention to $\mathcal{K}$, we protect ourselves from overestimation. In this sense, we refer to the criterion in (S.3) as the restricted BIC.

We test the performance of the restricted BIC on Models M1–M6 in Section 5. In Table S.1 we report the proportion of replicates in which we correctly select K . It can be seen that the restricted BIC in general produces good selection results for K . On most models the correct selection rate is around 60% or higher. Model M6 seems to be the most challenging one, with a correct selection rate of 22%. This is likely due to the fact that Model 6 has relatively large number of clusters ($K = 6$) and small within-cluster sample size.

D Simulations on tensor block model

As suggested by a referee, we also consider DEEM under a special case of the tensor block model (TBM; Wang & Zeng (2019)). In our context, TBM is constructed with a tensor $\mathcal{X} \in \mathbb{R}^{n \times p_1 \times p_2 \times \dots \times p_M}$. The tensor is clustered along all the modes. If an entry $X_{ij_1 \dots j_M}$ is from the k -th cluster on the first mode and the k_m -th cluster on the $(m + 1)$ -th mode, then we assume that

$$X_{ij_1 \dots j_M} = \mu_{k, k_1 \dots k_M}^* + \epsilon_{ij_1 \dots j_M},$$

where $\epsilon_{ij_1 \dots j_M}$ are independent statistical noises with mean zero. To best align TBM with TNMM, in what follows we assume that $\epsilon_{ij_1 \dots j_M}$ are normal, and we only consider clustering on the first mode.

We consider two TNMMs with the TBM structure. We consider $\mathbf{X}_i \in \mathbb{R}^{20 \times 20}$ for $i = 1, \dots, 150$. The 150 observations are evenly split into two clusters. We assume that $\boldsymbol{\mu}_1^* = \mathbf{0}$ and $\boldsymbol{\mu}_2^*$

	Optimal	K-means	SKM	DEEM	DTC	TBM	EM	AFPF	CHIME
TBM (a)	8.12 (0.20)	36.20 (0.87)	29.07 (1.18)	14.18 (1.11)	44.53 (0.44)	8.53 (0.29)	32.23 (0.84)	33.86 (0.82)	44.74 (0.42)
TBM (b)	7.74 (0.22)	35.70 (0.91)	23.13 (1.46)	12.04 (0.74)	43.64 (0.53)	10.74 (0.38)	31.78 (0.84)	32.82 (0.91)	44.79 (0.41)

Table S.2: Clustering results for TBMs (a) & (b). Reported are the means and standard errors (in parentheses) of clustering error rates based on 100 replicates.

is piecewise constant so that the model is a special case of TBM when we stack the observations $\mathbf{X}_i$ to form $\mathcal{X} \in \mathbb{R}^{150 \times 20 \times 20}$. We consider the following two settings for $\boldsymbol{\mu}_2^*$.

TBM (a): $\boldsymbol{\mu}_{2,[1:8,1]}^* = 1$ and all other entries are zero.

TBM (b): $\boldsymbol{\mu}_{2,[1:3,1]}^* = 1.5$, $\boldsymbol{\mu}_{2,[4:6,1]}^* = -0.5$, $\boldsymbol{\mu}_{2,[7:8,1]}^* = 0.5$ and all other entries are zero.

We continue to compare DEEM with all the competitors in Section 5.1 on these two models. The results are reported in Table S.2. Not surprisingly, the method TBM has the best performance on these two models, since it exploits all the model assumptions, such as the independence among variables and the piecewise-constant structure of $\boldsymbol{\mu}_k^*$. However, among all the other methods that do not assume TBM structure, DEEM is an obvious winner. It outperforms all the other methods by a large margin that assume independence among variables. It is slightly worse than TBM, likely because it does not use the information that the means are piece-wise constant, which greatly reduces the model complexity. Hence, it will be interesting in the future to extend DEEM to clustering problems with additional clustering structure beyond TNMM, such as TBM.

E Proofs of Lemmas 1, 2 and 3

We first list some simple propositions that will be repeatedly used.

Proposition S.2 (Kolda & Bader (2009) c.f Page 475). *For any tensor $\mathbf{C} \in \mathbb{R}^{p_1 \times \dots \times p_M}$ and matrices $\mathbf{A}_m \in \mathbb{R}^{p_m \times p_m}$, we have that*

$$\llbracket \mathbf{C}; \mathbf{A}_1, \dots, \mathbf{A}_M \rrbracket_{(m)} = \mathbf{A}_m \mathbf{C}_{(m)} (\otimes_{j=M, j \neq m}^{j=1} \mathbf{A}_j). \quad (\text{S.4})$$

Proposition S.3. *For matrices $\boldsymbol{\Omega}_m, \tilde{\boldsymbol{\Omega}}_m, m = 1, \dots, M$, we have*

$$\otimes_{m=M}^{m=1} \tilde{\boldsymbol{\Omega}}_m - \otimes_{m=M}^{m=1} \boldsymbol{\Omega}_m = \sum_{j=1}^M \{ \otimes_{m=M}^{m=j+1} \tilde{\boldsymbol{\Omega}}_m \} \otimes \{ \tilde{\boldsymbol{\Omega}}_j - \boldsymbol{\Omega}_j \} \otimes \{ \otimes_{m=j-1}^{m=1} \boldsymbol{\Omega}_m \}. \quad (\text{S.5})$$

Proof of Proposition S.3. The desired conclusion can be proved by straightforward calculation. $\square$

Proof of Lemma 1. The Q-function in (3.3) is equal to

$$Q_n(\boldsymbol{\theta} \mid \tilde{\boldsymbol{\theta}}^{(t)}) = \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} \left\{ \log \pi_k - \frac{1}{2} \sum_{m=1}^M q_m \log |\boldsymbol{\Sigma}_m| - \frac{1}{2} \langle \llbracket \mathbf{X}_i - \boldsymbol{\mu}_k; \boldsymbol{\Sigma}_1^{-1}, \dots, \boldsymbol{\Sigma}_M^{-1} \rrbracket, \mathbf{X}_i - \boldsymbol{\mu}_k \rangle \right\}. \quad (\text{S.6})$$

By Proposition S.2 we have

$$\llbracket \mathbf{X}_i - \boldsymbol{\mu}_k; \boldsymbol{\Sigma}_1^{-1}, \dots, \boldsymbol{\Sigma}_M^{-1} \rrbracket_{(m)} = \boldsymbol{\Sigma}_m^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1}. \quad (\text{S.7})$$

Hence,

$$\langle \llbracket \mathbf{X}_i - \boldsymbol{\mu}_k; \boldsymbol{\Sigma}_1^{-1}, \dots, \boldsymbol{\Sigma}_M^{-1} \rrbracket, \mathbf{X}_i - \boldsymbol{\mu}_k \rangle \quad (\text{S.8})$$

$$= \langle \llbracket \mathbf{X}_i - \boldsymbol{\mu}_k; \boldsymbol{\Sigma}_1^{-1}, \dots, \boldsymbol{\Sigma}_M^{-1} \rrbracket_{(m)}, (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \rangle \quad (\text{S.9})$$

$$= \text{tr}(\boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top \boldsymbol{\Sigma}_m^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}) \quad (\text{S.10})$$

$$= \text{tr}(\boldsymbol{\Sigma}_m^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top). \quad (\text{S.11})$$

Consequently, when we fix all other parameters except for $\boldsymbol{\Sigma}_m$, the Q-function reduces to

$$\begin{aligned} & Q_n^m(\boldsymbol{\Sigma}_m \mid \tilde{\boldsymbol{\theta}}^{(t)}) \\ &= \text{Const} - \frac{q_m}{2} \log |\boldsymbol{\Sigma}_m| \cdot \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} - \frac{1}{2} \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} \text{tr}\{\boldsymbol{\Sigma}_m^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top\} \\ &= \text{Const} - \frac{nq_m}{2} \log |\boldsymbol{\Sigma}_m| - \frac{1}{2} \text{tr}[\boldsymbol{\Sigma}_m^{-1} \{ \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top \}], \end{aligned}$$

where we use the fact that $\sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} = 1$ for any i .

By Petersen & Pedersen (2008), we have

$$\frac{\partial Q_n^m(\boldsymbol{\Sigma}_m \mid \tilde{\boldsymbol{\theta}}^{(t)})}{\partial \boldsymbol{\Sigma}_m} = -\frac{nq_m}{2} \cdot \frac{1}{|\boldsymbol{\Sigma}_m|} \cdot |\boldsymbol{\Sigma}_m| \cdot \boldsymbol{\Sigma}_m - \frac{1}{2} \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top. \quad (\text{S.12})$$

Set $\frac{\partial Q_n^m(\boldsymbol{\Sigma}_m \mid \tilde{\boldsymbol{\theta}}^{(t)})}{\partial \boldsymbol{\Sigma}_m} = 0$ and we have

$$\boldsymbol{\Sigma}_m = \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^K \tilde{\xi}_{ik}^{(t)} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1} (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)}^\top. \quad (\text{S.13})$$

By noting that $\mathbf{W}_{ik} = (\mathbf{X}_i - \boldsymbol{\mu}_k)_{(m)} \boldsymbol{\Sigma}_{-m}^{-1/2}$ we have the desired conclusion. $\square$

Proof of Lemma 2. Recall that

$$\xi_{ik} = \Pr(Y_i = k \mid \mathbf{X}_i, \boldsymbol{\theta}^*) = \frac{\pi_k^* f_k(\mathbf{X}_i; \boldsymbol{\theta}^*)}{\sum_{j=1}^K \pi_j^* f_j(\mathbf{X}_i; \boldsymbol{\theta}^*)}, \quad (\text{S.14})$$

where the density function f_k is given in (3.2). It follows that

$$\begin{aligned} \frac{\xi_{ik}}{\xi_{i1}} &= \frac{\pi_k^* f_k(\mathbf{X}_i; \boldsymbol{\theta}^*)}{\pi_1^* f_1(\mathbf{X}_i; \boldsymbol{\theta}^*)} \\ &= \frac{\pi_k^*}{\pi_1^*} \exp\left\{-\frac{1}{2}(\langle [\mathbf{X}_i - \boldsymbol{\mu}_k^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_k^* \rangle \right. \\ &\quad \left. - \langle [\mathbf{X}_i - \boldsymbol{\mu}_1^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_1^* \rangle)\right\}. \end{aligned}$$

Note that $\langle [\mathbf{X}_i - \boldsymbol{\mu}_k^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_k^* \rangle = \text{vec}^\top(\mathbf{X}_i - \boldsymbol{\mu}_k^*)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\mathbf{X}_i - \boldsymbol{\mu}_k^*)$, $\langle [\mathbf{X}_i - \boldsymbol{\mu}_1^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_1^* \rangle = \text{vec}^\top(\mathbf{X}_i - \boldsymbol{\mu}_1^*)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\mathbf{X}_i - \boldsymbol{\mu}_1^*)$, where $\boldsymbol{\Sigma}^* = \otimes_{m=1}^M \boldsymbol{\Sigma}_m^*$.

It follows that

$$\begin{aligned} &\langle [\mathbf{X}_i - \boldsymbol{\mu}_k^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_k^* \rangle - \langle [\mathbf{X}_i - \boldsymbol{\mu}_1^*; (\boldsymbol{\Sigma}_1^*)^{-1}, \dots, (\boldsymbol{\Sigma}_M^*)^{-1}], \mathbf{X}_i - \boldsymbol{\mu}_1^* \rangle \\ &= \text{vec}^\top(\mathbf{X}_i)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\mathbf{X}_i) - 2\text{vec}^\top(\mathbf{X}_i)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\boldsymbol{\mu}_k^*) + \text{vec}^\top(\boldsymbol{\mu}_k^*)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\boldsymbol{\mu}_k^*) \\ &\quad - \text{vec}^\top(\mathbf{X}_i)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\mathbf{X}_i) + 2\text{vec}^\top(\mathbf{X}_i)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\boldsymbol{\mu}_1^*) - \text{vec}^\top(\boldsymbol{\mu}_1^*)(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\boldsymbol{\mu}_1^*) \\ &= -2\text{vec}^\top\left\{\mathbf{X}_i - \frac{1}{2}(\boldsymbol{\mu}_k^* + \boldsymbol{\mu}_1^*)\right\}(\boldsymbol{\Sigma}^*)^{-1} \text{vec}(\boldsymbol{\mu}_k^* - \boldsymbol{\mu}_1^*) \\ &= -2\langle \mathbf{X}_i - \frac{1}{2}(\boldsymbol{\mu}_k^* + \boldsymbol{\mu}_1^*), \mathbf{B}_k^* \rangle. \end{aligned}$$

Hence,

$$\frac{\xi_{ik}}{\xi_{i1}} = \frac{\pi_k^*}{\pi_1^*} \exp\left\{\langle \mathbf{X}_i - \frac{1}{2}(\boldsymbol{\mu}_k^* + \boldsymbol{\mu}_1^*), \mathbf{B}_k^* \rangle\right\}, \quad k > 1. \quad (\text{S.15})$$

Solve the equation system in (S.15) along with $\sum_{k=1}^K \xi_{ik} = 1$ and we have the desired conclusion. $\square$

To prove Lemma 3, we need the following proposition.

Proposition S.4. [Pan et al. (2019) c.f. Lemma 2] *If an M -way random tensor $\mathbf{W} \sim TN(0; \boldsymbol{\Omega}_1, \dots, \boldsymbol{\Omega}_M)$, then $\mathbb{E}\{\mathbf{W}_{(j)} \mathbf{W}_{(j)}^\top\} = \boldsymbol{\Omega}_j \cdot \prod_{l \neq j} \text{tr}(\boldsymbol{\Omega}_l)$, where $\mathbf{W}_{(j)}$ is the mode- j matricization of $\mathbf{W}$.*

Now we prove Lemma 3.

Proof of Lemma 3. For any k , we have that

$$\frac{1}{q_m} E[\xi_{ik} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.16})$$

$$= \frac{1}{q_m} E[\Pr(Y_i = k \mid \mathbf{X}_i, \boldsymbol{\theta}^*) (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.17})$$

$$= \frac{1}{q_m} E[E[\mathbf{1}(Y_i = k) \mid \mathbf{X}_i, \boldsymbol{\theta}^*] (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.18})$$

$$= \frac{1}{q_m} E[E[\mathbf{1}(Y_i = k) (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top \mid \mathbf{X}_i, \boldsymbol{\theta}^*]] \quad (\text{S.19})$$

$$= \frac{1}{q_m} E_{\boldsymbol{\theta}^*}[\mathbf{1}(Y = k) (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.20})$$

$$\propto \pi_k^* \boldsymbol{\Sigma}_m^*, \quad (\text{S.21})$$

where in the last equation we use Proposition S.4. Hence, $\frac{1}{q_m} E[\sum_{k=1}^K \xi_{ik} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)} \{\mathbf{X}_i - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \propto \boldsymbol{\Sigma}_m^*$. $\square$

F Proofs of Theorems 1 and 2

F.1 Notations and propositions

Some of the following notations are already defined in the main body of the paper. For better readability, we repeat them here. For two numbers $x, y \in \mathbb{R}$, we denote $x \wedge y = \min\{x, y\}$, $x \vee y = \max\{x, y\}$. For a vector $\mathbf{x} \in \mathbb{R}^p$, we denote

$$\|\mathbf{x}\|_0 = \sum_{j=1}^p \mathbf{1}(x_j \neq 0), \|\mathbf{x}\|_1 = \sum_{j=1}^p |x_j|, \|\mathbf{x}\|_2 = \sqrt{\sum_{j=1}^p x_j^2}, \|\mathbf{x}\|_\infty = \max_{j=1}^p |x_j|. \quad (\text{S.22})$$

For a matrix $\mathbf{X} \in \mathbb{R}^{p \times q}$, we let $\|\mathbf{X}\|_1 = \max_{i=1}^p \{\sum_{j=1}^q |x_{ij}|\}$ and $\|\mathbf{X}\|_{\max} = \max_{i,j} |x_{ij}|$. For any positive integer $j \leq q$, $\mathbf{X}_{\cdot j}$ is the j -th column of $\mathbf{X}$. If $\mathbf{X}$ is symmetric, $\lambda_{\min}(\mathbf{X})$ is its smallest eigenvalue, and $\lambda_{\max}(\mathbf{X})$ is its largest eigenvalue. For a vector $\mathbf{u} \in \mathbb{R}^q$ and a matrix $\mathbf{A} \in \mathbb{R}^{p \times q}$, we denote

$$\|\mathbf{u}\|_{2,s} = \sup_{\|\mathbf{x}\|_2=1, \mathbf{x} \in \Gamma(s)} |\mathbf{u}^\top \mathbf{x}|, \quad \|\mathbf{A}\|_{2,s} = \sup_{\|\mathbf{x}\|_2=1, \mathbf{x} \in \Gamma(s)} \|\mathbf{A}\mathbf{x}\|_2. \quad (\text{S.23})$$

For an M -way tensor $\mathbf{X} \in \mathbb{R}^{p_1 \times p_2 \times \dots \times p_M}$, we denote $\|\mathbf{X}\| = \sqrt{\sum_{\mathcal{J}} X_{\mathcal{J}}^2}$ and $\|\mathbf{X}\|_0 = \sum_{\mathcal{J}} \mathbf{1}(X_{\mathcal{J}} \neq 0)$. We also need the following two propositions. We include the proof for the first one for completeness.

Proposition S.5. For matrices $\mathbf{A} \in \mathbb{R}^{p \times p}$, $\mathbf{B} \in \mathbb{R}^{q \times q}$, we have $\|\mathbf{A} \otimes \mathbf{B}\|_1 \leq \|\mathbf{A}\|_1 \cdot \|\mathbf{B}\|_1$.

Proof of Proposition S.5. For any $j = 1, \dots, pq$, there exist $l \in \{1, \dots, p\}, t \in \{1, \dots, q\}$ such that $(\mathbf{A} \otimes \mathbf{B})_{.j} = (a_{1l}\mathbf{B}_{.t}^\top, a_{2l}\mathbf{B}_{.t}^\top, \dots, a_{pl}\mathbf{B}_{.t}^\top)^\top$. It follows that

$$\begin{aligned} \|\mathbf{A} \otimes \mathbf{B}\|_1 &= \max_j \|(\mathbf{A} \otimes \mathbf{B})_{.j}\|_1 \leq \max_{l,t} \{|a_{1l}|\|\mathbf{B}_{.t}\|_1 + |a_{2l}|\|\mathbf{B}_{.t}\|_1 + \dots + |a_{pl}|\|\mathbf{B}_{.t}\|_1\} \\ &\leq \max_l \left(\sum_v |a_{vl}| \right) \max_t \|\mathbf{B}_{.t}\|_1 \leq \|\mathbf{A}\|_1 \cdot \|\mathbf{B}\|_1. \end{aligned}$$

□

Proposition S.6 (Lyu et al. (2019) c.f Lemma S.1). Consider random matrices $\mathbf{W}_1, \dots, \mathbf{W}_n \in \mathbb{R}^{p \times q}$ iid from $TN(0; \mathbf{\Omega}_R, \mathbf{\Omega}_L)$. If the eigenvalues of $\mathbf{\Omega}_R, \mathbf{\Omega}_L$ are universally bounded below and above, then for any symmetric and positive definite matrix $\mathbf{\Omega} \in \mathbb{R}^{q \times q}$, we have

$$\left\| \frac{1}{nq} \sum_{i=1}^n \mathbf{W}_i \mathbf{\Omega} (\mathbf{W}_i)^\top - \frac{1}{q} E \mathbf{W}_i \mathbf{\Omega} (\mathbf{W}_i)^\top \right\|_{\max} = O_P\left(\sqrt{\frac{\log q}{np}}\right). \quad (\text{S.24})$$

F.2 Properties for a generic penalized EM algorithm

We introduce a generic penalized EM algorithm (PEM) that includes DEEM as a special case. Then we present a proposition to guarantee its consistency under proper conditions. The proposition is later used to prove Theorems 1 & 2. The proofs for our proposition are closely related to those in Cai et al. (2019), but we reach a more general conclusion. We also list some properties related to the initialization algorithm that will be used to prove Lemmas S.8 & 4.

Consider iid observations $\{\mathbf{U}_i, Y_i\}_{i=1}^n$, where $\mathbf{U}_i \in \mathbb{R}^p$ and $Y_i \in \{1, 2\}$ is unobservable. We make the following model assumption for $\{\mathbf{U}_i, Y_i\}$:

$$(A1') \quad \Pr(Y_i = k) = \pi_k^*, \mathbf{U}_i \mid Y_i = k \sim N(\phi_k^*, \mathbf{\Psi}^*), \text{ where } 0 < \pi_k^* < 1, \phi_k^* \in \mathbb{R}^p, \mathbf{\Psi}^* \in \mathbb{R}^{p \times p}.$$

Note that Model (A1') is equivalent to Model (A1) in Section G. We consider a generic PEM to recover Y_i . Denote $\boldsymbol{\theta}^{\text{vec}} = (\pi_k, \phi_k, k = 1, 2; \mathbf{\Psi})$, $\xi_{ik}(\boldsymbol{\theta}^{\text{vec}}) = \Pr(Y_i = k \mid \mathbf{U}_i, \boldsymbol{\theta}^{\text{vec}})$ and $\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}) \in \mathbb{R}^{n \times 2}$ with the i -th row being $(\xi_{i1}(\boldsymbol{\theta}^{\text{vec}}), \xi_{i2}(\boldsymbol{\theta}^{\text{vec}}))$. Given $\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}})$, we estimate the parameters in the M-step by

$$\hat{\pi}_k(\boldsymbol{\theta}^{\text{vec}}) = g_{\pi_k}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}, \hat{\phi}_k(\boldsymbol{\theta}^{\text{vec}}) = g_{\phi_k}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}, \hat{\mathbf{\Psi}}(\boldsymbol{\theta}^{\text{vec}}) = g_{\mathbf{\Psi}}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}, \quad (\text{S.25})$$

where

$$g_{\pi_k} : \mathbb{R}^{n \times 2} \times \mathbb{R}^{n \times p} \mapsto \mathbb{R}; \quad g_{\phi_k} : \mathbb{R}^{n \times 2} \times \mathbb{R}^{n \times p} \mapsto \mathbb{R}^p; \quad g_{\Psi} : \mathbb{R}^{n \times 2} \times \mathbb{R}^{n \times p} \mapsto \mathbb{R}^{p \times p} \quad (\text{S.26})$$

are user-specified functions. We further use the shorthand notation

$$\begin{aligned} g\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\} &= \{\hat{\pi}_k(\boldsymbol{\theta}^{\text{vec}}), \hat{\phi}_k(\boldsymbol{\theta}^{\text{vec}}), k = 1, 2; \hat{\Psi}(\boldsymbol{\theta}^{\text{vec}})\} \\ &= [g_{\pi_k}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}, g_{\phi_k}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}, k = 1, 2; g_{\Psi}\{\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}\}]. \end{aligned}$$

On the other hand, in the E-step, given the value of $\boldsymbol{\theta}^{\text{vec}} = (\pi_k, \phi_k, k = 1, 2; \Psi)$, $\xi_{ik}(\boldsymbol{\theta}^{\text{vec}})$ is estimated by projecting $\mathbf{U}_i$ to $\boldsymbol{\beta} = \Psi^{-1}(\phi_2 - \phi_1)$. We obtain a sparse estimate for $\boldsymbol{\beta}$ through a penalized optimization problem. Such a PEM algorithm is summarized in Algorithm S.3. Note that PEM is generic since we do not restrict g to a specific form.

To derive the theoretical properties of the PEM algorithm, consider the parameter space

$$\begin{aligned} \Theta^{\text{vec}}(s, c_{\pi}, C_0, C_b, \Delta_0) &= \{\boldsymbol{\theta}^{\text{vec}} = (\pi_k, \phi_k, k = 1, 2; \Psi) : \\ &\quad \phi_k \in \mathbb{R}^p, \Psi \in \mathbb{R}^p, C_0^{-1} \leq \lambda_{\min}(\Psi) \leq \lambda_{\max}(\Psi) \leq C_0, \\ &\quad \|\boldsymbol{\beta}\|_0 \leq s, \|\boldsymbol{\beta}\|_1 \leq C_b, c_{\pi} \leq \pi_1, \pi_2 \leq 1 - c_{\pi}, \Delta > \Delta_0\}, \end{aligned}$$

where s is a positive integer, $0 < c_{\pi} < 1$, $C_0, C_b > 0$ and

$$\Delta = \sqrt{(\phi_1 - \phi_2)^{\top}(\Psi)^{-1}(\phi_1 - \phi_2)}.$$

Recall $\Gamma(s)$ defined in (4.2). For the true parameter $\boldsymbol{\theta}^{\text{vec}*} \in \Theta^{\text{vec}}$, we define

$\Delta^* = \sqrt{(\phi_1^* - \phi_2^*)^{\top}(\Psi^*)^{-1}(\phi_1^* - \phi_2^*)}$ and the contraction basin of $\boldsymbol{\theta}^{\text{vec}*}$ to be

$$\begin{aligned} \mathcal{B}_{\text{con}}(\boldsymbol{\theta}^{\text{vec}*}; a_{\pi}, a_{\Delta}, a_b, s) &= \{\boldsymbol{\theta}^{\text{vec}} = (\pi_k, \phi_k, k = 1, 2; \Psi) : \\ &\quad \phi_k \in \mathbb{R}^p, \Psi \in \mathbb{R}^{p \times p}, \Psi \succ 0, a_{\pi} < \pi_k < 1 - a_{\pi}, \\ &\quad (1 - a_{\Delta})(\Delta^*)^2 < |\delta_k(\boldsymbol{\beta})|, \sigma^2(\boldsymbol{\beta}) < (1 + a_{\Delta})(\Delta^*)^2, \\ &\quad \boldsymbol{\beta} - \boldsymbol{\beta}^* \in \Gamma(s), \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_1 \leq a_b \Delta^*, \|\phi_k\|_{2,s} \leq a_b \Delta^*\}, \end{aligned}$$

where $a_{\pi}, a_{\Delta}, a_b > 0$, $a_{\pi} \leq c_{\pi} < 1$, $\delta_k(\boldsymbol{\beta}) = \boldsymbol{\beta}^{\top}(\phi_k^* - \frac{\phi_1 + \phi_2}{2})$ and $\sigma^2(\boldsymbol{\beta}) = \boldsymbol{\beta}^{\top} \Psi^* \boldsymbol{\beta}$.

Now we introduce a few assumptions under which PEM is consistent. Define the distances:

$$d_2(\boldsymbol{\theta}^{\text{vec}}, \tilde{\boldsymbol{\theta}}^{\text{vec}}) = (\bigvee_{k=1}^2 |\pi_k - \tilde{\pi}_k|) \vee (\bigvee_{k=1}^2 \|\phi_k - \tilde{\phi}_k\|_2) \vee \|(\Psi - \tilde{\Psi})\tilde{\boldsymbol{\beta}}\|_2, \quad (\text{S.32})$$

$$d_{2,s}(\boldsymbol{\theta}^{\text{vec}}, \tilde{\boldsymbol{\theta}}^{\text{vec}}) = (\bigvee_{k=1}^2 |\pi_k - \tilde{\pi}_k|) \vee (\bigvee_{k=1}^2 \|\phi_k - \tilde{\phi}_k\|_{2,s}) \vee \|(\Psi - \tilde{\Psi})\tilde{\boldsymbol{\beta}}\|_{2,s}, \quad (\text{S.33})$$

$$d_{2,s}^{(1)}(\boldsymbol{\theta}^{\text{vec}}, \tilde{\boldsymbol{\theta}}^{\text{vec}}) = (\bigvee_{k=1}^2 |\pi_k - \tilde{\pi}_k|) \vee (\bigvee_{k=1}^2 \|\phi_k - \tilde{\phi}_k\|_{2,s}). \quad (\text{S.34})$$

Algorithm S.3 A generic PEM algorithm

1. Initialize $\theta^{\text{vec}(0)}$, $0 < \kappa < 1$. Set

$$\hat{\beta} = \arg \min_{\mathbb{R}^p} \left\{ \frac{1}{2} \beta^\top \hat{\Psi} \beta - \beta^\top (\hat{\phi}_2^{(0)} - \hat{\phi}_1^{(0)}) + \lambda^{(0)} \|\beta\|_1 \right\}, \quad (\text{S.27})$$

where $\lambda_n^{(0)} = C_d(|\hat{\pi}_2| \vee \|\hat{\phi}_1^{(0)} - \hat{\phi}_2^{(0)}\|_{2,s} \vee \|\hat{\Psi}^{(0)}\|_{2,s}) + C_\lambda \sqrt{\log p/n}$.

2. For $t = 0, \dots$, do

(a) Update

$$(\hat{\pi}_k^{(t+1)}, \hat{\phi}_k^{(t+1)}, k = 1, 2; \hat{\Psi}^{(t+1)}) = g(\xi^{(t)}, \mathbf{U}); \quad (\text{S.28})$$

(b) Update $\hat{\xi}_{i1}^{(t+1)}, \hat{\xi}_{i1}^{(t+1)}$ by

$$\hat{\xi}_{i1}^{(t+1)} = \frac{\pi_1^{(t+1)}}{\hat{\pi}_1^{(t+1)} + \hat{\pi}_2^{(t+1)} \exp[\{\mathbf{U}_i - (\hat{\phi}_1^{(t+1)} + \hat{\phi}_2^{(t+1)})/2\}^\top \hat{\beta}^{(t+1)}]}, \quad \hat{\xi}_{i2}^{(t+1)} = 1 - \hat{\xi}_{i1}^{(t+1)}, \quad (\text{S.29})$$

where

$$\hat{\beta}^{(t+1)} = \arg \min_{\beta \in \mathbb{R}^p} \{ \beta^\top \hat{\Psi}^{(t+1)} \beta - 2(\hat{\phi}_2^{(t+1)} - \hat{\phi}_1^{(t+1)})^\top \beta + \lambda^{(t+1)} \|\beta\|_1 \}, \quad (\text{S.30})$$

with

$$\lambda^{(t+1)} = \kappa \lambda^{(t)} + C_\lambda \sqrt{\frac{\log p}{n}}. \quad (\text{S.31})$$

3. Output $\hat{\xi}_{ik}$ at convergence.

Ideally, we hope that $g(\xi, \mathbf{U})$ converges in probability as $n \rightarrow \infty$. We formally write this assumption as follows:

(A2) There exists $g^{\text{lim}}(\theta^{\text{vec}}) = \{\pi_k(\theta^{\text{vec}}), \phi_k(\theta^{\text{vec}}), k = 1, 2; \Psi(\theta^{\text{vec}})\}$ such that $|g_{\pi_k}\{\xi(\theta), \mathbf{U}\} - \pi_k(\theta^{\text{vec}})|, \|g_{\phi_k}\{\xi(\theta), \mathbf{U}\} - \phi_k(\theta^{\text{vec}})\|_{\max}, \|g_{\Psi}\{\xi(\theta^{\text{vec}}), \mathbf{U}\} - \Psi(\theta^{\text{vec}})\|_{\max} \rightarrow 0$ in probability.

Under (A2), we further make the following assumptions concerning g and g^{lim} . For simplicity, we write the value of our estimate at $(t+1)$ -th iteration, i.e., $g(\xi(\hat{\theta}^{\text{vec}(t+1)}), \mathbf{U})$, as $\hat{\theta}^{\text{vec}(t+1)}$, and its limit $g^{\text{lim}}(\hat{\theta}^{\text{vec}(t+1)})$ as $\theta^{\text{vec}(t+1)}$.

(A3) With a probability tending to 1, we have

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} d_{2,s}(\hat{\boldsymbol{\theta}}^{\text{vec}(t+1)}, \boldsymbol{\theta}^{\text{vec}(t+1)}) \leq C_{con} \sqrt{\frac{s \log p}{n}}, \quad (\text{S.35})$$

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} \|(\hat{\boldsymbol{\Psi}}^{(t+1)} - \boldsymbol{\Psi}^{(t+1)}) \boldsymbol{\beta}^{\text{vec}*}\|_{\infty} \leq C_{con} \sqrt{\frac{\log p}{n}}. \quad (\text{S.36})$$

(A4) There exists a sufficiently large Δ_0 and $0 < \kappa_0 < \frac{1}{2\sqrt{16C_0}}$ that does not depend on n, p , such that for any $\boldsymbol{\theta}^{\text{vec}} \in \mathcal{B}_{con}$ we have

$$d_{2,s}[g^{\text{lim}}\{\boldsymbol{\theta}\}, \boldsymbol{\theta}^{\text{vec}*}] \leq \kappa_0 \cdot (d_{2,s}(\boldsymbol{\theta}^{\text{vec}}, \boldsymbol{\theta}^{\text{vec}*}) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2), \quad (\text{S.37})$$

where g^{lim} is defined as in (A2).

(A5) For the initial value $\hat{\boldsymbol{\theta}}^{\text{vec}(0)}$, we have

$$d_{2,s}(\hat{\boldsymbol{\theta}}^{\text{vec}(0)}, \boldsymbol{\theta}^{\text{vec}*}) \vee \|\hat{\boldsymbol{\beta}}^{(0)} - \boldsymbol{\beta}^*\|_2 \leq r\Delta^*, \quad \hat{\boldsymbol{\beta}}^{(0)} - \boldsymbol{\beta}^* \in \Gamma(s), \quad (\text{S.38})$$

with $r < \frac{|a_{\pi} - c_{\pi}|}{\Delta} \wedge \frac{\sqrt{9C_0 + 16c_{\Delta}} - \sqrt{9C_0}}{4} \wedge \frac{a_{\Delta}}{C_0} \wedge \frac{a_b}{5\sqrt{s}}$.

Proposition S.7. Suppose that Assumptions (A1')–(A5) hold. If $\sqrt{s \log p/n} \ll r$, there exist $C_d, C_{\lambda} > 0, 0 < \kappa < 1/2$ such that with a probability greater than $1 - O(p^{-1})$, we have

$$\|\hat{\boldsymbol{\beta}}^{(t)} - \boldsymbol{\beta}^*\|_2 \lesssim \kappa^t d_{2,s}(\hat{\boldsymbol{\theta}}^{\text{vec}(0)}, \boldsymbol{\theta}^{\text{vec}*}) + \sqrt{\frac{s \log p}{n}}. \quad (\text{S.39})$$

Moreover, for $t \gtrsim (-\log(\kappa))^{-1} \log n \cdot d_{2,s}(\hat{\boldsymbol{\theta}}^{\text{vec}(0)}, \boldsymbol{\theta}^{\text{vec}*})$, then

$$\|\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*\|_2 \lesssim \sqrt{\frac{s \log p}{n}}, \quad (\text{S.40})$$

$$R(\text{PEM}) - R(\text{Opt}) \lesssim \frac{s \log p}{n}. \quad (\text{S.41})$$

Proof of Proposition S.7. Proposition S.7 can be proved following the same line of Theorems 3.1 & 3.2 in Cai et al. (2019) by replacing their Lemma 3.1 with Assumption (A4) and their Lemma 3.2 with Assumption (A3). $\square$

The following two propositions are important for our proof as well.

Proposition S.8. [Cai et al. (2019) c.f Lemmas 3.1 & 3.2] Consider data drawn from the model in (A1'). In the PEM, let $g_{\omega_k}(\boldsymbol{\xi}, \mathbf{U}) = \frac{1}{n} \sum_{i=1}^n \xi_{ik}$ and $g_{\phi_k}(\boldsymbol{\xi}, \mathbf{U}) = \frac{\sum_{i=1}^n \xi_{ik} \mathbf{U}_i}{\sum_{i=1}^n \xi_{ik}}$. We have the following conclusions.

1. For any $\boldsymbol{\theta}^{\text{vec}*} \in \boldsymbol{\Theta}^{\text{vec}}$, there exists a constant $C_{\text{con}} > 0$ such that with probability at least $1 - O(p^{-1})$, we have

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{\text{con}}} d_{2,s}^{(1)}(g(\boldsymbol{\xi}(\boldsymbol{\theta}^{\text{vec}}), \mathbf{U}), g^{\text{lim}}(\boldsymbol{\theta}^{\text{vec}})) \leq C_{\text{con}} \sqrt{\frac{s \log p}{n}}. \quad (\text{S.42})$$

2. Under Assumption (A5), there exists $\kappa_0^{(1)}(\Delta)$ such that for any $\boldsymbol{\theta}^{\text{vec}} \in \mathcal{B}_{\text{con}}$,

$$d_2^{(1)}[g^{\text{lim}}(\boldsymbol{\theta}^{\text{vec}}), \boldsymbol{\theta}^{\text{vec}*}] \leq \kappa_0^{(1)}(\Delta) \cdot \{d_{2,s}^{(1)}(\boldsymbol{\theta}^{\text{vec}}, \boldsymbol{\theta}^{\text{vec}*}) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2\}, \quad (\text{S.43})$$

where $\kappa^{(1)}(\Delta) = o(1/\Delta^2)$ as $\Delta \rightarrow \infty$.

Proposition S.9 (Cai et al. (2019) c.f Proof of Theorem 3.3). *Define*

$$\tilde{\mathcal{A}}_s = \{\mathbf{u} \in \{0, 1\}^p : \mathbf{u}^\top \mathbf{e}_1 = 0, \|\mathbf{u}\|_0 = s\}, \quad (\text{S.44})$$

where $e_{11} = 1$ and $e_{ij} = 0$ for $j > 1$. There exists

$$\boldsymbol{\theta}_i^{\text{vec}} = \{\pi_1 = \pi_2 = \frac{1}{2}, \phi_1 = \epsilon \mathbf{u}_i + \eta \mathbf{e}_i, \phi_2 = -\phi_1, \boldsymbol{\Sigma} = \sigma^2 \mathbf{I}_p\}, i = 1, \dots, N, \quad (\text{S.45})$$

where $\mathbf{u} \in \tilde{\mathcal{A}}_s, \epsilon = \sigma \sqrt{\log p/n}, \sigma^2 = O(1), \eta = O(1), c_0 \leq \sigma^2 \leq C_0$ and $(\phi_2 - \phi_1)^\top \boldsymbol{\Sigma}^{-1}(\phi_2 - \phi_1) \geq c_L$ for some constant $c_L > 0$, such that

$$\inf_{\hat{Y}} \sup_{i=1, \dots, N} \{R(\hat{Y} \mid \boldsymbol{\theta}_i^{\text{vec}}) - R(\text{Opt} \mid \boldsymbol{\theta}_i^{\text{vec}})\} \succeq \frac{\log p}{n}. \quad (\text{S.46})$$

We further consider the following initialization condition.

(IC) For some permutation $\Pi : \{1, 2\} \mapsto \{1, 2\}$,

$$\max_{k=1,2} \{\|\hat{\boldsymbol{\phi}}_k^{(0)} - \boldsymbol{\phi}_{\Pi(k)}\|_\infty\} \preceq \frac{1}{s}, \quad \|\hat{\boldsymbol{\Psi}}^{(0)} - \boldsymbol{\Psi}\|_{\max} \preceq \frac{1}{s}. \quad (\text{S.47})$$

Proposition S.10 (Cai et al. (2019) c.f Lemma 3.3). *If the initial estimates $\hat{\boldsymbol{\phi}}_k^{(0)}, \hat{\boldsymbol{\Psi}}^{(0)}$ satisfy Condition (IC), then the resulting $\hat{\boldsymbol{\theta}}^{\text{vec}(0)}$ produced by PEM satisfies Condition (A5).*

F.3 Proofs for Theorems 1 & 2

For simplicity, we present the proofs under the assumption that $\text{diag}(\boldsymbol{\Sigma}_m^*) = \mathbf{I}$ so that we do not need to rescale our covariance estimates. This assumption can be straightforwardly relaxed with more tedious calculation. To simplify the notation, we use $\mathbf{X}_{i(m)}$ to denote the matricization $(\mathbf{X}_i)_{(m)}$.

Lemma S.5. Consider $\mathbf{W}_1, \dots, \mathbf{W}_n \in \mathbb{R}^{p_1 \times \dots \times p_M}$ iid from $TN(0; \boldsymbol{\Omega}_1, \dots, \boldsymbol{\Omega}_M)$, where diagonal elements of $\boldsymbol{\Omega}_m, m = 1, \dots, M$ are all ones, and $\lambda_{\max}(\boldsymbol{\Omega}_m)$ are universally bounded below and above. Define

$$\hat{\boldsymbol{\Omega}}_m = \frac{1}{np_{-m}} \sum_{i=1}^n (\mathbf{W}_i)_{(m)} (\mathbf{W}_i)_{(m)}^\top. \quad (\text{S.48})$$

Then there exists a constant C such that with a probability tending to 1, we have

$$\|\hat{\boldsymbol{\Omega}}_m - \boldsymbol{\Omega}_m\|_{\max} \leq C \sqrt{\frac{\log p_m}{nq_m}}. \quad (\text{S.49})$$

Proof of Lemma S.5. Because $\mathbf{W}_i \sim TN(0; \boldsymbol{\Omega}_1, \dots, \boldsymbol{\Omega}_M)$, there exists $\mathbf{Z} \sim TN(0; \mathbf{I}_{p_1}, \dots, \mathbf{I}_{p_M})$ such that

$$\mathbf{W}_i = [\mathbf{Z}; \boldsymbol{\Omega}_1^{1/2}, \dots, \boldsymbol{\Omega}_M^{1/2}]. \quad (\text{S.50})$$

It follows that $\mathbf{W}_{i(m)} = [\mathbf{Z}_{(m)}; \boldsymbol{\Omega}_m, \otimes_{j=M, j \neq m}^{j=1} \boldsymbol{\Omega}_M]$ and hence

$$\mathbf{W}_{i(m)} \sim TN(0; \boldsymbol{\Omega}_m, \otimes_{j=M, j \neq m}^{j=1} \boldsymbol{\Omega}_j).$$

By Proposition S.4, we have that $\frac{1}{q_m} E(\mathbf{W}_{i(m)})(\mathbf{W}_{i(m)})^\top = \boldsymbol{\Omega}_m$ since $\prod_{j \neq m} \text{tr}(\boldsymbol{\Omega}_j) = q_m$. Then the desired conclusion follows from Proposition S.6. $\square$

Recall that we have data $\{\mathbf{X}_i, Y_i\}_{i=1}^n$, where $\Pr(Y_i = k) = \pi_k^*$, $\mathbf{X}_i \mid Y_i = k \sim TN(\boldsymbol{\mu}_k^*; \boldsymbol{\Sigma}_1^*, \dots, \boldsymbol{\Sigma}_M^*)$. We consider the parameter space $\boldsymbol{\Theta}(c_\pi, C_b, s, \{C_m\}_{m=1}^M, C_b, \Delta_0)$ in (4.1). For ease of presentation, we write $\boldsymbol{\Theta}$ instead of $\boldsymbol{\Theta}(c_\pi, C_b, s, \{C_m\}_{m=1}^M, C_b, \Delta_0)$ when there is no confusion.

Let $\mathbf{U}_i = \text{vec}(\mathbf{X}_i)$, $\boldsymbol{\phi}_k = \text{vec}(\boldsymbol{\mu}_k)$, $\boldsymbol{\Psi} = \otimes_{m=1}^M \boldsymbol{\Sigma}_m$. For any $\boldsymbol{\theta} = (\pi_k, \boldsymbol{\mu}_k, k = 1, 2; \boldsymbol{\Sigma}_m, m = 1, \dots, M)$, we denote its vectorized counterpart as $\boldsymbol{\theta}^{\text{vec}} = (\pi_k, \boldsymbol{\phi}_k, k = 1, 2, \boldsymbol{\Psi})$. It is easy to see that, if $\boldsymbol{\theta}^* \in \boldsymbol{\Theta}$, we must have $\boldsymbol{\theta}^{\text{vec}*} \in \boldsymbol{\Theta}^{\text{vec}}$, with $C_0 = \prod_{m=1}^M C_m$ and $p = \prod_{m=1}^M p_m$. Moreover,

$$\text{vec}(\hat{\mathbf{B}}^{(t+1)}) = \arg \min_{\boldsymbol{\beta} \in \mathbb{R}^p} \left[\frac{1}{2} \boldsymbol{\beta}^\top \left\{ \bigotimes_{m=M}^{m=1} \hat{\boldsymbol{\Sigma}}^{(t+1)} \right\} \boldsymbol{\beta} - \boldsymbol{\beta}^\top (\hat{\boldsymbol{\mu}}_1^{(t+1)} - \hat{\boldsymbol{\mu}}_2^{(t+1)}) + \lambda_n^{(t+1)} \|\boldsymbol{\beta}\|_1 \right]. \quad (\text{S.51})$$

Hence, DEEM is a special case of PEM with

$$g_{\pi_k}(\boldsymbol{\xi}, \mathbf{U}) = \frac{1}{n} \sum_{i=1}^n \xi_{ik}, \quad g_{\boldsymbol{\phi}_k}(\boldsymbol{\xi}, \mathbf{U}) = \frac{\sum_{i=1}^n \xi_{ik} \mathbf{U}_i}{\sum_{i=1}^n \xi_{ik}}, \quad (\text{S.52})$$

$$g_{\boldsymbol{\Psi}}(\boldsymbol{\xi}, \mathbf{U}) = \otimes \left\{ \frac{1}{np_{-m}} \sum_{i=1}^n \xi_{ik} (\mathbf{X}_i - g_{\boldsymbol{\phi}_k}(\boldsymbol{\xi}, \mathbf{U}))_{(m)} (\mathbf{X}_i - g_{\boldsymbol{\phi}_k}(\boldsymbol{\xi}, \mathbf{U}))_{(m)}^\top \right\}. \quad (\text{S.53})$$

With a little abuse of notation, we define

$$d_{2,s}(\boldsymbol{\theta}, \tilde{\boldsymbol{\theta}}) = (\vee_k |\pi_k - \tilde{\pi}_k|) \vee (\vee_k \|\text{vec}(\boldsymbol{\mu}_k - \tilde{\boldsymbol{\mu}}_k)\|_{2,s}) \vee \|(\boldsymbol{\Psi} - \tilde{\boldsymbol{\Psi}})\tilde{\boldsymbol{\beta}}\|_{2,s}. \quad (\text{S.54})$$

We have two lemmas concerning the theoretical properties of this particular choice of g .

Lemma S.6. *Suppose $\boldsymbol{\theta}^* \in \boldsymbol{\Theta}$ with $0 < c_\pi < 1$ and g defined as in DEEM. We have the following conclusions.*

1. *If $\frac{\log p}{n} = o(1)$, then we must have*

$$g^{lim}(\boldsymbol{\theta}) = \{\pi_k(\boldsymbol{\theta}), \boldsymbol{\mu}_k(\boldsymbol{\theta}), k = 1, 2; \boldsymbol{\Psi}(\boldsymbol{\theta})\}, \quad (\text{S.55})$$

where

$$\pi_k(\boldsymbol{\theta}) = E[\xi_{ik}(\boldsymbol{\theta})], \quad \boldsymbol{\mu}_k(\boldsymbol{\theta}) = \frac{E\xi_{ik}\mathbf{U}_i}{E\xi_{ik}}, \quad (\text{S.56})$$

$$\boldsymbol{\Psi}(\boldsymbol{\theta}) = \bigotimes_{m=M}^{m=1} \left[\frac{1}{p-m} E\{\xi_{i1}(\mathbf{X}_i - \boldsymbol{\mu}_k(\boldsymbol{\theta}))_{(m)}(\mathbf{X}_i - \boldsymbol{\mu}_k(\boldsymbol{\theta}))_{(m)}^\top\} \right]. \quad (\text{S.57})$$

2. *There exists a constant $C_{con} > 0$, such that with a probability greater than $1 - O(p^{-1})$, we have*

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} d_{2,s}(g\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\}, g^{lim}(\boldsymbol{\theta})) \leq C_{con} \sqrt{\frac{s \log p}{n}}, \quad (\text{S.58})$$

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} \| [g_{\boldsymbol{\Psi}}\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\} - \boldsymbol{\Psi}(\boldsymbol{\theta})] \boldsymbol{\beta}^* \|_\infty \leq C_{con} \sqrt{\frac{\log p}{n}}. \quad (\text{S.59})$$

Proof of Lemma S.6. By the first conclusion in Proposition S.8, we have that

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} d_{2,s}^{(1)}(g\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\}, g^{lim}\{\boldsymbol{\theta}\}) \leq C_{con} \sqrt{\frac{s \log p}{n}}. \quad (\text{S.60})$$

Hence, to prove the desired conclusions, it suffices to focus on $\boldsymbol{\Psi}$ and verify that

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} \|g_{\boldsymbol{\Psi}}\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\} - \boldsymbol{\Psi}(\boldsymbol{\theta})\|_{\max} = o_P(1), \quad (\text{S.61})$$

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} \| [g_{\boldsymbol{\Psi}}\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\} - \boldsymbol{\Psi}(\boldsymbol{\theta})] \boldsymbol{\beta}^* \|_{2,s} \leq C_{con} \sqrt{\frac{s \log p}{n}}, \quad (\text{S.62})$$

$$\sup_{\boldsymbol{\theta} \in \mathcal{B}_{con}} \| [g_{\boldsymbol{\Psi}}\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\} - \boldsymbol{\Psi}(\boldsymbol{\theta})] \boldsymbol{\beta}^* \|_\infty \leq C_{con} \sqrt{\frac{\log p}{n}}. \quad (\text{S.63})$$

For simplicity, we write $\widehat{\Psi}(\boldsymbol{\theta}) = g_{\Psi}\{\boldsymbol{\xi}(\boldsymbol{\theta}), \mathbf{U}\}$. We start by bounding $\sup_{\boldsymbol{\theta} \in \mathcal{B}_{\text{con}}} \|\widehat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max}$. Note that $\widehat{\Psi}(\boldsymbol{\theta})$ and $\Psi(\boldsymbol{\theta})$ can be written as

$$\widehat{\Psi}(\boldsymbol{\theta}) = \otimes_{m=M}^{m=1} \widehat{\Sigma}_m(\boldsymbol{\theta}), \quad \Psi(\boldsymbol{\theta}) = \otimes_{m=M}^{m=1} \Sigma_m(\boldsymbol{\theta}). \quad (\text{S.64})$$

By Proposition S.3 we have

$$\begin{aligned} & \|\widehat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max} \\ &= \left\| \otimes_{m=M}^{m=1} \widehat{\Sigma}_m(\boldsymbol{\theta}) - \otimes_{m=M}^{m=1} \Sigma_m(\boldsymbol{\theta}) \right\|_{\max} \\ &\leq \sum_{j=1}^M \left[\prod_{m=M}^{m=j+1} (\|\widehat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max} + \|\Sigma_m(\boldsymbol{\theta})\|_{\max}) \cdot \|\widehat{\Sigma}_j(\boldsymbol{\theta}) - \Sigma_j(\boldsymbol{\theta})\|_{\max} \cdot \prod_{m=j-1}^{m=1} \|\Sigma_m(\boldsymbol{\theta})\|_{\max} \right] \\ &\leq \sum_{j=1}^M \left[\prod_{m=M}^{m=j+1} (\|\widehat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max} + \|\Sigma_m(\boldsymbol{\theta})\|_{\max}) \cdot \|\widehat{\Sigma}_j(\boldsymbol{\theta}) - \Sigma_j(\boldsymbol{\theta})\|_{\max} \right]. \end{aligned}$$

It can be seen that a bound on $\|\widehat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max}$ is implied by bounds on $\|\widehat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max}$. We provide bounds on $\|\widehat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max}$ in what follows. Note that

$$\begin{aligned} \widehat{\Sigma}_m(\boldsymbol{\theta}) &= \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} (\mathbf{X}_i - \widehat{\boldsymbol{\mu}}_k)_{(m)} (\mathbf{X}_i - \widehat{\boldsymbol{\mu}}_k)_{(m)}^{\top} \\ &= \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} (\mathbf{X}_i)_{(m)} (\mathbf{X}_i)_{(m)}^{\top} - \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\mathbf{X}_{i(m)}\}^{\top} \\ &\quad - \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} \{\mathbf{X}_{i(m)}\} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top} + \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top}. \end{aligned}$$

For any i , we have $\sum_{k=1}^2 \xi_{ik} = 1$. Hence,

$$\frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} (\mathbf{X}_i)_{(m)} (\mathbf{X}_i)_{(m)}^{\top} = \frac{1}{nq_m} \sum_{i=1}^n (\mathbf{X}_i)_{(m)} (\mathbf{X}_i)_{(m)}^{\top}.$$

Moreover, for any k , we have $\frac{1}{n} \sum_{i=1}^n \xi_{ik} = \widehat{\pi}_k(\boldsymbol{\theta})$, $\sum_{i=1}^n \xi_{ik} \mathbf{X}_i = \{\sum_{i=1}^n \xi_{ik}\} \widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta}) = n\widehat{\pi}_k(\boldsymbol{\theta}) \widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})$ by the definitions of $\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})$ and $\widehat{\pi}_k(\boldsymbol{\theta})$, which implies

$$\begin{aligned} & \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\mathbf{X}_{i(m)}\}^{\top} = \frac{1}{nq_m} \sum_{k=1}^2 \left\{ \sum_{i=1}^n \xi_{ik} \right\} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top} \\ &= \frac{1}{q_m} \sum_{k=1}^2 \widehat{\pi}_k(\boldsymbol{\theta}) \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top}, \\ & \frac{1}{nq_m} \sum_{i=1}^n \sum_{k=1}^2 \xi_{ik} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top} = \frac{1}{q_m} \sum_{k=1}^2 \widehat{\pi}_k(\boldsymbol{\theta}) \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^{\top}. \end{aligned}$$

Consequently,

$$\widehat{\Sigma}_m(\boldsymbol{\theta}) = \frac{1}{nq_m} \sum_{i=1}^n \mathbf{X}_{i(m)} \{\mathbf{X}_{i(m)}\}^\top - \frac{1}{q_m} \sum_{k=1}^2 \widehat{\pi}_k(\boldsymbol{\theta}) \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^\top.$$

Similarly, we have

$$\Sigma_m(\boldsymbol{\theta}) = \frac{1}{q_m} E\{\mathbf{X}_{i(m)}\} \{\mathbf{X}_{i(m)}\}^\top - \frac{1}{q_m} \sum_{k=1}^2 \pi_k(\boldsymbol{\theta}) \{\boldsymbol{\mu}_k(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_k(\boldsymbol{\theta})\}_{(m)}^\top.$$

As a result,

$$\begin{aligned} & \|\widehat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max} \\ \leq & \left\| \frac{1}{nq_m} \sum_{i=1}^n \mathbf{X}_{i(m)} \{\mathbf{X}_{i(m)}\}^\top - \frac{1}{q_m} E\{\mathbf{X}_{i(m)}\} \{\mathbf{X}_{i(m)}\}^\top \right\|_{\max} \\ & + \frac{1}{q_m} \sum_{k=1}^2 \left\| \widehat{\pi}_k(\boldsymbol{\theta}) \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)} \{\widehat{\boldsymbol{\mu}}_k(\boldsymbol{\theta})\}_{(m)}^\top - \pi_k(\boldsymbol{\theta}) \{\boldsymbol{\mu}_k(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_k(\boldsymbol{\theta})\}_{(m)}^\top \right\|_{\max} \\ \equiv & L_1 + \sum_{k=1}^2 L_2^{(k)}. \end{aligned}$$

For L_1 , we note that

$$\frac{1}{nq_m} \sum_{i=1}^n \mathbf{X}_{i(m)} (\mathbf{X}_{i(m)})^\top \tag{S.65}$$

$$\begin{aligned} = & \frac{1}{nq_m} \sum_{k=1}^2 \sum_{Y_i=k} (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)}^\top + \frac{1}{q_m} \sum_{k=1}^2 \{ \tilde{\pi}_k(\boldsymbol{\mu}_k^*)_{(m)} (\tilde{\boldsymbol{\mu}}_k)_{(m)}^\top \\ & + \tilde{\pi}_k(\tilde{\boldsymbol{\mu}}_k)_{(m)} (\boldsymbol{\mu}_k^*)_{(m)}^\top - \tilde{\pi}_k(\boldsymbol{\mu}_k^*)_{(m)} (\boldsymbol{\mu}_k^*)_{(m)}^\top \}, \end{aligned} \tag{S.66}$$

where $\tilde{\pi}_k = \frac{1}{n} \sum_{i=1}^n 1(Y_i = k)$, $\tilde{\boldsymbol{\mu}}_k = \frac{1}{n\tilde{\pi}_k} \sum_{Y_i=k} \mathbf{X}_i$.

By Proposition S.6, we have that

$$\begin{aligned} & \left\| \frac{1}{nq_m} \sum_{k=1}^2 \sum_{Y_i=k} (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)}^\top - \frac{1}{nq_m} \sum_{k=1}^2 \sum_{Y_i=k} E(\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)} (\mathbf{X}_i - \boldsymbol{\mu}_k^*)_{(m)}^\top \right\|_{\max} \\ = & O_P \left(\sqrt{\frac{\log p_m}{nq_m}} \right). \end{aligned}$$

For the remaining three terms in (S.66), we focus on the term $\frac{\tilde{\pi}_k}{q_m} (\boldsymbol{\mu}_k^*)_{(m)} (\tilde{\boldsymbol{\mu}}_k)_{(m)}^\top$, as all the other terms can be bounded similarly. Note that

$$\frac{\tilde{\pi}_k}{q_m} (\boldsymbol{\mu}_k^*)_{(m)} (\tilde{\boldsymbol{\mu}}_k)_{(m)}^\top - \frac{\pi_k^*}{p-m} (\boldsymbol{\mu}_k^*)_{(m)} (\boldsymbol{\mu}_k^*)_{(m)}^\top \tag{S.67}$$

$$= \frac{\tilde{\pi}_k}{q_m} (\boldsymbol{\mu}_k^*)_{(m)} (\tilde{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k^*)_{(m)}^\top + \frac{1}{q_m} (\tilde{\pi}_k - \pi_k) (\boldsymbol{\mu}_k^*)_{(m)} (\boldsymbol{\mu}_k^*)_{(m)}^\top \tag{S.68}$$

$$\equiv L_3 + L_4. \tag{S.69}$$

By Hoeffding's inequality, we have $|\tilde{\pi}_k - \pi_k| = O_P(\sqrt{\frac{1}{n}})$ and thus $\|L_4\|_{\max} = O_P(\sqrt{\frac{1}{n}})$.

For L_3 , we note that by the Chernoff bound $\|\tilde{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k^*\|_{\max} = O_P(\sqrt{\frac{\log p}{n}})$ and thus

$$\|L_3\|_{\max} = \left\| \frac{\tilde{\pi}_k}{q_m} (\boldsymbol{\mu}_k^*)_{(m)} (\tilde{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k^*)_{(m)}^\top \right\|_{\max} \leq \frac{1}{q_m} \|(\boldsymbol{\mu}_k^*)_{(m)}\|_{\infty} \cdot \|\tilde{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k^*\|_{\max} \quad (\text{S.70})$$

$$= \frac{1}{q_m} \cdot O_P\left(\sqrt{\frac{\log p}{n}}\right). \quad (\text{S.71})$$

It follows that $\|\hat{\Sigma}_m(\boldsymbol{\theta}) - \Sigma_m(\boldsymbol{\theta})\|_{\max} = O_P(\sqrt{\frac{\log p}{n}})$ and

$$\|\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max} = O_P\left(\sqrt{\frac{\log p}{n}}\right). \quad (\text{S.72})$$

To verify (S.61)–(S.63), we note that

$$\begin{aligned} & \|(\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta}))\boldsymbol{\beta}^*\|_{2,s} = \arg \max_{\mathbf{u} \in \Gamma(s) \cap \mathbb{S}^{p-1}} |\langle \mathbf{u}, (\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta}))\boldsymbol{\beta}^* \rangle| \\ & \leq \|\mathbf{u}\|_1 \|\boldsymbol{\beta}^*\|_1 \|\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max} \leq 7C\sqrt{s} \|\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max} \\ & \|(\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta}))\boldsymbol{\beta}^*\|_{\infty} \leq \|\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max} \|\boldsymbol{\beta}^*\|_1 \leq C_b \|\hat{\Psi}(\boldsymbol{\theta}) - \Psi(\boldsymbol{\theta})\|_{\max}. \end{aligned}$$

By (S.72) we have (S.61)–(S.63). □

Lemma S.7. Assume that $\boldsymbol{\theta}^* \in \Theta$. For sufficiently large Δ_0 there exists $0 < \kappa_0 < \frac{1}{2\sqrt{(16C_0)}}$ and $\Delta_0 > 0$ that does not depend on n, p such that for any $\boldsymbol{\theta}^{\text{vec}} \in \mathcal{B}_{\text{con}}$, we have

$$d_2(g^{\text{lim}}(\boldsymbol{\theta}^{\text{vec}}), \boldsymbol{\theta}^{*\text{vec}}) \leq \kappa_0 \cdot (d_{2,s}(\boldsymbol{\theta}^{\text{vec}}, \boldsymbol{\theta}^{*\text{vec}}) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2). \quad (\text{S.73})$$

Proof of Lemma S.7. By the first conclusion in Proposition S.8, we have that

$$d_2^{(1)}(g^{\text{lim}}(\boldsymbol{\theta}^{\text{vec}}), \boldsymbol{\theta}^{*\text{vec}}) \leq \kappa_0^{(1)}(\Delta) \cdot (d_{2,s}(\boldsymbol{\theta}^{\text{vec}}, \boldsymbol{\theta}^{*\text{vec}}) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2). \quad (\text{S.74})$$

Thus, it only remains to verify that

$$\|(\Psi(\boldsymbol{\theta}) - \Psi^*)\boldsymbol{\beta}^*\|_2 \leq \kappa_0 \cdot (d_{2,s}(\boldsymbol{\theta}^{\text{vec}}, \boldsymbol{\theta}^{*\text{vec}}) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2), \quad (\text{S.75})$$

where $\Psi(\boldsymbol{\theta}) = \bigotimes_{m=M}^{m=1} \Sigma_m(\boldsymbol{\theta})$. We start by comparing $\Sigma_m(\boldsymbol{\theta})$ with Σ_m^* . For any k , we have that

$$\frac{1}{q_m} E[\xi_{ik}(\boldsymbol{\theta}^*) \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)} \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.76})$$

$$= \frac{1}{q_m} E[\Pr(Y = k \mid \mathbf{X}, \boldsymbol{\theta}^*) (\mathbf{X} - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.77})$$

$$= \frac{1}{q_m} E[E[\mathbf{1}(Y = k) \mid \mathbf{X}, \boldsymbol{\theta}^*] (\mathbf{X} - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.78})$$

$$= \frac{1}{q_m} E[E[\mathbf{1}(Y = k) (\mathbf{X} - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)}^\top \mid \mathbf{X}, \boldsymbol{\theta}^*]] \quad (\text{S.79})$$

$$= \frac{1}{q_m} E_{\boldsymbol{\theta}^*}[\mathbf{1}(Y = k) (\mathbf{X} - \boldsymbol{\mu}_k^*)_{(m)} \{\mathbf{X} - \boldsymbol{\mu}_k^*\}_{(m)}^\top] \quad (\text{S.80})$$

$$= \pi_k^* \Sigma_m^*, \quad (\text{S.81})$$

where the last equation follows from Proposition S.4. Hence, $\Sigma_m(\boldsymbol{\theta}^*) = \sum_k \pi_k^* \Sigma_m^* = \Sigma_m^*$. It follows that

$$\begin{aligned} & q_m \{\Sigma_m(\boldsymbol{\theta}) - \Sigma_m^*\} \\ = & \pi_1(\boldsymbol{\theta}) [\boldsymbol{\mu}_1(\boldsymbol{\theta})]_{(m)} [\boldsymbol{\mu}_1(\boldsymbol{\theta})]_{(m)}^\top - \pi_1^* [\boldsymbol{\mu}_1^*]_{(m)} [\boldsymbol{\mu}_1^*]_{(m)}^\top + \pi_2(\boldsymbol{\theta}) \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)}^\top - \pi_2^* \{\boldsymbol{\mu}_2^*\}_{(m)} \{\boldsymbol{\mu}_2^*\}_{(m)}^\top. \end{aligned}$$

We note that $\pi_1(\boldsymbol{\theta}) \boldsymbol{\mu}_1(\boldsymbol{\theta}) + \pi_2(\boldsymbol{\theta}) \boldsymbol{\mu}_2(\boldsymbol{\theta}) = E\mathbf{X} = \pi_1^* \boldsymbol{\mu}_1^* + \pi_2^* \boldsymbol{\mu}_2^*$. Hence, let $\boldsymbol{\mu}^* = \pi_1^* \boldsymbol{\mu}_1^* + \pi_2^* \boldsymbol{\mu}_2^*$ and we have

$$\begin{aligned} & \pi_1(\boldsymbol{\theta}) \{\boldsymbol{\mu}_1(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_1(\boldsymbol{\theta})\}_{(m)}^\top + \pi_2(\boldsymbol{\theta}) \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)}^\top \\ = & \pi_1(\boldsymbol{\theta}) \{\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)}^\top + \pi_1(\boldsymbol{\theta}) \{\boldsymbol{\mu}_1(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_{(m)}^*\}^\top - \pi_1(\boldsymbol{\theta}) \boldsymbol{\mu}_{(m)}^* \{\boldsymbol{\mu}_{(m)}^*\}^\top \\ & + \pi_1(\boldsymbol{\theta}) \boldsymbol{\mu}_{(m)}^* \{\boldsymbol{\mu}_1(\boldsymbol{\theta})\}_{(m)}^\top \\ & + \pi_2(\boldsymbol{\theta}) \{\boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)}^\top + \pi_2(\boldsymbol{\theta}) \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)} \{\boldsymbol{\mu}_{(m)}^*\}^\top - \pi_2(\boldsymbol{\theta}) \boldsymbol{\mu}_{(m)}^* \{\boldsymbol{\mu}_{(m)}^*\}^\top \\ & + \pi_2(\boldsymbol{\theta}) \boldsymbol{\mu}_{(m)}^* \{\boldsymbol{\mu}_2(\boldsymbol{\theta})\}_{(m)}^\top \\ = & \pi_1(\boldsymbol{\theta}) \{\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)}^\top + \pi_2(\boldsymbol{\theta}) \{\boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^*\}_{(m)}^\top \\ & + \{\boldsymbol{\mu}_{(m)}^*\} \{\boldsymbol{\mu}_{(m)}^*\}^\top. \end{aligned}$$

Similarly,

$$\begin{aligned} & \pi_1^* [\boldsymbol{\mu}_1^*]_{(m)} [\boldsymbol{\mu}_1^*]_{(m)}^\top + \pi_2^* \{\boldsymbol{\mu}_2^*\}_{(m)} \{\boldsymbol{\mu}_2^*\}_{(m)}^\top \\ = & \pi_1^* \{\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*\}_{(m)}^\top + \pi_2^* \{\boldsymbol{\mu}_2^* - \boldsymbol{\mu}^*\}_{(m)} \{\boldsymbol{\mu}_2^* - \boldsymbol{\mu}^*\}_{(m)}^\top + \{\boldsymbol{\mu}_{(m)}^*\} \{\boldsymbol{\mu}_{(m)}^*\}^\top. \end{aligned}$$

It follows that

$$\begin{aligned}
& q_m \{ \Sigma_m(\boldsymbol{\theta}) - \Sigma_m^* \} \\
&= \pi_1(\boldsymbol{\theta}) \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)}^\top + \pi_2(\boldsymbol{\theta}) \{ \boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)}^\top \\
&\quad - \pi_1^* \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)}^\top - \pi_2^* \{ \boldsymbol{\mu}_2^* - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_2^* - \boldsymbol{\mu}^* \}_{(m)}^\top \\
&= [\pi_1(\boldsymbol{\theta}) \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)}^\top - \pi_1^* \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)}^\top] \\
&\quad + [\pi_2(\boldsymbol{\theta}) \{ \boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_2(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)}^\top - \pi_2^* \{ \boldsymbol{\mu}_2^* - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_2^* - \boldsymbol{\mu}^* \}_{(m)}^\top] \\
&= q_m \cdot L_5 + q_m \cdot L_6.
\end{aligned}$$

To bound $\| \Sigma_m(\boldsymbol{\theta}) - \Sigma_m^* \|_2$, we only provide a bound on L_5 , as L_6 can be bounded similarly.

$$\begin{aligned}
L_5 &= \frac{1}{q_m} \| \pi_1(\boldsymbol{\theta}) \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^* \}_{(m)}^\top - \pi_1^* \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)} \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)}^\top \|_2 \\
&\leq \frac{1}{q_m} | \pi_1^* - \pi_1(\boldsymbol{\theta}) | \cdot \| (\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*)_{(m)} (\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*)_{(m)}^\top \|_2 \\
&\quad + \frac{1}{q_m} | \pi_1^* | \cdot \| (\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*)_{(m)} (\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*)_{(m)}^\top \|_2 \\
&\quad + \frac{1}{q_m} | \pi_1(\boldsymbol{\theta}) | \cdot \| (\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*)_{(m)} (\boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}^*)_{(m)}^\top \|_2 \\
&\equiv L_7 + L_8 + L_9.
\end{aligned}$$

For L_7 , we note that $\| \text{vec}(\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*) \|_2^2 \leq C_0 \Delta^2$ by the definition of Δ . By the second conclusion in Proposition S.8, there exists $\kappa_0^{(1)}(\Delta) = o(\frac{1}{\Delta^2})$ such that

$$\begin{aligned}
L_7 &\leq \frac{1}{q_m} | \pi_1^* - \pi_1(\boldsymbol{\theta}) | \cdot \| \text{vec}(\boldsymbol{\mu}_1^* - \boldsymbol{\mu}^*) \|_2^2 \leq \frac{C_0}{q_m} | \pi_1^* - \pi_1(\boldsymbol{\theta}) | \cdot \Delta^2 \\
&\leq \frac{C_0}{q_m} \Delta^2 \cdot \kappa_0^{(1)}(\Delta) \cdot d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \| \boldsymbol{\beta} - \boldsymbol{\beta}^* \|_2.
\end{aligned}$$

Similarly,

$$\begin{aligned}
L_8 &\leq \frac{1}{q_m} \| [\{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \}_{(m)} + \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)}] \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \}_{(m)}^\top \|_2 \\
&\leq \frac{1}{q_m} [\| \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \}_{(m)} \|_2 + \| \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)} \|_2] \cdot \| \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \}_{(m)}^\top \|_2.
\end{aligned}$$

We note that

$$\begin{aligned}
\| \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \}_{(m)} \|_2 &\leq \| \text{vec} \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \} \|_F = \| \text{vec} \{ \boldsymbol{\mu}_1(\boldsymbol{\theta}) - \boldsymbol{\mu}_1^* \} \|_2 \\
&\leq \kappa_0^{(1)}(\Delta) \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \| \boldsymbol{\beta} - \boldsymbol{\beta}^* \|_2 \}, \\
\| \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \}_{(m)} \|_2 &\leq \| \text{vec} \{ \boldsymbol{\mu}_1^* - \boldsymbol{\mu}^* \} \|_2 \leq \sqrt{C_0} \Delta.
\end{aligned}$$

Further noting that for any $\boldsymbol{\theta} \in \mathcal{B}_{con}$, we have $d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \leq (a_b + \sqrt{C_0})\Delta$. As a result,

$$L_8 \leq \frac{1}{q_m} \{ \kappa^{(1)} \cdot d_{2,s}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) + \sqrt{C_0}\Delta \} \kappa^{(1)} \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \} \quad (\text{S.82})$$

$$\leq \frac{1}{q_m} \{ a_b + \sqrt{C_0} \} \Delta \cdot \kappa_0^{(1)}(\Delta) \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \}. \quad (\text{S.83})$$

It is easy to see that L_9 has a similar bound. Therefore, let $\kappa_{\Sigma,0} = \frac{1}{q_m} \{ a_b + \sqrt{C_0} \} \Delta \cdot \kappa^{(1)}(\Delta)$ and we have $\kappa_{\Sigma,0} = o(1/\Delta)$ and

$$\|\Sigma_m(\boldsymbol{\theta}) - \Sigma_m^*\|_2 \leq \frac{1}{p-m} \kappa_{\Sigma,0} \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \}. \quad (\text{S.84})$$

By Proposition S.3 we have

$$\begin{aligned} & \left\| \bigotimes_{m=M}^{m=1} \Sigma_m(\boldsymbol{\theta}) - \bigotimes_{m=M}^{m=1} \Sigma_m^* \right\|_2 \\ & \leq \sum_{j=1}^M \left[\prod_{m=M}^{m=j} (\|\Sigma_m(\boldsymbol{\theta}) - \Sigma_m^*\|_2 + \|\Sigma_m^*\|_2) \cdot \|\Sigma_j(\boldsymbol{\theta}) - \Sigma_j^*\|_2 \cdot \prod_{m=j-1}^{m=1} \|\Sigma_m^*\|_2 \right] \\ & \leq M \left\{ \frac{\kappa_{\Sigma,0}}{p-m} \cdot d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 + C_0 \right\}^{M-1} \cdot \frac{\kappa_{\Sigma,0}}{p-m} \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \} \\ & \leq M \left\{ \frac{\kappa_{\Sigma,0}}{p-m} (a_b + \sqrt{C_0})\Delta + C_0 \right\}^{M-1} \cdot \frac{\kappa_{\Sigma,0}}{p-m} \cdot \{ d_{2,s}^{(1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2 \}. \end{aligned}$$

Moreover,

$$\|\{\Sigma(\boldsymbol{\theta}) - \Sigma^*\}\boldsymbol{\beta}^*\|_2 \leq \|\Sigma(\boldsymbol{\theta}) - \Sigma^*\|_2 \cdot \|\boldsymbol{\beta}^*\|_2 \quad (\text{S.85})$$

$$\leq M \left\{ \frac{\kappa_{\Sigma,0}}{p-m} (a_b + \sqrt{C_0})\Delta + C_0 \right\}^{M-1} \cdot \frac{\kappa_{\Sigma,0}}{p-m} \cdot d_{2,s}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \cdot \sqrt{C_0}\Delta \quad (\text{S.86})$$

$$\equiv \kappa_{\Sigma} \cdot d_{2,s}(\boldsymbol{\theta}, \boldsymbol{\theta}^*). \quad (\text{S.87})$$

Since $\kappa_{\Sigma,0} = o(1/\Delta)$, we have $\kappa_{\Sigma} = o(1)$. It follows that there exists $\Delta_0 > 0$ such that when $\Delta > \Delta_0$, we have that $\kappa_0 = \min\{\kappa_0^{(1)}(\Delta_0), \kappa_{\Sigma}\} < \frac{1}{2\sqrt{16C_0}}$ and

$$d_{2,s}(g^{lim}(\boldsymbol{\theta}), \boldsymbol{\theta}^*) \leq \kappa_0 \cdot d_{2,s}(\boldsymbol{\theta}, \boldsymbol{\theta}^*) \vee \|\boldsymbol{\beta} - \boldsymbol{\beta}^*\|_2. \quad (\text{S.88})$$

□

We are now ready to prove Theorems 1 & 2. Recall that we generate the sequence of tuning parameters $\lambda_n^{(t)}$ as follows. For some constant $C_d, C_\lambda > 0$, we let

$$\lambda^{(0)} = C_d \cdot (|\hat{\pi}_2| \vee \|\text{vec}(\hat{\boldsymbol{\mu}}_1^{(0)} - \hat{\boldsymbol{\mu}}_2^{(0)})\|_{2,s} \vee \|\bigotimes_{m=M}^{m=1} \hat{\Sigma}_m^{(0)}\|_{2,s}) / \sqrt{s} + C_\lambda \sqrt{\sum_{m=1}^M \log p_m / n}, \quad (\text{S.89})$$

and $\lambda^{(t+1)} = \kappa\lambda^{(t)} + C_\lambda\sqrt{\frac{\log p}{n}}$. Note this choice coincides with that in the generic PEM algorithm.

Proof of Theorem 1. We prove Theorem 1 by verifying the five assumptions for Proposition S.7. Assumption (A1') apparently holds when $\mathbf{X}_i$ is drawn from the mixture tensor normal model. Condition (C1) implies (A5), while (A2) & (A3) are implied by Lemma S.6. By Lemma S.7, (A4) is also true. Hence, the desired conclusion follows from the first conclusion in Proposition S.7. $\square$

Proof of Theorem 2. Similar to the proof of Theorem 1, (A1')–(A5) holds. By the second conclusion in Proposition S.7, we have the first desired conclusion. To show the second conclusion, we consider

$$\boldsymbol{\theta}_i = \{\pi_1 = \pi_2, \boldsymbol{\mu}_1 = \epsilon\mathbf{U}_i + \eta\mathbf{E}_i, \boldsymbol{\mu}_2 = -\boldsymbol{\mu}_1, \boldsymbol{\Sigma}_1 = \sigma^2\mathbf{I}_{p_1}, \boldsymbol{\Sigma}_m = \mathbf{I}_{p_m}, m > 1\}, i = 1, \dots, N \quad (\text{S.90})$$

such that the corresponding $\boldsymbol{\theta}_i^{\text{vec}}$ is equal to that defined in (S.45). Note that this is possible because $\bigotimes_{m=1}^M \boldsymbol{\Sigma}_m = \sigma^2\mathbf{I}_p$. It immediately follows from Proposition S.9 that

$$\inf_{\hat{Y}} \sup_{i=1}^N \{R(\hat{Y} \mid \boldsymbol{\theta}_i) - R(\text{Opt} \mid \boldsymbol{\theta}_i)\} \succeq \frac{\log p}{n}. \quad (\text{S.91})$$

Since $\boldsymbol{\theta}_i \in \boldsymbol{\Theta}$, we have

$$\inf_{\hat{Y}} \sup_{\boldsymbol{\theta} \in \boldsymbol{\Theta}} \{R(\hat{Y} \mid \boldsymbol{\theta}) - R(\text{Opt} \mid \boldsymbol{\theta})\} \succeq \frac{\log p}{n}. \quad (\text{S.92})$$

Combine this result with the first conclusion and we have that DEEM is minimax optimal. $\square$

G An Initialization Algorithm

We propose an initialization algorithm that satisfies Condition (C1) with a high probability. For ease of presentation, we refer to Algorithm $\mathcal{B}$ in Hardt & Price (2015) as the HP Algorithm. The HP Algorithm considers $\mathbf{U}_i, i = 1, \dots, n$, where $\mathbf{U}_i \in \mathbb{R}^p$ is drawn from Gaussian mixtures. In other words, we assume that

$$(A1) \quad \mathbf{U}_i \sim \pi_1^* N(\boldsymbol{\phi}_1^*, \boldsymbol{\Psi}^*) + \pi_2^* N(\boldsymbol{\phi}_1^*, \boldsymbol{\Psi}^*), \text{ where } 0 < \pi_k^* < 1, \boldsymbol{\phi}_k^* \in \mathbb{R}^p, \boldsymbol{\Psi}^* \in \mathbb{R}^{p \times p}.$$

The HP Algorithm finds estimates for $\boldsymbol{\phi}_k^*, \boldsymbol{\Psi}^*$. We have the following result for the HP Algorithm.

Proposition S.11 (Hardt & Price (2015), Cai et al. (2019) c.f Proposition 3.1). *Suppose that we observe n i.i.d samples $\mathbf{U}_i$ from Model (A1). Given $\epsilon, \delta > 0$, if $\frac{1}{\epsilon\delta} \log(\frac{p}{\delta} \log(\frac{1}{\epsilon})) = o(n)$, then with probability at least $1 - \delta$, the HP algorithm produces estimates $\hat{\phi}_k^{(0)}$ and $\hat{\Psi}^{(0)}$ such that for some permutation $\Pi : \{1, 2\} \mapsto \{1, 2\}$, we have*

$$\max\{\|\hat{\phi}_1^{(0)} - \phi_{\Pi(1)}^*\|_\infty^2, \|\hat{\phi}_2^{(0)} - \phi_{\Pi(2)}^*\|_\infty^2, \|\hat{\Psi}^{(0)} - \Psi^*\|_{\max}\} \leq \epsilon(\frac{1}{4}\|\phi_1^* - \phi_2^*\|_\infty^2 + \|\Psi^*\|_{\max}). \quad (\text{S.93})$$

Obviously, a main difference in our problem of interest is that we consider clustering for tensors instead of vectors. But we make two observations to connect tensor clustering to the HP Algorithm. First, for any two indices $\mathcal{J} = (j_1, \dots, j_M)$ and $\mathcal{J}' = (j'_1, \dots, j'_M)$, let $\mathbf{U}_i = (X_{i,\mathcal{J}}, X_{i,\mathcal{J}'})^\top \in \mathbb{R}^2$. Then $\mathbf{U}_i$ follows Model (A1), and the HP algorithm can be applied to $\mathbf{U}_i$ to estimate the within-cluster means of $X_{i,\mathcal{J}}$ and $X_{i,\mathcal{J}'}$. Repeat this procedure for all the indices and we can find an initial value for $\hat{\mu}_k$. We do note that the HP Algorithm could be applied to each $X_{\mathcal{J}}$ as well, but doing so could potentially lead to mismatched mean estimates; see our discussion after we propose our algorithm. Hence, we instead consider a pair of variables at a time to match the estimates.

Second, for Σ_m^* , we average $\mathbf{X}_i$ along the m -th mode to reduce them to vectors that follow Model (A1). More explicitly, let

$$\mathbf{X}_i^m = \frac{1}{\sqrt{q_m}} \mathbf{X}_{(m)} \mathbf{1}_{q_m} \in \mathbb{R}^{p_m}. \quad (\text{S.94})$$

We have the following lemma.

Lemma S.8. *For $\mathbf{X}_i$ following TNMM and $\mathbf{X}_i^m$ defined in (S.94), we have $\mathbf{X}_i^m \sim \pi_1^* N(\mu_1^m, a_m \Sigma_m^*) + \pi_2^* N(\mu_2^m, a_m \Sigma_m^*)$, where $\mu_k^m = (\mu_k^*)_{(m)} \cdot \frac{1}{\sqrt{q_m}} \mathbf{1}_{q_m}$ and $a_m = \frac{1}{q_m} \mathbf{1}_{q_m}^\top \Sigma_{-m} \mathbf{1}_{q_m}$ with $\Sigma_{-m} = \Sigma_M^* \otimes \dots \otimes \Sigma_{m+1}^* \otimes \Sigma_{m-1}^* \otimes \dots \otimes \Sigma_1^*$.*

The proof of Lemma S.8 is presented in Section G.1. Lemma S.8 indicates that the within-cluster covariance of $\mathbf{X}_i^m$ is proportional to Σ_m^* . Moreover, recall that we assume that $\sigma_{m,11}^* = 1, m > 1$. Hence, with proper scaling, we can apply the HP algorithm to $\mathbf{X}_i^m$ to find an initial estimate for Σ_m^* .

With these two observations, we propose the initialization algorithm for tensor data in Algorithm S.4. Algorithm S.4 finds $\hat{\mu}_k^{(0)}, \hat{\Sigma}_m^{(0)}$ in two mostly separate steps. The step for finding $\hat{\Sigma}_m^{(0)}$ (i.e, Step 3) is relatively more straightforward; it converts our tensor observations in the same

way suggested in Lemma S.8 to obtain estimates for $a_m \Sigma_m^*$. Then these estimates are rescaled according to our identifiability assumptions in Section 2.2.

The step for finding $\hat{\mu}_k^{(0)}$ (i.e, Step 2) is more involved. As briefly mentioned earlier, a major obstacle in this step is that, if we estimate $\mu_{k\mathcal{J}}^*$ separately, the estimates could be mismatched. For example, if we consider two indices $\mathcal{J}$ and $\mathcal{J}'$ separately to obtain $\hat{\xi}_{k\mathcal{J}}$ and $\hat{\xi}_{k\mathcal{J}'}$, it is possible that $\hat{\xi}_{1\mathcal{J}}$ is close to $\mu_{1\mathcal{J}}^*$, while $\hat{\xi}_{1\mathcal{J}'}$ is close to $\mu_{2\mathcal{J}'}^*$, as the cluster labels are not identifiable themselves. Consequently, when we put them together, $(\hat{\xi}_{1\mathcal{J}}, \hat{\xi}_{1\mathcal{J}'})$ is neither close to $\mu_{1,(\mathcal{J},\mathcal{J}')}^*$, nor $\mu_{2,(\mathcal{J},\mathcal{J}')}^*$. To resolve this issue, in Step 2(a) we find an element $X_{\mathcal{J}_0}$ such that the clusters are best separated so that estimates for all the other elements can be matched according to this element. Indeed, in Step 2(b), we repeatedly apply the HP Algorithm to $(X_{\mathcal{J}_0}, X_{\mathcal{J}})$ to obtain $(\hat{\xi}_{k\mathcal{J}_0}^{\mathcal{J}}, \hat{\xi}_{k\mathcal{J}})$. The estimate $\hat{\xi}_{k\mathcal{J}_0}^{\mathcal{J}}$ is compared to $\tilde{\mu}_{k\mathcal{J}_0}$ to match $\hat{\xi}_{k\mathcal{J}}$ in Step 2(c).

In principle, one could also vectorize the whole tensor $\mathbf{X}_i$ to $\mathbf{U}_i \in \mathbb{R}^{\prod_{m=1}^M p_m}$ and directly apply the HP Algorithm to $\mathbf{U}_i$ to find initial values for μ_k^* . Such an approach does not have the potential issue of mismatched estimates, but it will result in drastic increase in storage. This is because the HP algorithm produces estimates for the mean and the covariance matrix simultaneously. Hence, in order to estimate μ_k^* , the HP Algorithm needs to estimate $\text{cov}(\mathbf{U}_i | Y_i)$ at the same time. The latter has an intimidating number of parameters at the order of $O(\prod_{m=1}^M p_m^2)$, and significantly inflates the storage and computation costs. However, $\text{cov}(\mathbf{U}_i | Y_i)$ is not needed for later tensor clustering, and thus the additional costs do not facilitate our analysis in any way. In contrast, Algorithm S.4 avoids estimating $\text{cov}(\mathbf{U}_i | Y_i)$ and keeps the storage and computation costs at a low level. This fact demonstrates the potential benefit of exploiting the tensor structure in initialization as well.

G.1 Proof of Lemma S.8

Define Y_i as the latent variable such that $\mathbf{X}_i | (Y_i = k) \sim TN(\mu_k^*; \Sigma_1^*, \dots, \Sigma_M^*)$. It suffices to show that $\mathbf{X}_i^m | (Y_i = k) \sim N(\mu_k^m, a_m \Sigma_m^*)$. By the definition of tensor normal distribution, there exists $\mathbf{Z} \in \mathbb{R}^{p_1 \times \dots \times p_M}$ with iid $N(0, 1)$ elements such that, when $Y_i = k$,

$$\mathbf{X}_i = \mu_k^* + \llbracket \mathbf{Z}; (\Sigma_1^*)^{1/2}, \dots, (\Sigma_M^*)^{1/2} \rrbracket. \quad (\text{S.95})$$

By Proposition S.2, we have that, conditional on $Y_i = k$,

$$\mathbf{X}_i^m = (\mathbf{X}_i)_{(m)} \cdot \frac{1}{\sqrt{q_m}} \mathbf{1}_{q_m} = (\mu_k^m)_{(m)} \cdot \frac{1}{q_m} \mathbf{1}_{q_m} + (\Sigma_m^*)^{1/2} \mathbf{Z}_{(m)} \left(\frac{1}{\sqrt{q_m}} \Sigma_{-m} \mathbf{1}_{q_m} \right). \quad (\text{S.96})$$

Algorithm S.4 An initialization algorithm for tensor data.

1. Input $\mathbf{X}_i, i = 1, \dots, n$.
 2. Find $\hat{\boldsymbol{\mu}}_k^{(0)}$ as follows:
 - (a) For each $\mathcal{J}$, apply the HP algorithm to $\{X_{i,\mathcal{J}}\}_{i=1}^n$ to obtain $\tilde{\mu}_{k,\mathcal{J}}$ as estimates of $\mu_{k,\mathcal{J}}$. Find $\mathcal{J}_0 = \arg \max_{\mathcal{J}} |\tilde{\mu}_{2,\mathcal{J}} - \tilde{\mu}_{1,\mathcal{J}}|$. Let $\hat{\mu}_{k,\mathcal{J}_0}^{(0)} = \tilde{\mu}_{k,\mathcal{J}_0}$. Also record $\widehat{\text{var}}(X_{i,1\dots 1} | Y)$ obtained from the HP Algorithm on $\{X_{i,1\dots 1}\}_{i=1}^n$.
 - (b) For each $\mathcal{J}$, apply the HP Algorithm to $(X_{i,\mathcal{J}_0}, X_{i,\mathcal{J}})_{i=1}^n$ to find $\hat{\nu}_k^{\mathcal{J}} = (\hat{\nu}_{k,\mathcal{J}_0}^{\mathcal{J}}, \hat{\nu}_{k,\mathcal{J}}^{\mathcal{J}})$ such that $\hat{\nu}_{k,\mathcal{J}_0}^{\mathcal{J}}$ is an estimate for the within-cluster mean of $X_{\mathcal{J}_0}$ and $\hat{\nu}_{k,\mathcal{J}}^{\mathcal{J}}$ estimates the within-cluster mean of $X_{\mathcal{J}}$.
 - (c) For each $\mathcal{J} \neq \mathcal{J}_0$, find a permutation $\Lambda_{\mathcal{J}} : \{1, 2\} \mapsto \{1, 2\}$ to minimize $\max_k \{|\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0}^{\mathcal{J}} - \tilde{\mu}_{k,\mathcal{J}_0}|\}$ and set $\hat{\mu}_{k,\mathcal{J}}^{(0)} = \hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}}^{\mathcal{J}}$.
 3. Find $\hat{\boldsymbol{\Sigma}}_m^{(0)}$ as follows:
 - (a) For $m = 1, \dots, M$, do
 - i. Let $\mathbf{X}_i^m = \frac{1}{\sqrt{p-m}} \mathbf{X}_{(m)} \mathbf{1}_{p-m}$. Apply the HP algorithm to $\mathbf{X}_i^m$ to estimate the within-cluster covariance of $\mathbf{X}_i^m$, denoted as $\hat{\boldsymbol{\Psi}}_m^{(0)}$.
 - ii. Set $\tilde{\boldsymbol{\Psi}}_m^{(0)} \propto \hat{\boldsymbol{\Psi}}_m^{(0)}$ with $\tilde{\Psi}_{m,11}^{(0)} = 1$.
 - (b) Let $\hat{\boldsymbol{\Sigma}}_1^{(0)} = \widehat{\text{var}}(X_{i,1\dots 1} | Y) \cdot \hat{\boldsymbol{\Psi}}_1^{(0)}$, where $\widehat{\text{var}}(X_{i,1\dots 1} | Y)$ is found in Step 2(a). For $m > 1$, let $\hat{\boldsymbol{\Sigma}}_m^{(0)} = \hat{\boldsymbol{\Psi}}_m^{(0)}$.
 4. Output $\hat{\boldsymbol{\mu}}_k^{(0)}, \hat{\boldsymbol{\Sigma}}_m^{(0)}$.
-

Note that $\mathbf{X}_i^m$ is a vector and $\text{vec}\{(\boldsymbol{\Sigma}_m^*)^{1/2}\mathbf{Z}_{(m)}(\frac{1}{\sqrt{q_m}}\boldsymbol{\Sigma}_{-m}\mathbf{1}_{q_m})\} = \{(\frac{1}{\sqrt{q_m}}\boldsymbol{\Sigma}_{-m}\mathbf{1}_{q_m})^\top \otimes (\boldsymbol{\Sigma}_m^*)^{1/2}\}\text{vec}(\mathbf{Z}_{(m)})$. It follows that $\mathbf{X}_i^m \mid (Y_i = k) \sim N(\boldsymbol{\mu}_k^m, a_m \boldsymbol{\Sigma}_m^*)$.

G.2 Proof of Lemma 4

We start by verifying that the estimates $\hat{\boldsymbol{\phi}}_k^{(0)} = \text{vec}(\hat{\boldsymbol{\mu}}_k^{(0)})$, $\hat{\boldsymbol{\Psi}}^{(0)} = \otimes_{m=1}^M \hat{\boldsymbol{\Sigma}}_m^{(0)}$ satisfy Condition (IC). We consider $\hat{\boldsymbol{\mu}}_k^{(0)}$ and $\otimes_{m=1}^M \hat{\boldsymbol{\Sigma}}_m^{(0)}$ separately.

The bound on $\hat{\boldsymbol{\mu}}_k^{(0)}$. For $\hat{\boldsymbol{\mu}}_k^{(0)}$ and some constant $C_{\text{init}} > 0$, we consider the event

$$\begin{aligned} \mathcal{A} = & \{ \text{For each } \mathcal{J}, \text{ there exist a permutations } \Pi_{\mathcal{J}}, \Upsilon_{\mathcal{J}} \text{ such that } |\tilde{\mu}_{\Pi_{\mathcal{J}}(k), \mathcal{J}_0} - \mu_{k, \mathcal{J}}^*| \leq \frac{C_{\text{init}}}{s} \\ & \text{and } \|\hat{\boldsymbol{\nu}}_{\Upsilon_{\mathcal{J}}(k)}^{\mathcal{J}} - \boldsymbol{\mu}_{k, (\mathcal{J}, \mathcal{J}_0)}^*\|_{\max} \leq \frac{C_{\text{init}}}{s} \}, \end{aligned} \quad (\text{S.97})$$

where $\tilde{\mu}_{k, \mathcal{J}_0}$ and $\hat{\boldsymbol{\nu}}_k^{\mathcal{J}}$ are defined in Step 2(a) & 2(b) of Algorithm S.4, respectively.

We claim that $\mathcal{A}$ implies the existence of a permutation Π such that $\|\hat{\boldsymbol{\mu}}_{\Pi(k)} - \boldsymbol{\mu}_k^*\|_{\max} \preceq \frac{1}{s}$. We now prove this claim. We consider two cases determined by the magnitude of $\|\boldsymbol{\mu}_1^* - \boldsymbol{\mu}_2^*\|_{\max}$.

Case 1: $\|\boldsymbol{\mu}_1^* - \boldsymbol{\mu}_2^*\|_{\max} \leq \frac{8C_{\text{init}}}{s}$. For the permutation $\Pi_{\mathcal{J}}$ in $\mathcal{A}$, we have

$$|\hat{\nu}_{1, \mathcal{J}} - \hat{\nu}_{2, \mathcal{J}}| = |\hat{\nu}_{\Pi_{\mathcal{J}}(1), \mathcal{J}} - \hat{\nu}_{\Pi_{\mathcal{J}}(2), \mathcal{J}}| \quad (\text{S.98})$$

$$\leq |\hat{\nu}_{\Pi_{\mathcal{J}}(1), \mathcal{J}} - \mu_{1, \mathcal{J}}^*| + |\hat{\nu}_{\Pi_{\mathcal{J}}(2), \mathcal{J}} - \mu_{2, \mathcal{J}}^*| + \|\boldsymbol{\mu}_1^* - \boldsymbol{\mu}_2^*\|_{\max} \quad (\text{S.99})$$

$$\leq \frac{10C_{\text{init}}}{s}. \quad (\text{S.100})$$

It follows that, for any permutation Π , we have

$$|\hat{\nu}_{\Pi(k), \mathcal{J}} - \mu_{k, \mathcal{J}}^*| \leq |\hat{\nu}_{\Pi(k), \mathcal{J}} - \hat{\nu}_{\Pi_{\mathcal{J}}(k), \mathcal{J}}| + |\hat{\nu}_{\Pi_{\mathcal{J}}(k), \mathcal{J}} - \mu_{k, \mathcal{J}}^*| \quad (\text{S.101})$$

$$\leq |\hat{\nu}_{1, \mathcal{J}} - \hat{\nu}_{2, \mathcal{J}}| + |\hat{\nu}_{\Pi_{\mathcal{J}}(k), \mathcal{J}} - \mu_{k, \mathcal{J}}^*| \quad (\text{S.102})$$

$$\leq \frac{11C_{\text{init}}}{s}. \quad (\text{S.103})$$

And thus our claim holds in this case with an arbitrary Π .

Case 2: $\|\boldsymbol{\mu}_1^* - \boldsymbol{\mu}_2^*\|_{\max} > \frac{8C_{\text{init}}}{s}$. Without loss of generality, assume that there exists a unique $\mathcal{J}_0^*$ such that $|\mu_{1, \mathcal{J}_0^*} - \mu_{2, \mathcal{J}_0^*}| = \|\boldsymbol{\mu}_1^* - \boldsymbol{\mu}_2^*\|_{\max}$ and for any $\mathcal{J} \neq \mathcal{J}_0^*$, we have $|\mu_{1, \mathcal{J}_0^*} - \mu_{2, \mathcal{J}_0^*}| - |\mu_{1, \mathcal{J}} - \mu_{2, \mathcal{J}}| > \frac{4C_{\text{init}}}{s}$.

Under $\mathcal{A}$, for any $\mathcal{J} \neq \mathcal{J}_0^*$, we have

$$|\tilde{\mu}_{1\mathcal{J}} - \tilde{\mu}_{2\mathcal{J}}| \leq |\tilde{\mu}_{\Pi_{\mathcal{J}}(1),\mathcal{J}} - \mu_{1\mathcal{J}}^*| + |\mu_{1\mathcal{J}}^* - \mu_{2\mathcal{J}}^*| + |\tilde{\mu}_{\Pi_{\mathcal{J}}(2),\mathcal{J}} - \mu_{2\mathcal{J}}^*| \quad (\text{S.104})$$

$$\leq |\mu_{1\mathcal{J}}^* - \mu_{2\mathcal{J}}^*| + \frac{2C_{\text{init}}}{s} \quad (\text{S.105})$$

$$< |\mu_{1\mathcal{J}_0^*} - \mu_{2\mathcal{J}_0^*}| - \frac{2C_{\text{init}}}{s}. \quad (\text{S.106})$$

Meanwhile,

$$|\tilde{\mu}_{1\mathcal{J}_0^*} - \tilde{\mu}_{2\mathcal{J}_0^*}| \geq |\mu_{1\mathcal{J}_0^*} - \mu_{2\mathcal{J}_0^*}| - \sum_k |\tilde{\mu}_{\Pi_{\mathcal{J}_0^*}(k),\mathcal{J}_0^*} - \mu_{k\mathcal{J}_0^*}| \geq |\mu_{1\mathcal{J}_0^*} - \mu_{2\mathcal{J}_0^*}| - \frac{2C_{\text{init}}}{s} \quad (\text{S.107})$$

It follows that $\mathcal{J}_0 = \mathcal{J}_0^*$. Also note that $|\hat{\mu}_{\Pi_{\mathcal{J}_0}(k),\mathcal{J}_0} - \mu_{k\mathcal{J}_0}^*| \leq \frac{C_{\text{init}}}{s}$.

For $\mathcal{J} \neq \mathcal{J}_0$, we have

$$|\hat{\mu}_{k\mathcal{J}}^{(0)} - \mu_{k\mathcal{J}}^*| = |\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}} - \mu_{k\mathcal{J}}^*| \quad (\text{S.108})$$

$$\leq |\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}} - \hat{\nu}_{\Upsilon_{\mathcal{J}}(k),\mathcal{J}}| + |\hat{\nu}_{\Upsilon_{\mathcal{J}}(k),\mathcal{J}} - \mu_{k\mathcal{J}}^*| \quad (\text{S.109})$$

$$\equiv L_1 + L_2. \quad (\text{S.110})$$

Under $\mathcal{A}$, we have $L_2 \leq \frac{C_{\text{init}}}{s}$. We further show that L_1 by proving that $\Lambda_{\mathcal{J}} = \Upsilon_{\mathcal{J}}$. Because $\Lambda_{\mathcal{J}}$ minimizes $\max_k \{|\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}|\}$, it suffices to show that

$$\max_k \{|\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}|\} < \max_k \{|\hat{\nu}_{\bar{\Lambda}_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}|\}, \quad (\text{S.111})$$

where $\bar{\Lambda}_{\mathcal{J}} = 3 - \Lambda_{\mathcal{J}}$ is the only other permutation from $\{1, 2\}$ to $\{1, 2\}$. Now note that (S.111) is implied by

$$|\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}| < |\hat{\nu}_{\bar{\Lambda}_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}| \text{ for } k = 1, 2. \quad (\text{S.112})$$

We show that (S.112) is true. Without loss of generality, we assume that there is no label switching in $\tilde{\mu}_{k\mathcal{J}}$, i.e., $|\tilde{\mu}_{k\mathcal{J}} - \mu_{k\mathcal{J}}^*| < |\tilde{\mu}_{k'\mathcal{J}} - \mu_{k'\mathcal{J}}^*|$ for $k' \neq k$. Note that the left-hand side

$$|\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}| \leq |\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \mu_{k\mathcal{J}_0}^*| + |\mu_{k\mathcal{J}_0}^* - \tilde{\mu}_{k\mathcal{J}_0}| \leq \frac{2C_{\text{init}}}{s}. \quad (\text{S.113})$$

For the right-hand side, we note that for any $k' \neq k$, we must have $\bar{\Lambda}_{\mathcal{J}}(k) = \Lambda_{\mathcal{J}}(k')$. Hence,

$$|\hat{\nu}_{\bar{\Lambda}_{\mathcal{J}}(k),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}| \quad (\text{S.114})$$

$$= |\hat{\nu}_{\Lambda_{\mathcal{J}}(k'),\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0}| = |(\hat{\nu}_{\Lambda_{\mathcal{J}}(k'),\mathcal{J}_0} - \mu_{k'\mathcal{J}_0}^*) + (\mu_{k'\mathcal{J}_0}^* - \mu_{k\mathcal{J}_0}) + (\mu_{k\mathcal{J}_0} - \tilde{\mu}_{k\mathcal{J}_0})| \quad (\text{S.115})$$

$$\geq |\mu_{1\mathcal{J}_0^*} - \mu_{2\mathcal{J}_0^*}| - \sum_k |\hat{\nu}_{\Lambda_{\mathcal{J}}(k),\mathcal{J}_0} - \mu_{k\mathcal{J}_0}^*| \quad (\text{S.116})$$

$$< \frac{2C_{\text{init}}}{s}. \quad (\text{S.117})$$

As a result, $\Lambda_{\mathcal{J}} = \Upsilon_{\mathcal{J}}$, and our claim is true in this case, too.

By Proposition S.11, we have $\Pr(\mathcal{A}) \geq 1 - \frac{1}{\prod_m p_m}$. Consequently, there exists a permutation Π such that

$$\Pr(\|\widehat{\boldsymbol{\mu}}_{\Pi(k)}^{(0)} - \boldsymbol{\mu}_k^*\|_{\max} \preceq \frac{1}{s}) \geq 1 - \frac{1}{\prod_m p_m}. \quad (\text{S.118})$$

The bound on $\widehat{\Sigma}_m^{(0)}$. By Proposition S.3, we have that

$$\left\| \bigotimes_{m=M}^{m=1} \widehat{\Sigma}_m^{(0)} - \bigotimes_{m=M}^{m=1} \Sigma_m^* \right\|_{\max} \preceq O(\|\widehat{\Sigma}_m^{(0)} - \Sigma_m\|_{\max}). \quad (\text{S.119})$$

Hence, it suffices to find a bound for the probability that $\|\widehat{\Sigma}_m - \Sigma_m\|_{\max} \preceq \frac{1}{s}$. Consider the event

$$\mathcal{B} = \{|\widehat{\text{var}}(X_{i,1\dots 1} \mid Y) - \text{var}(X_{1\dots 1} \mid Y)| \preceq \frac{1}{s}, \|\widehat{\Psi}_m^{(0)} - a_m \Sigma_m^*\|_{\max} \preceq \frac{1}{s}\}, \quad (\text{S.120})$$

where $a_m = \frac{1}{q_m} \mathbf{1}_{q_m}^\top \Sigma_{-m}^* \mathbf{1}_{q_m}$. By Lemma S.8,

$$\frac{1}{q_m} \mathbf{X}_{(m)} \mathbf{1}_{q_m} \sim \pi_1^* N(\phi_1, a_m \Sigma_m^*) + \pi_2^* N(\phi_2, a_m \Sigma_m^*), \quad (\text{S.121})$$

where $a_m = \frac{1}{q_m} \mathbf{1}_{q_m}^\top \Sigma_{-m}^* \mathbf{1}_{q_m}$. In other words, $\frac{1}{q_m} \mathbf{X}_{(m)} \mathbf{1}_{q_m}$ follows Model (A1). By letting $\epsilon = \frac{C}{s^2}$ and $\delta = \prod_m p_m^{-3}$ in Proposition S.11, with probability greater than $1 - C \prod_m p_m^{-1}$, we have $\Pr(\mathcal{B}) \preceq \frac{1}{s}$.

We claim that event $\mathcal{B}$ implies that $\|\widehat{\Sigma}_m - \Sigma_m^*\|_{\max} \preceq \frac{1}{s}$. We proceed to prove this claim. Because eigenvalues of Σ_m^* are bounded below and above, a_m is also bounded below and above. For ease of presentation, we let $\Psi_m = a_m \Sigma_m^*$. When we rescale $\widehat{\Psi}_m^{(0)}$ to $\widetilde{\Psi}_m$ such that the upper-left corner is equal to 1, we must have

$$\|\widetilde{\Psi}_m^{(0)} - (\Psi_{m,11})^{-1} \Psi_m\|_{\max} \quad (\text{S.122})$$

$$= \|(\widehat{\Psi}_{m,11})^{-1} \widehat{\Psi}_m^{(0)} - (\Psi_{m,11})^{-1} \Psi_m\|_{\max} \quad (\text{S.123})$$

$$\leq \|(\widehat{\Psi}_{m,11})^{-1} (\widehat{\Psi}_m^{(0)} - \Psi_m)\|_{\max} + |\widehat{\Psi}_{m,11}^{(0)} - \Psi_{m,11}^{-1}| \cdot \|\Psi_m\|_{\max} \quad (\text{S.124})$$

$$\preceq \|\widehat{\Psi}_m^{(0)} - \Psi_m\|_{\max} \preceq \frac{1}{s}. \quad (\text{S.125})$$

Because $\widehat{\Sigma}_m^{(0)} = (\widehat{\Psi}_{m,11}^{(0)})^{-1} \widehat{\Psi}_m^{(0)}$, $\Sigma_m^* = (\Psi_{m,11})^{-1} \Psi_m$ for $m > 1$, we have $\|\widehat{\Sigma}_m - \Sigma_m^*\|_{\max} \preceq \frac{1}{s}$ for $m > 1$.

For $m = 1$, we let $\sigma = \text{var}(X_{i,1\dots 1} \mid Y)$ and $\hat{\sigma} = \widehat{\text{var}}(X_{i,1\dots 1} \mid Y)$. Then

$$\|\hat{\Sigma}_1^{(0)} - \Sigma_1^*\|_{\max} = \|\hat{\sigma}\tilde{\Psi}_1^{(0)} - \Sigma_1^*\|_{\max} \quad (\text{S.126})$$

$$\leq |\hat{\sigma} - \sigma| \cdot \|\tilde{\Psi}_1^{(0)}\|_{\max} + \sigma \|\tilde{\Psi}_1^{(0)} - \sigma^{-1}\Sigma_1^*\|_{\max} \quad (\text{S.127})$$

$$\preceq \max\{|\hat{\sigma} - \sigma|, \|\tilde{\Psi}_1^{(0)} - \sigma^{-1}\Sigma_1^*\|_{\max}\} \preceq \frac{1}{s}. \quad (\text{S.128})$$

Therefore, with a probability greater than $1 - O(\prod_m p_m^{-1})$, we have that $\hat{\mu}_k^{(0)}, \bigotimes_{m=M}^{m=1} \hat{\Sigma}_m^{(0)}$ satisfy Condition (IC). By Proposition S.10 we have the desired conclusion.

References

- Cai, T. T., Ma, J. & Zhang, L. (2019), ‘Chime: Clustering of high-dimensional gaussian mixtures with em algorithm and its optimality’, *The Annals of Statistics* **47**(3), 1234–1267.
- Hardt, M. & Price, E. (2015), Tight bounds for learning a mixture of two gaussians, in ‘Proceedings of the forty-seventh annual ACM symposium on Theory of computing’, ACM, pp. 753–760.
- Kolda, T. G. & Bader, B. W. (2009), ‘Tensor decompositions and applications’, *SIAM Review* **51**(3), 455–500.
- Lyu, X., Sun, W. W., Wang, Z., Liu, H., Yang, J. & Cheng, G. (2019), ‘Tensor graphical model: Non-convex optimization and statistical inference’, *IEEE transactions on pattern analysis and machine intelligence*.
- Pan, Y., Mai, Q. & Zhang, X. (2019), ‘Covariate-adjusted tensor classification in high dimensions’, *Journal of the American statistical association* **114**(527), 1305–1319.
- Petersen, K. B. & Pedersen, M. S. (2008), ‘The matrix cookbook’, *Technical University of Denmark* **7**(15), 510.
- Wang, M. & Zeng, Y. (2019), Multiway clustering via tensor block models, in ‘Advances in Neural Information Processing Systems’, pp. 715–725.