

Supplementary Materials for “Covariate-Adjusted Tensor Classification in High Dimensions” Yuqing Pan, Qing Mai, and Xin Zhang

A Additional Simulation Results

Recall that we include various popular and state-of-the-art classification methods as competitors in Sections 6 & 7. We used the following R packages that implement our competing methods in numerical studies: *penalizedLDA* for ℓ_1 -FDA, *sparseLDA* for SOS, *glmnet* for ℓ_1 -GLM, *randomForest* for Random Forest and *penalizedSVM* for ℓ_1 -SVM. For other competitors, we downloaded R code from <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC4714801/#SD1> for MDA and PMDA, MATLAB code from <https://hua-zhou.github.io/TensorReg/> for CP-GLM, and MATLAB code from <http://www.scholat.com/portalPaperInfo.html?paperID=25520&Entry=laizhihui> for STDA. The author of (Li and Schonfeld, 2014) kindly provides the code for CMDA and DGTDA.

For methods with only one tuning parameter, we select the tuning parameter based on the error rates on validation set. For methods with multiple tuning parameters, we select them as follows. For random forest, we tuned the number of variables randomly selected as candidate in each split and used 400 ensembled trees. For STDA, we did a grid search on both the reduced tensor dimension and penalty parameter. For CMDA and DGTDA, we let the reduced tensor have the same dimension on each mode in simulation and ADHD, and we fixed the reduced dimension to be 10×2 in the CSA data.

A.1 Variable Selection for results for Models (C1–C3i)

In Table A.1 we compare the variable selection results of CATCH, ℓ_1 -FDA, SOS as well as ℓ_1 -GLM and ℓ_1 -SVM. For each methods, two tasks are considered: one with only tensor predictor $\mathbf{X}$ (or $\text{vec}(\mathbf{X})$), and one with both tensor $\mathbf{X}$ and covariates $\mathbf{U}$. It turns out that CATCH results in higher TPR and relatively lower FPR than others. Comparing the performance of methods with and without covariates, we found that CATCH with covariates has much higher TPR compared to CATCH without covariates, except in (C3a), where there is no dependence relationship between the covariates and the tensor. Meanwhile, in (C3b), the covariates have not influence on the classification by themselves, but help with the variable selection in $\mathbf{X}$ nonetheless.

A.2 Higher dimension tensors

To demonstrate the applicability of CATCH in higher dimensional tensor data, we further considered variants of models (C1)–(C3) where we increase the tensor dimensions from $30 \times 36 \times 30$ to $80 \times 80 \times 80$, that is 512,000 voxels in total. Parallel to models (C1)–(C3) in the paper, we now have three new models (C1H)–(C3H) where the sample size are also increased slightly from $n_k = 75$ to $n_k = 100$ in each class. Many methods (CMDA, DGTDA, RF, ℓ_1 -SVM) become practically inapplicable because of the prohibitive computational costs. Therefore, we only compare CATCH with ℓ_1 -FDA, ℓ_1 -logistic regression, GP-GLM and DGTDA. The classification errors are summarized in Table A.2, and the variable selection in Table A.3. Not surprisingly, CATCH continues to achieve better accuracy and variable selection than the competitors.

B The Autism Spectrum disorder (ASD) dataset

In this section we study a real data example similar to the ADHD data in the paper, where we have a three-way tensor predictor from structural magnetic resonance imaging (MRI) and additional covariates (age and sex). The original data is from the Autism Brain Imaging Data

Rate(%)	C1		C2		C3		C3a		C3b		C3i	
	TPR	FPR	TPR	FPR	TPR	FPR	TPR	FPR	TPR	FPR	TPR	FPR
CATCH	49.8	0.0	53.4	0.0	60.6	0.1	69.8	0.1	64.3	0.1	58.1	0.1
without U	(2.0)	(0.0)	(1.9)	(0.0)	(1.4)	(0.0)	(1.2)	(0.0)	(1.0)	(0.0)	(1.4)	(0.0)
CATCH	92.1	0.1	66.9	0.1	69.4	0.1	61.2	0.1	73.2	0.1	70.0	0.1
	(0.7)	(0.0)	(0.1)	(0.0)	(1.0)	(0.0)	(1.0)	(0.0)	(1.2)	(0.0)	(1.1)	(0.0)
ℓ_1 -FDA	96.0	33.9	100	25.8	100	24.0	100	21.8	100	20.9	97.9	28.8
without U	(2.0)	(2.0)	(0)	(1.5)	(0)	(1.5)	(0)	(0.8)	(0)	(0.9)	(1.4)	(2.5)
ℓ_1 -FDA	96.0	34.7	100	25.4	100	24.0	100	21.7	100	20.9	97.9	28.8
with U	(2.0)	(2.1)	(0)	(1.5)	(0)	(1.5)	(0)	(0.8)	(0)	(0.9)	(1.4)	(2.5)
SOS	26.8	0.0	55.8	0.1	53.4	0.1	61.3	0.1	54.8	0.1	58.0	0.1
without U	(1.0)	(0.0)	(1.0)	(0.0)	(1.1)	(0.0)	(1.1)	(0.0)	(1.0)	(0.0)	(1.1)	(0.0)
SOS	26.9	0.0	55.9	0.1	53.4	0.1	60.4	0.1	54.8	0.0	57.9	0.1
with U	(1.0)	(0.0)	(1.0)	(0.0)	(1.1)	(0.0)	(1.1)	(0.0)	(1.0)	(0.0)	(1.1)	(0.0)
ℓ_1 logistic	20.8	0.0	54.1	0.0	51.1	0.0	57.1	0.0	52.1	0.1	58.9	0.1
without U	(0.1)	(0.0)	(1.7)	(0.0)	(1.2)	(0.0)	(1.1)	(0.1)	(1.1)	(0.1)	(1.1)	(0.0)
ℓ_1 logistic	21.4	0.1	53.9	0.0	50.6	0.0	59.5	0.1	50.1	0.1	58.8	0.1
with U	(0.9)	(0.0)	(1.1)	(0.0)	(1.2)	(0.0)	(1.0)	(0.0)	(1.1)	(0.1)	(1.1)	(0.0)
ℓ_1 SVM	43.9	0.1	55.9	0.1	31.3	0.1	56.1	0.3	43.6	0.1	56.5	0.12
without U	(0.2)	(0.0)	(0.3)	(0.0)	(0.0)	(0.0)	(0.1)	(0.0)	(0.1)	(0.0)	(0.25)	(0.00)
ℓ_1 SVM	81.0	0.5	55.9	0.1	31.4	0	68.7	0.4	74.7	0.2	56.5	0.2
with U	(0.3)	(0.0)	(0.3)	(0.0)	(0.1)	(0.0)	(0.1)	(0.0)	(0.3)	(0.0)	(0.25)	(0.0)

Table A.1: Variable selection results for Models (C1–C3) & (C3a), (C3b), (C3i). True Positive Rate (TPR) and False Positive Rate (FPR) are defined in 6.1. Means and standard errors (in parentheses) are based on 100 replicates.

Exchange (ABIDE) (available at http://fcon_1000.projects.nitrc.org/indi/abide/), where the goal is to study the characterizations of Autism Spectrum disorder (ASD) among 23 laboratories. We use the dataset from New York University that contains $n = 182$ observations of 78 ASD subjects and 104 normal controls. For each subject, a $91 \times 109 \times 91$ tensor predictor is provided. We remove the fibers with all zero variables or limited nonzero variables, which lie on the outer shell of the tensor, so that the (nonzero) tensor predictor has dimension $71 \times 89 \times 71$. We record the 5-fold and 10-fold cross validation errors in Table B.1. Given the large dimension and slow speed of SOS, we complete 100 replicates for CATCH but only 20 replicates for SOS.

Error rate (%)	C1H		C2H		C3H	
	(X, U)	X	(X, U)	X	(X, U)	X
Bayes	5.58 (0.22)	NA NA	11.22 (0.32)	NA NA	8.34 (0.30)	NA NA
CATCH	10.03 (0.36)	30.33 (0.44)	16.11 (0.37)	20.01 (0.65)	10.58 (0.35)	17.12 (0.62)
CP-GLM	43.33 (1.01)	43.84 (1.20)	29.23 (1.15)	29.02 (1.28)	16.71 0.57	17.22 0.73
DGTDA	49.59 (0.50)	49.51 (0.48)	50.39 (0.52)	50.58 (0.52)	46.44 (0.45)	46.47 (0.45)
ℓ_1 -FDA	43.75 (0.47)	43.72 (0.47)	34.59 (0.51)	34.59 (0.51)	35.6 (0.52)	35.57 (0.51)
ℓ_1 -GLM	33.75 (0.53)	33.75 (0.53)	16.63 (0.38)	16.63 (0.38)	13.62 (0.35)	13.61 (0.35)

Table A.2: Higher dimension simulations. We apply these methods on using only tensor and both tensor and covariates two scenarios. Due to memory limitation, results are recorded from 100 replicates and each replicate has 100 test observations.

X	C1H		C2H		C3H	
	TPR	FPR	TPR	FPR	TPR	FPR
CATCH	69.5 (2.82)	0.01 (0.00)	56.63 (1.97)	0.00 (0.00)	57.8 (1.93)	0.00 (0.00)
ℓ_1 -FDA	98 (1.41)	51.18 (3.21)	100 (0.00)	24.46 (1.25)	100 (0.00)	45.61 (3.01)
ℓ_1 -GLM	19.5 (0.90)	0.00 (0.00)	58.63 (1.04)	0.00 (0.00)	55.81 (0.98)	0.00 (0.00)
(X, U)	C1H		C2H		C3H	
	TPR	FPR	TPR	FPR	TPR	FPR
CATCH	92.69 (0.73)	0.01 (0.00)	70.00 (1.12)	0.01 (0.00)	75.38 (1.12)	0.01 (0.00)
ℓ_1 -FDA	98 (1.41)	51.81 (3.22)	100 (0.00)	24.46 (1.25)	100 (0.00)	46.22 (3.02)
ℓ_1 -logistic	19.5 (0.90)	0.00 (0.00)	58.63 (1.04)	0.00 (0.00)	55.88 (0.98)	0.00 (0.00)

Table A.3: Variable selection results.

56 Classification of ASD in this dataset is more difficult than ADHD study, given the subtle
57 relationship between brain anatomy and autism. Classification of ASD has error rates over 40%

Error rate (%)	CATCH	SOS	Random guess
5 fold	43.89 (0.31)	48.57 (0.29)	49.53 (0.04)
10 fold	43.77 (0.29)	48.91 (0.31)	49.06 (0.04)

Table B.1: ASD data analysis. Error rates and standard error (in parentheses) of 5-fold and 10-fold cross validation.

across literature. It is shown that the error rate of SOS is very close to random guessing rate, while CATCH performs much better.

C Technical Details for Example 1 and some related results

First, we prove the following Proposition C.1. Although Proposition C.1 is a well-known result in the LDA literature, we still present its proof for completeness.

Proposition C.1. *For $Y \in \{1, 2\}$, $\mathbf{W} \in \mathbb{R}^p$, assume that $(\mathbf{W}, Y)$ follows a linear discriminant analysis model in the sense that $\mathbf{W} \mid Y = k \sim N(\boldsymbol{\zeta}_k, \boldsymbol{\Omega})$. Further assume that $\pi_1 = \pi_2 = 1/2$. Then the lowest classification error rate possible is*

$$1 - \Phi\left(\frac{\sqrt{(\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1)^\top \boldsymbol{\Omega}^{-1} (\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1)}}{2}\right) \quad (\text{C.1})$$

where Φ is the cumulative distribution function of the standard normal random variable.

Proof of Proposition C.1. It is well-known that the best classifier under this model classifies an observation to Class 2 if and only if $(\mathbf{W} - \frac{\boldsymbol{\zeta}_1 + \boldsymbol{\zeta}_2}{2})^\top \boldsymbol{\Omega}^{-1} (\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1) > 0$. It is easy to show that

$$\Pr((\mathbf{W} - \frac{\boldsymbol{\zeta}_1 + \boldsymbol{\zeta}_2}{2})^\top \boldsymbol{\Omega}^{-1} (\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1) > 0 \mid Y = 1) \quad (\text{C.2})$$

$$= \Pr((\mathbf{W} - \frac{\boldsymbol{\zeta}_1 + \boldsymbol{\zeta}_2}{2})^\top \boldsymbol{\Omega}^{-1} (\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1) < 0 \mid Y = 2) \quad (\text{C.3})$$

$$= 1 - \Phi\left(\frac{\sqrt{(\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1)^\top \boldsymbol{\Omega}^{-1} (\boldsymbol{\zeta}_2 - \boldsymbol{\zeta}_1)}}{2}\right) \quad (\text{C.4})$$

Hence, the classification error rate is

$$\Pr((\mathbf{W} - \frac{\zeta_1 + \zeta_2}{2})^\top \mathbf{\Omega}^{-1}(\zeta_2 - \zeta_1) > 0 \mid Y = 1) \Pr(Y = 1) \quad (\text{C.5})$$

$$+ \Pr((\mathbf{W} - \frac{\zeta_1 + \zeta_2}{2})^\top \mathbf{\Omega}^{-1}(\zeta_2 - \zeta_1) < 0 \mid Y = 2) \Pr(Y = 2) \quad (\text{C.6})$$

$$= 1 - \Phi\left(\frac{\sqrt{(\zeta_2 - \zeta_1)^\top \mathbf{\Omega}^{-1}(\zeta_2 - \zeta_1)}}{2}\right) \quad (\text{C.7})$$

□

Define $R(\mathbf{U}, \mathbf{X})$ as the lowest classification error rate possible if we build the classifier based on $\mathbf{X}$ and $\mathbf{U}$, and similarly, $R(\mathbf{X})$ as that based only on $\mathbf{X}$. For simplicity, we consider the binary classification case with equal class probability of 0.5, and assume that $\mu_1 \neq \mu_2$ because otherwise $(\mathbf{U}, \mathbf{X})$ are jointly independent of Y and the best classifier is random guessing.

Lemma C.1. For $K = 2$, $\phi_1 = \phi_2$, $\pi_1 = \pi_2$, $\alpha \neq 0$ and $\mu_1 \neq \mu_2$, we have $R(\mathbf{U}, \mathbf{X}) < R(\mathbf{X})$.

Proof of Lemma C.1. Straightforward calculation shows that

$$\begin{pmatrix} \mathbf{U} \\ \text{vec}(\mathbf{X}) \end{pmatrix} \mid (Y = k) \sim N\left(\begin{pmatrix} \phi_k \\ \mu_k + \alpha \bar{\mathbf{x}}_{(M+1)} \phi_k \end{pmatrix}, \begin{pmatrix} \mathbf{\Omega}_{11} & \mathbf{\Omega}_{12} \\ \mathbf{\Omega}_{21} & \mathbf{\Omega}_{22} \end{pmatrix}\right), \quad (\text{C.8})$$

where $\mathbf{\Omega}_{11} = \mathbf{\Psi}$, $\mathbf{\Omega}_{12} = \mathbf{\Omega}_{21}^\top = \mathbf{\Psi} \alpha_{(M+1)}$, $\mathbf{\Omega}_{22} = \otimes_{m=M}^{m=1} \mathbf{\Sigma}_m + \alpha_{(M+1)}^\top \mathbf{\Psi} \alpha_{(M+1)}$.

By Proposition C.1 and applying the block inverse formula, we have

$$R(\mathbf{U}, \mathbf{X}) = 1 - \Phi\left(\frac{\sqrt{\Delta_1 + \Delta_2}}{2}\right), \quad (\text{C.9})$$

where we have used the fact that $\phi_1 = \phi_2$, and,

$$\Delta_1 = (\mu_2 - \mu_1)^\top \mathbf{\Omega}_{22}^{-1} (\mu_2 - \mu_1) \quad (\text{C.10})$$

$$\Delta_2 = (\mu_2 - \mu_1)^\top \{\mathbf{\Omega}_{22}^{-1} \mathbf{\Omega}_{21} (\mathbf{\Omega}_{11} - \mathbf{\Omega}_{12} \mathbf{\Omega}_{22}^{-1} \mathbf{\Omega}_{21})^{-1} \mathbf{\Omega}_{12} \mathbf{\Omega}_{22}^{-1}\} (\mu_2 - \mu_1). \quad (\text{C.11})$$

It is easy to see that $\Delta_1 > 0$, $\Delta_2 > 0$ under our assumptions.

On the other hand, by (C.8), $(\text{vec}(\mathbf{X}), Y)$ also follows the LDA model. Hence, because of

$\phi_1 = \phi_2$,

$$R(\mathbf{X}) = 1 - \Phi\left(\frac{\sqrt{\Delta_1}}{2}\right) \quad (\text{C.12})$$

where Δ_1 is the same as defined in (C.10). Because $\Delta_2 > 0$, we must have $R(\mathbf{U}, \mathbf{X}) < R(\mathbf{X})$.

□

Technical Details for Example 1. From straightforward calculation, we have,

$$\Sigma_a \equiv \text{cov} \left(\begin{pmatrix} \text{vec}(\mathbf{X}) \\ \mathbf{U} \end{pmatrix} \mid Y \right) = \begin{pmatrix} 1 + \alpha^2 & \alpha^2 & 0 & 0 & \alpha \\ \alpha^2 & 1 + \alpha^2 & 0 & 0 & \alpha \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ \alpha & \alpha & 0 & 0 & 1 \end{pmatrix}.$$

Therefore, $\Sigma_b \equiv \text{cov}(\text{vec}(\mathbf{X}) \mid Y)$ is the upper left 4×4 block of Σ_a . It is thus easy to verify that,

$$\Sigma_a^{-1} = \begin{pmatrix} 1 & 0 & 0 & 0 & -\alpha \\ 0 & 1 & 0 & 0 & -\alpha \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ -\alpha & -\alpha & 0 & 0 & 1 + 2\alpha^2 \end{pmatrix}, \quad \Sigma_b^{-1} = \mathbf{I}_4 - (2 + \alpha^{-2})^{-1} \cdot \begin{pmatrix} 1 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Then from Proposition C.1 and Lemma C.1, we have, $R(\mathbf{X}, U) = R(X_{11}, U) = 1 - \Phi \left(\sqrt{(\Sigma_a^{-1})_{11}} \right) = 1 - \Phi(1)$; $R(X_{11}) = 1 - \Phi(\text{cov}^{-1}(X_{11} \mid Y)) = 1 - \Phi(1/\sqrt{1 + \alpha^2})$; because $\Sigma_b^{-1}(\text{E}(\text{vec}(\mathbf{X}) \mid Y = 2) - \text{E}(\text{vec}(\mathbf{X}) \mid Y = 1))$ contains two nonzero elements, $R(\mathbf{X}) = R(X_{11}, X_{21}) = 1 - \Phi \left(\sqrt{(\Sigma_b^{-1})_{11}} \right) = 1 - \Phi(\sqrt{2 + \alpha^2}/\sqrt{2 + 2\alpha^2})$. □

D Proofs for all Lemmas and Propositions

Proof of Proposition 1. First, note that (2.5) is equivalent to

$$\hat{Y} = \arg \max_{k=1, \dots, K} \{ \log(\pi_k/\pi_1) + \log \frac{f_k(\mathbf{U}, \mathbf{X})}{f_1(\mathbf{U}, \mathbf{X})} \} \quad (\text{D.1})$$

Under the model in (2.1), we have that, within Class k , the probability density function for $(\mathbf{U}, \mathbf{X})$ is

$$f_k(\mathbf{U}, \mathbf{X}) = \frac{1}{\det(2\pi\mathbf{\Psi})^{1/2}} \cdot \exp\left(-\frac{1}{2}(\mathbf{U} - \phi_k)^\top \mathbf{\Psi}^{-1}(\mathbf{U} - \phi_k)\right) \cdot \frac{1}{\det(2\pi \otimes_{m=M}^{m=1} \Sigma_m)^{1/2}} \cdot \exp\left(-\frac{1}{2}\text{vec}^\top(\mathbf{X} - \alpha \bar{\times}_{(M+1)} \mathbf{U} - \mu_k)(\otimes_{m=M}^{m=1} \Sigma_m)^{-1}\text{vec}(\mathbf{X} - \alpha \bar{\times}_{(M+1)} \mathbf{U} - \mu_k)\right) \quad (\text{D.2})$$

We have the desired conclusion by substituting (D.2) into (D.1). □

Proof of Lemma 1. Under the CATCH model, the negative log-likelihood up to additive and scaling constants can be written as

$$\begin{aligned} \ell(\boldsymbol{\alpha}, \boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_K, \boldsymbol{\Sigma}) \\ \simeq \log |\boldsymbol{\Sigma}| + n^{-1} \sum_{k=1}^K \sum_{Y^i=k} \text{tr} \left\{ \boldsymbol{\Sigma}^{-1} \text{vec} \left(\mathbf{X}^i - \boldsymbol{\mu}_k - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i \right) \text{vec}^\top \left(\mathbf{X}^i - \boldsymbol{\mu}_k - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i \right) \right\}, \end{aligned}$$

where $\boldsymbol{\Sigma} = \boldsymbol{\Sigma}_M \otimes \dots \otimes \boldsymbol{\Sigma}_1$. Then we can first see that the MLE for $\boldsymbol{\mu}_k$ (since it only appears in the log-likelihood of class k) is

$$\hat{\boldsymbol{\mu}}_k = \bar{\mathbf{X}}_k - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \bar{\mathbf{U}}_k,$$

then the partially minimized $\ell(\boldsymbol{\alpha}, \boldsymbol{\mu}_1, \dots, \boldsymbol{\mu}_K, \boldsymbol{\Sigma})$ becomes

$$\begin{aligned} \ell(\boldsymbol{\alpha}, \boldsymbol{\Sigma}) \simeq \log |\boldsymbol{\Sigma}| \\ + n^{-1} \sum_{i=1}^n \text{tr} \left\{ \boldsymbol{\Sigma}^{-1} \text{vec} \left(\tilde{\mathbf{X}}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \tilde{\mathbf{U}}^i \right) \text{vec}^\top \left(\tilde{\mathbf{X}}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \tilde{\mathbf{U}}^i \right) \right\}, \end{aligned}$$

where the second term no longer depends on k explicitly and the centered data is defined as $\tilde{\mathbf{X}}^i = \mathbf{X}^i - \bar{\mathbf{X}}_k$ and $\tilde{\mathbf{U}}^i = \mathbf{U}^i - \bar{\mathbf{U}}_k$ for observation i in class k . Also, notice that

$$\text{vec} \left(\tilde{\mathbf{X}}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \tilde{\mathbf{U}}^i \right) = \text{vec}(\tilde{\mathbf{X}}^i) - \boldsymbol{\alpha}_{(M+1)}^\top \tilde{\mathbf{U}}^i.$$

It is thus clear that the likelihood function under the CATCH model coincides with that of the tensor response regression model (Li and Zhang, 2017) with the modified $\tilde{\mathbf{X}}$ as response and the modified $\tilde{\mathbf{Z}}$ as predictor and with the same form of (partial) likelihood. Therefore, we have the desired conclusion. $\square$

Proof of Lemma 2. Write $\mathbf{W} = \llbracket \mathbf{Z}; \boldsymbol{\Omega}_1^{1/2}, \dots, \boldsymbol{\Omega}_m^{1/2} \rrbracket$. Then we have that all the entries in $\mathbf{Z}$ are independent of each other.

$$\mathbb{E}(\mathbf{W}_{(j)} \mathbf{W}_{(j)}^\top) = \boldsymbol{\Omega}_j^{1/2} \left\{ \mathbb{E}(\mathbf{Z}_{(j)} (\bigotimes_{l \neq j} \boldsymbol{\Omega}_l) \mathbf{Z}_{(j)}^\top) \right\} \boldsymbol{\Omega}_j^{1/2} \quad (\text{D.3})$$

107 Define $\mathbf{V} = \bigotimes_{l \neq j} \boldsymbol{\Omega}_l$ and $\mathbf{H} = \mathbb{E} \left\{ \mathbf{Z}_{(j)} (\bigotimes_{l \neq j} \boldsymbol{\Omega}_l) \mathbf{Z}_{(j)}^\top \right\}$. Then we have

$$h_{ij} = \mathbb{E} \left(\sum_{k,l} v_{kl} Z_{ik} Z_{jl} \right) = \begin{cases} \sum_i v_{ii} & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases} \quad (\text{D.4})$$

108 It follows that $\mathbf{H} = \text{tr}(\mathbf{V}) \cdot \mathbf{I}$. Then because of $\text{tr}(\mathbf{V}) = \prod_{l \neq j} \text{tr}(\boldsymbol{\Omega}_l)$, we have the desired
109 conclusion.

110 $\square$

111 **Lemma D.1.** For each $j = 1, \dots, p$, we can define a sequence $\{s_m, m = 2, \dots, M+1\}$ as
112 follows: $s_{M+1} = j$ and $s_m = s_{m+1} \bmod (\prod_{i=1}^{m-1} p_i)$ for $m = 2, \dots, M$. Then

$$\left(\bigotimes_{m=M}^1 \boldsymbol{\Sigma}_m \right)_{\cdot j} = \text{vec} \left(\bigotimes_{m=M}^1 \boldsymbol{\Sigma}_{m, j_m} \right) \quad (\text{D.5})$$

113 where $j_1 = s_2 \bmod p_1$ and $j_m = \lceil \frac{s_{m+1}}{\prod_{i=1}^{m-1} p_i} \rceil$ for $m = 2, \dots, M$.

114 **Proof of Lemma D.1.** By the definition of Kronecker product, we have that, for two matrices
115 $\boldsymbol{\Omega}_1 \in \mathbb{R}^{q_1}, \boldsymbol{\Omega}_2 \in \mathbb{R}^{q_2}$, we have

$$(\boldsymbol{\Omega}_2 \otimes \boldsymbol{\Omega}_1)_{\cdot j} = \text{vec}(\boldsymbol{\Omega}_{2, \lceil \frac{j}{p_1} \rceil} \otimes \boldsymbol{\Omega}_{1, (j \bmod p_1)}). \quad (\text{D.6})$$

116 where we restrict that $a \bmod b = b$ if a is divisible by b . Hence, we have

$$(\boldsymbol{\Sigma}_M \otimes \dots \otimes \boldsymbol{\Sigma}_1)_{\cdot j} = \{ \boldsymbol{\Sigma}_M \otimes (\boldsymbol{\Sigma}_{M-1} \otimes \dots \otimes \boldsymbol{\Sigma}_1) \}_{\cdot j} \quad (\text{D.7})$$

$$= \boldsymbol{\Sigma}_{M, j_M} \otimes (\boldsymbol{\Sigma}_{M-1} \otimes \dots \otimes \boldsymbol{\Sigma}_1)_{\cdot s_{M-1}} \quad (\text{D.8})$$

$$= \bigotimes_{m=M}^1 \boldsymbol{\Sigma}_{m, j_m}, \quad (\text{D.9})$$

117 where the last equality is obtained by recursively applying (D.6). $\square$

118 **Proof of Lemma 3.** For simplicity, we only present the proofs for no covariate case. Without
119 loss of generality, assume that the observations are arranged such that $Y^i = 1$ for $i = 1, \dots, n_1$;
120 $Y^i = 2$ for $i = n_1 + 1, \dots, n_1 + n_2$ and so on.

121 Define $\mathbf{X}_{(m)}$ as the $np_{-m} \times p_m$ matrix obtained by stacking the observations $\mathbf{X}_{(m)}^i, i =$
 122 $1, \dots, n$.

$$\mathbf{A} = \begin{pmatrix} (\mathbf{I}_{n_1} - \frac{1}{n_1} \mathbf{1}_{n_1} \mathbf{1}_{n_1}^\top) \otimes \mathbf{I}_{p-m} & \cdots & 0 \\ & \ddots & \\ 0 & \cdots & (\mathbf{I}_{n_K} - \frac{1}{n_K} \mathbf{1}_{n_K} \mathbf{1}_{n_K}^\top) \otimes \mathbf{I}_{p-m} \end{pmatrix} \quad (\text{D.10})$$

123 By (3.6), $\hat{\Sigma}_m \propto (\mathbf{X}_{(m)}^*)^\top \mathbf{X}_{(m)}^*$, where

$$\mathbf{X}_{(m)}^* = \begin{pmatrix} \mathbf{X}_{(m)}^1 - \hat{\boldsymbol{\mu}}_1 \\ \vdots \\ \mathbf{X}_{(m)}^K - \hat{\boldsymbol{\mu}}_K \end{pmatrix} = \mathbf{A} \mathbf{X}_{(m)} \quad (\text{D.11})$$

124 It follows that $\hat{\Sigma}_m$ is positive semi-definite, and $\text{rank}(\hat{\Sigma}) = \text{rank}(\mathbf{X}_{(m)}^*)$. It leaves to show that
 125 $\text{rank}(\mathbf{X}_{(m)}^*) = p_m$.

126 Because $\text{rank}(\mathbf{A}) = (n - K)p_{-m}$, while $\mathbf{X}_{(m)}$ is randomly sampled from a continuous
 127 distribution, we have that, when $(n - K)p_{-m} > p_m$, with a probability of 1, $\mathbf{X}_{(m)}^*$ is column
 128 full-rank, i.e., $\text{rank}(\mathbf{X}_{(m)}^*) = p_m$. Consequently, with a probability of 1, $\hat{\Sigma}_m$ is full rank. $\square$

129 **Proof of Lemma 4.** Set the partial derivative of $\mathcal{L}(\mathbf{C}_2, \dots, \mathbf{C}_K)$ with respect to $\mathbf{C}_k$ to be zero
 130 and we have

$$\frac{\partial \mathcal{L}(\mathbf{C}_2, \dots, \mathbf{C}_K)}{\partial \mathbf{C}_k} = 2\llbracket \mathbf{C}_k; \Sigma_2, \dots, \Sigma_K \rrbracket - 2(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) = 0 \quad (\text{D.12})$$

131 Solve this equation for $k = 2, \dots, K$ and we have that $(\mathbf{B}_2, \dots, \mathbf{B}_K)$ minimizes $\mathcal{L}(\mathbf{C}_2, \dots, \mathbf{C}_K)$.
 132 $\square$

133 **Proof of Lemma 5.** By Lemma 1 in Mai et al. (2017), we have that, given $\boldsymbol{\beta}_{j'}, j' \neq j$, the
 134 solution of $\boldsymbol{\beta}_{j.}$ to (4.1) is

$$\arg \min_{\boldsymbol{\beta}_{j.}} \sum_{k=2}^K \frac{1}{2} (\beta_{kj} - \tilde{\beta}_{kj})^2 + \frac{\lambda}{\hat{\sigma}_{jj}} \|\boldsymbol{\beta}_{j.}\| \quad (\text{D.13})$$

135 where $\tilde{\beta}_{kj} = \frac{(\hat{\nu}_{kj} - \hat{\nu}_{1j}) - \sum_{l \neq j} \hat{\sigma}_{lj} \beta_{kl}}{\hat{\sigma}_{jj}}$. Because $\hat{\Sigma} = \hat{\Sigma}_M \otimes \dots \otimes \hat{\Sigma}_1$, by Lemma D.1 we
 136 have

$$\sum_{l \neq j} \hat{\sigma}_{lj} \beta_{kl} = \llbracket \mathbf{B}_k^j; \hat{\Sigma}_{1,j_1}, \dots, \hat{\Sigma}_{M,j_M} \rrbracket, \hat{\sigma}_{jj} = \prod_{m=1}^M \hat{\sigma}_{m,j_m j_m} \quad (\text{D.14})$$

137 $j_m = \lceil \frac{s_{m+1}}{\prod_{i=1}^{m-1} p_i} \rceil$ for $m = 2, \dots, M$ and $j_1 = s_2 \bmod p_1$. And we have the first part of
 138 conclusion, (4.5, 4.6). By checking the subgradients we can easily obtain the second part of
 139 conclusion, (4.7). $\square$

140 **Proof of Lemma 6.** We first consider the computational cost for (4.8). For each j , by Lemma D.1,
 141 we need $M - 1$ operations to find $j_1, \dots, j_M$. Additional $M - 1$ operations are needed to find
 142 $\prod_{m=1}^M \hat{\sigma}_{m,j_m j_m}$. It leaves to calculate $\llbracket \mathbf{B}_k^j; \hat{\Sigma}_{1,j_1}, \dots, \hat{\Sigma}_{M,j_M} \rrbracket$, which is equal to $(\otimes_{m=M}^1 \hat{\Sigma}_{m,j_m}^\top) \text{vec}(\mathbf{B}_k^j)$.
 143 Since $\mathbf{B}_k^j$ is sparse with at most d_t nonzero entries, it suffices to find the d_t corresponding el-
 144 ements in $(\otimes_{m=M}^1 \hat{\Sigma}_{m,j_m}^\top)$. For any j' , it takes $M - 1$ operations to find $j'_m, m = 1, \dots, M$
 145 such that $\prod_{m=1}^M \hat{\sigma}_{m,j'_m j_m} = (\otimes_{m=M}^1 \hat{\Sigma}_{m,j_m}^\top)_{j'}$. The calculation of $\prod_{m=1}^M \hat{\sigma}_{m,j'_m j_m}$ is again
 146 $O(M - 1)$. Hence, the calculation of $\llbracket \mathbf{B}_k^j; \hat{\Sigma}_{1,j_1}, \dots, \hat{\Sigma}_{M,j_M} \rrbracket, k = 2, \dots, K$ is $O(MKd_t)$.
 147 The computational cost for soft-thresholding is lower. Hence, for each j , the computational
 148 cost is $O(MKd_t)$. The update for all j has the computational cost of $O(NKd_t p)$. $\square$

149 E Proof for Theorem 1 (Convergence analysis for Algorithm 1)

150 In this section, we present the results in Theorem 1.

151 Our proofs reply on the following definitions and propositions in Tseng (2001). For a
 152 generic objective function

$$f(\mathbf{x}_1, \dots, \mathbf{x}_N) = f_0(\mathbf{x}_1, \dots, \mathbf{x}_N) + \sum_{k=1}^N f_k(\mathbf{x}_k) \quad (\text{E.1})$$

153 where $f_0 : \mathbb{R}^{n_1 + \dots + n_N} \mapsto \mathbb{R} \cup \{\infty\}$ and some $f_k : \mathbb{R}^{n_k} \mapsto \mathbb{R} \cup \{\infty\}, k = 1, \dots, N$.

154 We have the following definitions:

155 **Definition E.1** (Gâteaux-differentiable (Bertsekas, 2016)). *For a function $F : \mathbb{R}^p \mapsto \mathbb{R}$, its*
 156 *Gâteaux derivative is defined as*

$$F'(\mathbf{x}; \mathbf{y}) = \lim_{\lambda \searrow 0} \frac{F(\mathbf{x} + \lambda \mathbf{y}) - F(\mathbf{x})}{\lambda} \quad (\text{E.2})$$

157 *If $F'(\mathbf{x}; \mathbf{y})$ exists for all $\mathbf{y}$ at $\mathbf{x}$, then F is Gâteaux-differentiable at $\mathbf{x}$.*

Definition E.2 (Stationary points (Tseng, 2001)). A point $\mathbf{z}$ is stationary for function h if $g(\mathbf{z}; \mathbf{d}) \geq 0$ for any $\mathbf{d}$, where

$$g(\mathbf{z}; \mathbf{d}) = \liminf_{\lambda \searrow 0} \frac{h(\mathbf{x} + \lambda \mathbf{d}) - h(\mathbf{x})}{\lambda} \quad (\text{E.3})$$

Definition E.3 (Regular (Tseng, 2001)). A function f is regular at $\mathbf{z}$ if

$$f'(\mathbf{z}; \mathbf{d}) \geq 0, \text{ for any } \mathbf{d} = (\mathbf{d}_1, \dots, \mathbf{d}_N) \text{ such that } f'(\mathbf{z}; (0, \dots, \mathbf{d}_k, \dots, 0)) \geq 0, k = 1, \dots, N. \quad (\text{E.4})$$

We will also use the following propositions.

Proposition E.1 (A simplified version of Lemma 3.1 in Tseng (2001)). If f_0 is Gâteaux-differentiable, f is regular at each $\mathbf{z}$.

Proposition E.2 (A simplified version of Theorem 4.1 in Tseng (2001)). Assume that the level set $\mathbf{X}^0 = \{\mathbf{x} : f(\mathbf{x}) \leq f(\mathbf{x}^0)\}$ is compact, where $\mathbf{x}^0$ is the initial value of the algorithm, and that f is continuous on $\mathbf{X}^0$. Further assume that $f(\mathbf{x}_1, \dots, \mathbf{x}_N)$ is pseudoconvex in $(\mathbf{x}_k, \mathbf{x}_i)$ for all $i, k \in \{1, \dots, N\}$, and if f is regular at every $\mathbf{x}$, then the solution generated by the cyclic coordinate descent method converges to a stationary point of f .

We make use of these results to prove Theorem 1.

Lemma E.1. If $\widehat{\Sigma}$ is positive definite, the set $\mathcal{B}^0 = \{\boldsymbol{\beta} : L(\boldsymbol{\beta}) \leq L(\boldsymbol{\beta}^0)\}$ is compact for any $\boldsymbol{\beta}^0$, where

$$L(\boldsymbol{\beta}) = \sum_{k=2}^K \boldsymbol{\beta}_k^\top \widehat{\Sigma} \boldsymbol{\beta}_k - 2\boldsymbol{\beta}_k^\top (\boldsymbol{\nu}_k - \boldsymbol{\nu}_1) + \lambda \sum_{j=1}^p \|\boldsymbol{\beta}_{\cdot j}\| \quad (\text{E.5})$$

Proof. Because $L(\boldsymbol{\beta})$ is continuous, $\mathcal{B}^0$ must be closed. It suffices to verify that $\mathcal{B}^0$ is bounded.

When $\widehat{\Sigma}$ is invertible, it is easy to see that, for any $\boldsymbol{\beta}_k \in \mathbb{R}^p$,

$$\boldsymbol{\beta}_k^\top \widehat{\Sigma} \boldsymbol{\beta}_k - 2(\widehat{\boldsymbol{\nu}}_k - \widehat{\boldsymbol{\nu}}_1)^\top \boldsymbol{\beta}_k \geq -(\widehat{\boldsymbol{\nu}}_k - \widehat{\boldsymbol{\nu}}_1)^\top \widehat{\Sigma}^{-1} (\widehat{\boldsymbol{\nu}}_k - \widehat{\boldsymbol{\nu}}_1) \quad (\text{E.6})$$

It follows that, for any $\beta \in \mathcal{B}^0$,

$$-\sum_{k=2}^K (\hat{\nu}_k - \hat{\nu}_1)^\top \hat{\Sigma}^{-1} (\hat{\nu}_k - \hat{\nu}_1) + \lambda \sum_{j=1}^p \|\beta_{\cdot j}\| \leq L(\beta) \leq L(\beta^0) \quad (\text{E.7})$$

which implies that

$$\sum_{j=1}^p \|\beta_{\cdot j}\| \leq \frac{1}{\lambda} \sum_{k=2}^K (\hat{\nu}_k - \hat{\nu}_1)^\top \hat{\Sigma}^{-1} (\hat{\nu}_k - \hat{\nu}_1) + L(\beta^0) \text{ for any } \beta \in \mathcal{B}^0. \quad (\text{E.8})$$

Therefore, $\mathcal{B}^0$ is bounded and hence compact. $\square$

Proof of Theorem 1. Let

$$f_0(\beta) = \sum_{k=2}^K \{\beta_k^\top \hat{\Sigma} \beta_k - 2(\nu_k - \nu_1)^\top \beta_k\}, f_j(\beta_{\cdot j}) = \lambda \|\beta_{\cdot j}\|, j = 1, \dots, p \quad (\text{E.9})$$

Apparently $L(\beta) = f_0(\beta) + \sum_{j=1}^p f_j(\beta_{\cdot j})$. Because $f_0(\beta)$ is differentiable, by Proposition E.1, we have that $L(\beta)$ is regular at each β . By Lemma 3, if $(n - K)p_{-m} > p_m$ for any m , with a probability of 1, $\hat{\Sigma}$ is positive definite. Combine this fact with Lemma E.1 and we have $\mathcal{B}^0 = \{\beta : L(\beta) \leq L(\beta^0)\}$ is compact. Further note that $L(\beta)$ is continuous and convex. By Proposition E.2, the blockwise coordinate descent algorithm converges to a stationary point of $L(\beta)$. Because $L(\beta)$ is strictly convex, the stationary point is the global minimizer. $\square$

F Technical Lemmas for the proof of Theorems 2 and 3

Our proofs rely on the following result that is well-known in the literature; see Bickel and Levina (2004) for example.

Lemma F.1. *For bivariate normal random variables $(Z_1^i, Z_2^i)_{i=1}^n$, i.i.d. with mean zero and bounded variances, there exists $\epsilon_0 > 0$ such that for any ϵ , $0 < \epsilon < \epsilon_0$, with a probability greater than $1 - 2 \exp(-Cn\epsilon^2)$, we have*

$$\left| \frac{1}{n} \sum_{i=1}^n Z_1^i - \mathbb{E} Z_1 \right| \leq \epsilon \quad (\text{F.1})$$

$$\left| \frac{1}{n} \sum_{i=1}^n Z_1^i Z_2^i - \mathbb{E}(Z_1^i Z_2^i) \right| \leq \epsilon \quad (\text{F.2})$$

190 In what follows, we prove some large deviation inequalities for matrix/tensor normal distri-
 191 butions.

192 **Lemma F.2.** For $\mathbf{Z}^i \in \mathbb{R}^{p_1 \times p_2}$, $\mathbf{Z}^i \sim TN(0, \mathbf{I}_{p_1}, \mathbf{I}_{p_2})$, $i = 1, \dots, n$, i.i.d., and $\mathbf{\Lambda} \in \mathbb{R}^{p_1 \times p_1}$ a
 193 diagonal matrix with diagonal elements $\lambda_1, \dots, \lambda_{p_1} \geq 0$, there exists $\epsilon_0 > 0$ such that for any
 194 $0 < \epsilon < \epsilon_0$, we have

$$\Pr(\|\frac{1}{np_1} \sum_{i=1}^n (\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i - \frac{1}{p_1} \mathbb{E}((\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i)\|_{\max} > \epsilon) \leq 2p_2^2 \exp(-\frac{Cnp_1^2\epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2}) \quad (\text{F.3})$$

195 **Proof of Lemma F.2.** Note that

$$\|\frac{1}{np_1} \sum_{i=1}^n (\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i - \frac{1}{p_1} \mathbb{E}((\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i)\|_{\max} \quad (\text{F.4})$$

$$= \max_{s,k} |\frac{1}{np_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i - \mathbb{E}(\frac{1}{p_1} \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i)| \equiv \max_{s,k} L(s, k) \quad (\text{F.5})$$

196 We analyze $L(s, k)$ under two scenarios: $s = k$ and $s \neq k$. We first consider the case $s = k$.

$$L(s, s) = |\frac{1}{np_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 - \mathbb{E}(\frac{1}{p_1} \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2)| \quad (\text{F.6})$$

197 Because $(Z_{ls}^i)^2 \sim \chi^2(1)$, we have, for any $t > 0$, $\mathbb{E}e^{t\lambda_l(Z_{ls}^i)^2} = (1 - 2\lambda_l t)^{-1/2}$. Thanks to
 198 the independence among Z_{ls}^i for (i, l) , we have

$$\mathbb{E}\{\exp(t \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2)\} = \prod_{l=1}^{p_1} (1 - 2\lambda_l t)^{-n/2} \quad (\text{F.7})$$

199 It follows from Chernoff bound that

$$\Pr\left(\frac{1}{np_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 - \mathbb{E}\left(\frac{1}{p_1} \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2\right) \geq \epsilon\right) \quad (\text{F.8})$$

$$\leq \left(\prod_{l=1}^{p_1} (1 - 2\lambda_l t)^{-n/2}\right) \exp(-tn \sum_{l=1}^{p_1} \lambda_l - tnp_1\epsilon) \quad (\text{F.9})$$

200 It can be verified that, for $0 < t < \frac{1}{4\lambda_l}$, we have

$$(1 - 2\lambda_l t)^{-1/2} \leq \exp(\lambda_l t + 4\lambda_l^2 t^2) \quad (\text{F.10})$$

201 It follows that

$$\Pr \left(\frac{1}{np_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 - E \left(\frac{1}{p_1} \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 \right) \geq \epsilon \right) \quad (\text{F.11})$$

$$\leq \exp \left(\sum_{l=1}^{p_1} 4\lambda_l^2 n t^2 - t n p_1 \epsilon \right) \quad (\text{F.12})$$

$$= \exp \left(-\frac{n p_1^2 \epsilon^2}{16 \sum_{l=1}^{p_1} \lambda_l^2} \right) \quad (\text{F.13})$$

202 where the last equality is obtained by letting $t = \frac{p_1 \epsilon}{8 \sum_{l=1}^{p_1} \lambda_l^2}$.

Similarly, we can show that

$$\Pr \left(\frac{1}{p_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 - E \left(\sum_{l=1}^{p_1} \lambda_l (Z_{ls}^i)^2 \right) \leq -\epsilon \right) \leq \exp \left(-\frac{C n p_1^2 \epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2} \right).$$

203 Hence, for any s , $\Pr(L(s, s) \geq \epsilon) \leq 2 \exp \left(-\frac{C n p_1^2 \epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2} \right)$.

204 Now we consider the case $s \neq k$, where $L(s, k) = \left| \frac{1}{n p_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i \right|$.

205 Note that

$$E[\exp(t \lambda_l Z_{ls}^i Z_{lk}^i)] = E \{ E[\exp(t \lambda_l Z_{ls}^i Z_{lk}^i) \mid Z_{lk}^i] \} \quad (\text{F.14})$$

$$= E \left\{ \exp \left(\frac{t^2 \lambda_l^2 (Z_{lk}^i)^2}{2} \right) \right\} = \frac{1}{\sqrt{1 - \lambda_l^2 t^2}} \quad (\text{F.15})$$

206 By the Chernoff bound, we have that

$$\Pr \left(\frac{1}{n p_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i \geq \epsilon \right) \leq \left\{ \prod_{l=1}^{p_1} (1 - \lambda_l^2 t^2)^{-n/2} \right\} \exp(-t n p_1 \epsilon) \quad (\text{F.16})$$

207 It can be verified that, if $t^2 \leq \frac{1}{2\lambda_l^2}$, then we have

$$(1 - \lambda_l^2 t^2)^{-1/2} \leq \exp(2\lambda_l^2 t^2) \quad (\text{F.17})$$

208 It follows that

$$\Pr \left(\frac{1}{n p_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i \geq \epsilon \right) \leq \exp(2n \sum_{l=1}^{p_1} \lambda_l^2 t^2 - t n p_1 \epsilon) = \exp \left(-\frac{n p_1^2 \epsilon^2}{8 \sum_{l=1}^{p_1} \lambda_l^2} \right) \quad (\text{F.18})$$

209 if we take $t = \frac{p_1 \epsilon}{4 \sum_{l=1}^{p_1} \lambda_l^2}$.

210 Similarly, we can show that

$$\Pr\left(\frac{1}{np_1} \sum_{i=1}^n \sum_{l=1}^{p_1} \lambda_l Z_{ls}^i Z_{lk}^i \leq -\epsilon\right) \leq \exp\left(-\frac{Cnp_1^2 \epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2}\right) \quad (\text{F.19})$$

211 Consequently, $\Pr(L(s, k) > \epsilon) \leq 2 \exp\left(-\frac{Cnp_1^2 \epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2}\right)$ for any $s \neq k$.

212 Therefore, apply the union bound and we have

$$\Pr\left(\left\|\frac{1}{np_1} \sum_{i=1}^n (\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i - \frac{1}{p_1} E((\mathbf{Z}^i)^\top \mathbf{\Lambda} \mathbf{Z}^i)\right\|_{\max} > \epsilon\right) \leq \Pr(\max_{s,k} L(s, k) > \epsilon) \quad (\text{F.20})$$

$$\leq 2p_2^2 \exp\left(-\frac{Cnp_1^2 \epsilon^2}{\sum_{l=1}^{p_1} \lambda_l^2}\right) \quad (\text{F.21})$$

□

214 **Lemma F.3.** Consider $\mathbf{W}^1, \dots, \mathbf{W}^n \in \mathbb{R}^{p_1 \times \dots \times p_M}$ i.i.d. with the distribution $\text{TN}(0, \Sigma_1, \Sigma_2, \dots, \Sigma_M)$,
 215 where diagonal elements of Σ_m , $m = 1, \dots, M$, are all ones. Define

$$\hat{\Sigma}_j = \frac{1}{np_{-j}} \sum_{i=1}^n (\mathbf{W}_{(j)}^i)(\mathbf{W}_{(j)}^i)^\top \quad (\text{F.22})$$

$$\hat{\boldsymbol{\mu}}_{(j)} = \frac{1}{n} \sum_{i=1}^n \mathbf{W}_{(j)}^i \quad (\text{F.23})$$

216 where $p_{-j} = \prod_{m \neq j} p_m$. Then there exists $\epsilon_0 > 0$ such that for any $0 < \epsilon < \epsilon_0$, with probability
 217 greater than $1 - 2p_j^2 \exp\left(-\frac{Cnp_{-j}^2 \epsilon^2}{\sum_{l=1}^{p_{-j}} \lambda_l^2}\right)$, we have

$$\|\hat{\Sigma}_j - \Sigma_j\|_{\max} \leq \|\Sigma_j^{1/2}\|_1^2 \epsilon \quad (\text{F.24})$$

218 In the meantime, for $2/n < \epsilon < \epsilon_0$, with a probability greater than $1 - 2p^2 \exp(-Cn^2 p_{-j} \epsilon^2)$,
 219 we have

$$\left\|\frac{1}{p_{-j}} \hat{\boldsymbol{\mu}}_{(j)} \hat{\boldsymbol{\mu}}_{(j)}^\top\right\|_{\max} \leq \lambda_{\max}^{M-1} \|\Sigma_j^{1/2}\|_1^2 \epsilon \quad (\text{F.25})$$

220 where $\lambda_{\max}$ is the largest eigenvalue of $\Sigma_1, \dots, \Sigma_M$.

221 **Proof of Lemma F.3.** By definition, we have

$$\mathbf{W}^i = \llbracket \mathbf{Z}^i; \Sigma_1^{1/2}, \dots, \Sigma_M^{1/2} \rrbracket \quad (\text{F.26})$$

222 where $\mathbf{Z}^i \sim TN(0, \mathbf{I}_{p_1}, \dots, \mathbf{I}_{p_M})$. Further decompose $\otimes_{m \neq j} \Sigma_m = \Gamma \Lambda \Gamma^\top$, where Λ is a diagonal matrix with diagonal elements $\lambda_l, l = 1, \dots, p_j$ and Γ is an orthonormal matrix. It is easy
223 to see that $\lambda_l \leq \lambda_{\max}^{M-1}$.

224 We first show (F.24). Note that

$$\hat{\Sigma}_j = \Sigma_j^{1/2} \left(\frac{1}{np-j} \sum_{i=1}^n (\mathbf{Z}_{(j)}^i) (\otimes_{m \neq j} \Sigma_m) (\mathbf{Z}_{(j)}^i)^\top \right) \Sigma_j^{1/2} \quad (\text{F.27})$$

$$= \Sigma_j^{1/2} \left(\frac{1}{np-j} \sum_{i=1}^n (\tilde{\mathbf{Z}}_{(j)}^i) \Lambda (\tilde{\mathbf{Z}}_{(j)}^i)^\top \right) \Sigma_j^{1/2} \quad (\text{F.28})$$

226 where $\tilde{\mathbf{Z}}_{(j)}^i = \mathbf{Z}_{(j)}^i \Gamma$. By properties of matrix normal distribution, we have $\tilde{\mathbf{Z}}_{(j)}^i \sim \mathbf{TN}(0, \mathbf{I}_{p_j}, \mathbf{I}_{p-j})$,
227 and hence is still a matrix with i.i.d standard normal random variables. Consequently,

$$\|\hat{\Sigma}_j - \Sigma_j\|_{\max} \leq \|\Sigma_j^{1/2}\|_1^2 \cdot \left\| \frac{1}{np-j} \sum_{i=1}^n (\tilde{\mathbf{Z}}_{(j)}^i) \Lambda (\tilde{\mathbf{Z}}_{(j)}^i)^\top - E(\tilde{\mathbf{Z}}_{(j)}^i) \Lambda (\tilde{\mathbf{Z}}_{(j)}^i)^\top \right\|_{\max} \quad (\text{F.29})$$

By Lemma F.2, we have

$$\Pr\left(\left\| \frac{1}{np-j} \sum_{i=1}^n (\tilde{\mathbf{Z}}_{(j)}^i) \Lambda (\tilde{\mathbf{Z}}_{(j)}^i)^\top - \frac{1}{p-j} E(\tilde{\mathbf{Z}}_{(j)}^i) \Lambda (\tilde{\mathbf{Z}}_{(j)}^i)^\top \right\|_{\max} > \epsilon\right) \leq 2p_j^2 \exp\left(-\frac{Cnp_j^2 \epsilon^2}{\sum_{l=1}^{p_j} \lambda_l^2}\right)$$

228 And the conclusion follows.

229 Now we show (F.25). Note that

$$\hat{\boldsymbol{\mu}}_{(j)} \hat{\boldsymbol{\mu}}_{(j)}^\top = \Sigma_j^{1/2} \left(\frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{Z}}_{(j)}^i \right) (\otimes_{m \neq j} \Sigma_m) \left(\frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{Z}}_{(j)}^i \right)^\top \Sigma_j^{1/2} \quad (\text{F.30})$$

$$= \Sigma_j^{1/2} \left(\frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{Z}}_{(j)}^i \right) \Lambda \left(\frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{Z}}_{(j)}^i \right)^\top \Sigma_j^{1/2} \quad (\text{F.31})$$

$$= \Sigma_j^{1/2} \tilde{\boldsymbol{\mu}}_{(j)} \Lambda \tilde{\boldsymbol{\mu}}_{(j)}^\top \Sigma_j^{1/2} \quad (\text{F.32})$$

230 where $\tilde{\boldsymbol{\mu}}_{(j)} = \frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{z}}_{(j)}^i$. Therefore,

$$\|\hat{\boldsymbol{\mu}}_{(j)} \hat{\boldsymbol{\mu}}_{(j)}^\top\|_{\max} \leq \|\boldsymbol{\Sigma}_j^{1/2}\|_1^2 \|\tilde{\boldsymbol{\mu}}_{(j)} \boldsymbol{\Lambda} \tilde{\boldsymbol{\mu}}_{(j)}^\top\|_{\max} = \|\boldsymbol{\Sigma}_1^{1/2}\|_1^2 \max_{s,k} \left| \sum_{l=1}^{p-j} \lambda_l \tilde{\mu}_{(j),sl} \tilde{\mu}_{(j),kl} \right| \quad (\text{F.33})$$

$$\leq \|\boldsymbol{\Sigma}_j^{1/2}\|_1^2 \max_{s,k} \frac{1}{2} \left(\sum_{l=1}^{p-j} \lambda_l \tilde{\mu}_{(j),sl}^2 + \sum_{l=1}^{p-j} \lambda_l \tilde{\mu}_{(j),kl}^2 \right) \quad (\text{F.34})$$

$$\leq \|\boldsymbol{\Sigma}_j^{1/2}\|_1^2 \max_s \sum_{l=1}^{p-j} \lambda_l \tilde{\mu}_{(j),sl}^2 \leq \lambda_{\max}^{M-1} \|\boldsymbol{\Sigma}_j^{1/2}\|_1^2 \max_s \sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2 \quad (\text{F.35})$$

231 For the term $\sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2$, note that $n \sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2 \sim \chi^2(p-j)$. Then by the Chernoff bound,
 232 there exists $\epsilon_0 > 0$ such that for $2/n < \epsilon < \epsilon_0$, we have

$$\Pr\left(\frac{1}{p-j} \sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2 > \epsilon\right) = \Pr\left(n \sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2 - p-j > np-j\epsilon - p-j\right) \quad (\text{F.36})$$

$$\leq \Pr\left(n \sum_{l=1}^{p-j} \tilde{\mu}_{(j),sl}^2 - p-j > np-j\epsilon/2\right) \leq \exp(-Cn^2 p-j \epsilon^2) \quad (\text{F.37})$$

233 And the conclusion follows. $\square$

234 **Lemma F.4.** For matrices $\boldsymbol{\Sigma}_m, \hat{\boldsymbol{\Sigma}}_m \in \mathbb{R}^{p_m}$, define $\sigma_{\max} = \max_{m,j,k} |\sigma_{m,j,k}|$. If $\|\hat{\boldsymbol{\Sigma}}_m -$
 235 $\boldsymbol{\Sigma}_m\|_{\max} \leq \epsilon \leq \max_{m,j,k} |\sigma_{m,j,k}|$ for $m = 1, \dots, M$, then we have

$$\|\hat{\boldsymbol{\Sigma}} - \boldsymbol{\Sigma}\|_{\max} \leq M 2^M \sigma_{\max}^{M-1} \epsilon \quad (\text{F.38})$$

236 where $\hat{\boldsymbol{\Sigma}} = \otimes_{m=M}^1 \hat{\boldsymbol{\Sigma}}_m$ and $\boldsymbol{\Sigma} = \otimes_{m=M}^1 \boldsymbol{\Sigma}_m$.

Proof of Lemma F.4.

$$\|\hat{\boldsymbol{\Sigma}} - \boldsymbol{\Sigma}\|_{\max} = \max_{j_m, k_m, m=1, \dots, M} \left| \prod_{m=1}^M \hat{\sigma}_{m, j_m, k_m} - \prod_{m=1}^M \sigma_{m, j_m, k_m} \right| \quad (\text{F.39})$$

$$= \max_{j_m, k_m, m=1, \dots, M} \left| \sum_{m=1}^M \left(\prod_{l=1}^{m-1} \sigma_{l, j_l, k_l} \right) (\hat{\sigma}_{m, j_m, k_m} - \sigma_{m, j_m, k_m}) \left(\prod_{l=m+1}^M \hat{\sigma}_{l, j_l, k_l} \right) \right| \quad (\text{F.40})$$

$$\leq \max_{j_m, k_m, m=1, \dots, M} \sum_{m=1}^M \left| \prod_{l \neq m} \max\{|\hat{\sigma}_{l, j_l, k_l}|, |\sigma_{l, j_l, k_l}|\} (\hat{\sigma}_{m, j_m, k_m} - \sigma_{m, j_m, k_m}) \right| \quad (\text{F.41})$$

$$\leq M 2^M \sigma_{\max}^{M-1} \epsilon \quad (\text{F.42})$$

237 $\square$

238 **Lemma F.5.** Consider random variables $(W, \mathbf{V})$ where $W \in \mathbb{R}$, $\mathbf{V} \in \mathbb{R}^q$ and $\mathbf{V} \sim N(0, \mathbf{\Omega})$,
 239 $W \mid \mathbf{V} = \mathbf{v} \sim N(\boldsymbol{\alpha}^\top \mathbf{v}, \sigma^2)$. Given n independent samples $\{W^i, \mathbf{V}^i\}_{i=1}^n$, define

$$\hat{\boldsymbol{\alpha}} = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{V}^i (\mathbf{V}^i)^\top \right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{V}^i W^i \right) \quad (\text{F.43})$$

240 Then

$$\Pr \left(\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}\|_{\max} > 2\|\mathbf{\Omega}^{-1}\|_1 (1 + \|\mathbf{\Omega}^{-1}\|_1 \cdot \|\boldsymbol{\delta}\|_\infty) \epsilon \right) < 2q^2 \exp\left(-\frac{Cn\epsilon^2}{q^2}\right) + 2q \exp(-2Cn\epsilon^2) \quad (\text{F.44})$$

241 where $\boldsymbol{\delta} = E(\mathbf{V}W)$.

242 **Proof Lemma F.5.** Write $\hat{\mathbf{\Omega}} = \frac{1}{n} \sum_{i=1}^n \mathbf{V}^i (\mathbf{V}^i)^\top$ and $\hat{\boldsymbol{\delta}} = \frac{1}{n} \sum_{i=1}^n \mathbf{V}^i W^i$. By Lemma F.1, for
 243 any (j, k) , we have

$$\Pr(|\hat{\omega}_{jk} - \omega_{jk}| > \epsilon) \leq 2 \exp(-Cn\epsilon^2) \quad (\text{F.45})$$

$$\Pr(|\hat{\delta}_k - \delta_k| > \epsilon) \leq 2 \exp(-Cn\epsilon^2) \quad (\text{F.46})$$

244 Consequently,

$$\Pr(\|\hat{\mathbf{\Omega}} - \mathbf{\Omega}\|_1 > \epsilon) \leq 2q^2 \exp\left(-\frac{Cn\epsilon^2}{q^2}\right) \quad (\text{F.47})$$

$$\Pr(\|\hat{\boldsymbol{\delta}} - \boldsymbol{\delta}\|_\infty > \epsilon) \leq 2q \exp(-Cn\epsilon^2) \quad (\text{F.48})$$

245 It follows that

$$\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}\|_\infty = \|\hat{\mathbf{\Omega}}^{-1} \hat{\boldsymbol{\delta}} - \mathbf{\Omega}^{-1} \boldsymbol{\delta}\|_\infty = \|\hat{\mathbf{\Omega}}^{-1} (\hat{\boldsymbol{\delta}} - \boldsymbol{\delta}) + (\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}) \boldsymbol{\delta}\|_\infty \quad (\text{F.49})$$

$$\leq \|\hat{\mathbf{\Omega}}^{-1}\|_1 \|\hat{\boldsymbol{\delta}} - \boldsymbol{\delta}\|_\infty + \|\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}\|_1 \cdot \|\boldsymbol{\delta}\|_\infty \quad (\text{F.50})$$

$$\leq (\|\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}\|_1 + \|\mathbf{\Omega}^{-1}\|_1) \|\hat{\boldsymbol{\delta}} - \boldsymbol{\delta}\|_\infty + \|\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}\|_1 \|\boldsymbol{\delta}\|_\infty \quad (\text{F.51})$$

246 Now note that

$$\|\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}\|_1 = \|\hat{\mathbf{\Omega}}^{-1} (\mathbf{\Omega} - \hat{\mathbf{\Omega}}) \mathbf{\Omega}^{-1}\|_1 \leq \|\hat{\mathbf{\Omega}}^{-1}\|_1 \cdot \|\mathbf{\Omega} - \hat{\mathbf{\Omega}}\|_1 \cdot \|\mathbf{\Omega}^{-1}\|_1 \quad (\text{F.52})$$

$$\leq (\|\hat{\mathbf{\Omega}}^{-1} - \mathbf{\Omega}^{-1}\|_1 + \|\mathbf{\Omega}^{-1}\|_1) \|\mathbf{\Omega} - \hat{\mathbf{\Omega}}\|_1 \cdot \|\mathbf{\Omega}^{-1}\|_1 \quad (\text{F.53})$$

247 It follows that if $\|\widehat{\Omega} - \Omega\|_1 < \epsilon$ for $\epsilon < \frac{1}{2\|\Omega^{-1}\|_1}$, we must have

$$\|\widehat{\Omega}^{-1} - \Omega^{-1}\|_1 \leq 2\|\Omega^{-1}\|_1^2 \epsilon \quad (\text{F.54})$$

248 Hence, under the event $\|\widehat{\Omega} - \Omega\|_1 < \epsilon$, $\|\widehat{\delta} - \delta\|_\infty < \epsilon$, we have

$$\|\widehat{\alpha} - \alpha\|_{\max} \leq 2\|\Omega^{-1}\|_1(1 + \|\Omega^{-1}\|_1 \cdot \|\delta\|_\infty)\epsilon. \quad (\text{F.55})$$

249

□

250 **Lemma F.6.** Consider $(\mathbf{V}, \mathbf{W})$ where $\mathbf{V} \in \mathbb{R}^q$, $\mathbf{W} \in \mathbb{R}^{p_1 \times \dots \times p_M}$. Assume that $\mathbf{V} \sim N(0, \Omega)$,
 251 $\mathbf{W} \mid \mathbf{V} = \mathbf{v} \sim TN(\alpha \bar{\times}_{(M+1)} \mathbf{v}, \Sigma_1, \dots, \Sigma_M)$. Given n independent samples $\{\mathbf{V}^i, \mathbf{W}^i\}_{i=1}^n$,
 252 define

$$\widehat{\alpha}_{j_1 \dots j_M} = \left(\frac{1}{n} \sum_{i=1}^n \mathbf{V}^i (\mathbf{V}^i)^\top \right)^{-1} \left(\frac{1}{n} \sum_{i=1}^n \mathbf{V}^i W_{j_1 \dots j_M}^i \right) \quad (\text{F.56})$$

$$\widehat{\Sigma}_j(\widehat{\alpha}) = \frac{1}{np-j} \sum_{i=1}^n (\mathbf{W}^i - \widehat{\alpha} \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\mathbf{W}^i - \widehat{\alpha} \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \quad (\text{F.57})$$

$$\widehat{\Sigma}_j(\alpha) = \frac{1}{np-j} \sum_{i=1}^n (\mathbf{W}^i - \alpha \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\mathbf{W}^i - \alpha \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \quad (\text{F.58})$$

253 Then we have

$$\Pr(\|\widehat{\Sigma}_j(\widehat{\alpha}) - \widehat{\Sigma}_j(\alpha)\|_{\max} > q^2 \epsilon) \leq C p q^2 \exp(-C \frac{n\epsilon}{q^2}) \quad (\text{F.59})$$

254 **Proof of Lemma F.6.** Define events:

$$\mathcal{E}_1 = \{\|\widehat{\alpha}_{j_1 \dots j_M} - \alpha_{j_1 \dots j_M}\|_\infty \leq \epsilon \text{ for all } \{j_1, \dots, j_M\}\} \quad (\text{F.60})$$

$$\mathcal{E}_2 = \{\max_{s,l} \frac{1}{n} \sum_{i=1}^n (W_{(j),sl}^i)^2 \leq 1 + \epsilon\} \quad (\text{F.61})$$

$$\mathcal{E}_3 = \{\max_l \frac{1}{n} \sum_{i=1}^n (V_l^i)^2 \leq 1 + \epsilon\} \quad (\text{F.62})$$

255 By Lemma F.5, $\Pr(\mathcal{E}_1) \geq 1 - 2p(q^2 \exp(-Cn \frac{\epsilon^2}{q^2}) + q \exp(-2Cn\epsilon^2))$. By Lemma F.1, we
 256 have $\Pr(\mathcal{E}_2) \geq 1 - p \exp(-Cn\epsilon^2)$, $\Pr(\mathcal{E}_3) \geq 1 - q \exp(-Cn\epsilon^2)$. Hence,

$$\Pr(\mathcal{E}_1 \cap \mathcal{E}_2 \cap \mathcal{E}_3) \geq 1 - 2C p q^2 \exp(-Cn \frac{\epsilon^2}{q^2}) \quad (\text{F.63})$$

257 In what follows we assume that $\mathcal{E}_1 \cap \mathcal{E}_2 \cap \mathcal{E}_3$ has happened.

258 Straightforward calculation shows that

$$\begin{aligned}
& \widehat{\Sigma}_j(\widehat{\alpha}) - \widehat{\Sigma}_j(\alpha) \\
= & \frac{1}{np_{-j}} \sum_{i=1}^n \mathbf{W}_{(j)}^i ((\alpha - \widehat{\alpha}) \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \\
& + \frac{1}{np_{-j}} \sum_{i=1}^n ((\alpha - \widehat{\alpha}) \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\mathbf{W}_{(j)}^i)^\top \\
& + \left(\frac{1}{np_{-j}} \sum_{i=1}^n (\widehat{\alpha} \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\widehat{\alpha} \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top - \frac{1}{np_{-j}} \sum_{i=1}^n (\alpha \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\alpha \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \right) \\
\equiv & L_1 + L_2 + L_3
\end{aligned}$$

259 We analyze these three terms one by one.

$$\begin{aligned}
\|L_1\|_{\max} &= \max_{s,k} \left| \frac{1}{np_{-j}} \sum_{i=1}^n \sum_{l=1}^{p-j} W_{(j),sl}^i [(\alpha - \widehat{\alpha}) \bar{\times}_{(M+1)} \mathbf{V}^i]_{(j),kl} \right| \\
&\leq \max_{k_1, \dots, k_M} \|\alpha_{k_1 \dots k_M} - \widehat{\alpha}_{k_1 \dots k_M}\|_\infty \max_{s,l} \frac{1}{np_{-j}} \sum_{i=1}^n \sum_{l=1}^{p-j} |W_{(j),sl}^i| \cdot \|\mathbf{V}^i\|_1 \\
&\leq \max_{k_1, \dots, k_M} \|\alpha_{k_1 \dots k_M} - \widehat{\alpha}_{k_1 \dots k_M}\|_\infty \max_{s,l} \frac{1}{np_{-j}} \sum_{i=1}^n \sum_{l=1}^{p-j} |W_{(j),sl}^i| \cdot \sqrt{q} \|\mathbf{V}^i\|_2 \\
&\leq \max_{k_1, \dots, k_M} \|\alpha_{k_1 \dots k_M} - \widehat{\alpha}_{k_1 \dots k_M}\|_\infty \max_{s,l} \frac{1}{np_{-j}} \sum_{i=1}^n \sum_{l=1}^{p-j} \{(W_{(j),sl}^i)^2/2 + q \|\mathbf{V}^i\|_2^2\} \\
&\leq \epsilon \left(\max_{s,l} \frac{1}{n} \sum_{i=1}^n (W_{(j),sl}^i)^2 + \max_l \frac{q^2}{n} \sum_{i=1}^n (V_l^i)^2 \right) \\
&\leq Cq^2\epsilon
\end{aligned}$$

260 Hence, $\|L_1\|_{\max} \leq Cq^2\epsilon$. Meanwhile, $L_2 = L_1^\top$. Hence, $\|L_2\|_{\max} = \|L_1\|_{\max} \leq Cq^2\epsilon$.

For L_3 , note that

$$L_3 = \frac{1}{np_{-j}} \sum_{i=1}^n ((\hat{\alpha} - \alpha) \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} (\hat{\alpha} \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \quad (\text{F.64})$$

$$+ \frac{1}{np_{-j}} \sum_{i=1}^n (\alpha \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)} ((\hat{\alpha} - \alpha) \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j)}^\top \quad (\text{F.65})$$

$$\equiv L_4 + L_5 \quad (\text{F.66})$$

$$\|L_4\|_{\max} = \max_{k,s} \frac{1}{np_{-j}} \left| \sum_{i=1}^n \sum_{l=1}^{p-j} ((\hat{\alpha} - \alpha) \bar{\times}_{(M+1)} \mathbf{V}^i)_{(j),kl} \cdot \hat{\alpha}_{(j),sl}^\top \mathbf{V}^i \right| \quad (\text{F.67})$$

$$\leq \max_{k_1, \dots, k_M} \|\hat{\alpha}_{k_1, \dots, k_M} - \alpha_{k_1, \dots, k_M}\|_\infty \cdot \left(\max_{k_1, \dots, k_M} \|\alpha_{k_1, \dots, k_M}\|_\infty + \epsilon \right) \quad (\text{F.68})$$

$$\cdot \frac{1}{np_{-j}} \sum_{i=1}^n \sum_{l=1}^{p-j} \|\mathbf{V}^i\|_1^2 \quad (\text{F.69})$$

$$\leq \epsilon \left(\max_{k_1, \dots, k_M} \|\alpha_{k_1, \dots, k_M}\|_\infty + \epsilon \right) \max_l \frac{q^2}{n} \sum_{i=1}^n (V_l^i)^2 \quad (\text{F.70})$$

$$\leq 2q^2 \epsilon \left(\max_{k_1, \dots, k_M} \|\alpha_{k_1, \dots, k_M}\|_\infty + \epsilon \right) \quad (\text{F.71})$$

$$\leq Cq^2 \epsilon \quad (\text{F.72})$$

Similarly we can provide a bound for $\|L_5\|_{\max}$. Putting all the pieces together, we have the desired conclusion. $\square$

G Proofs for Theorems 1 and 2

Now we consider the CATCH model in (2.1, 2.2). Throughout this section, we assume that $n_k > \frac{\pi_k n}{2}$, where $n_k = \sum_{i=1}^n 1(Y^i = k)$. Note that this event happens with a probability greater than $1 - 2K \exp(-Cn\pi_{\min}^2)$, as

$$\Pr(n_k < \frac{\pi_k n}{2}, k = 1, \dots, K) \leq \sum_{k=1}^K \Pr\left(\left|\frac{n_k}{n} - \pi_k\right| > \frac{\pi_k}{2}\right) \quad (\text{G.1})$$

$$\leq 2 \sum_{k=1}^K \exp(-Cn\pi_k^2) \leq 2K \exp(-Cn\pi_{\min}^2) \quad (\text{G.2})$$

268 where we apply Hoeffding's inequality.

269 Define $\hat{\boldsymbol{\mu}}_k = \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i)$, $\hat{\boldsymbol{\Sigma}}_j^{(k)} = \frac{1}{n_k p_{-j}} \sum_{Y^i=k} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i -$
 270 $\hat{\boldsymbol{\mu}}_k)_{(j)} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i - \hat{\boldsymbol{\mu}}_k)_{(j)}^\top$. We have the following lemma.

271 **Lemma G.1.** For $0 < \epsilon < \min\{\epsilon_0, \frac{\pi_{\min}}{2}\}$, we have

$$\Pr(\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}\|_{\max} > \epsilon) \leq 2q^2 p \exp(-\frac{Cn\epsilon^2}{q^2}) + 2K \exp(-Cn\pi_{\min}^2) \quad (\text{G.3})$$

272 **Proof of Lemma G.1.** The proof is similar to that of Lemma F.5 and is hence omitted here. $\square$

273 **Lemma G.2.** With a probability greater than $1 - Cq^2 p \exp(-C\frac{n\epsilon^2}{q^2}) - 2K \exp(-Cn\pi_{\min}^2)$, we
 274 have

$$\|\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k\|_{\max} \leq (2 + \|\boldsymbol{\phi}_k\|_{\infty})q\epsilon \text{ for } k = 1, \dots, K \quad (\text{G.4})$$

275 **Proof of Lemma G.2.** Note that

$$\|\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k\|_{\max} \quad (\text{G.5})$$

$$\leq \left\| \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i) - \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i) \right\|_{\max} \quad (\text{G.6})$$

$$+ \left\| \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i) - \boldsymbol{\mu}_k \right\|_{\max} \quad (\text{G.7})$$

$$\leq \|(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}) \bar{\times}_{(M+1)} \boldsymbol{\phi}_k\|_{\max} + \|(\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}) \bar{\times}_{(M+1)} (\boldsymbol{\phi}_k - \hat{\boldsymbol{\phi}}_k)\|_{\max} \quad (\text{G.8})$$

$$+ \left\| \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i) - \boldsymbol{\mu}_k \right\|_{\max} \quad (\text{G.9})$$

$$\leq (2 + \|\boldsymbol{\phi}_k\|_{\infty})q\epsilon \quad (\text{G.10})$$

276 under the event $n_k > \frac{\pi_k n}{2}$, $\|\hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha}\|_{\max} \leq \epsilon$, $\|\hat{\boldsymbol{\phi}}_k - \boldsymbol{\phi}_k\|_{\max} \leq \epsilon$, $\left\| \frac{1}{n_k} \sum_{Y^i=k} (\mathbf{X}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i) - \right.$
 277 $\left. \boldsymbol{\mu}_k \right\|_{\max} \leq \epsilon$. Combine Lemma G.1 with Lemma F.1 and we have the desired conclusion.

278 $\square$

279 **Lemma G.3.** There exists $\epsilon_0 > 0$ such that for $\frac{4}{\pi_k n} < \epsilon < \epsilon_0$, we have

$$\Pr(\|\hat{\boldsymbol{\Sigma}}_j^{(k)} - \boldsymbol{\Sigma}_j\|_{\max} > Cq^2 \epsilon) \leq 2p_j^2 \exp(-Cnp_{-j}\epsilon^2) + Cq^2 p \exp(-\frac{Cn\epsilon^2}{q^2}) + 2K \exp(-Cn\pi_{\min}^2) \quad (\text{G.11})$$

280 **Proof of Lemma G.3.** For simplicity, write

$$\tilde{\mathbf{E}}^i = \mathbf{X}^i - \boldsymbol{\alpha} \bar{\times}_{(M+1)} \mathbf{U}^i - \boldsymbol{\mu}_k \quad (\text{G.12})$$

281 Throughout this proof, we assume that the following events have happened:

$$\mathcal{E}_4 = \left\{ \left\| \frac{1}{np_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i (\tilde{\mathbf{E}}_{(j)}^i)^\top - \boldsymbol{\Sigma}_j \right\|_{\max} < \epsilon \right\} \quad (\text{G.13})$$

$$\mathcal{E}_5 = \left\{ \left\| \hat{\boldsymbol{\alpha}} - \boldsymbol{\alpha} \right\|_{\max} < \epsilon \right\} \quad (\text{G.14})$$

$$\mathcal{E}_6 = \left\{ \frac{1}{n_k} \sum_{Y^i=k} (\tilde{\mathbf{E}}_{(j),lt}^i)^2 < 1 + \epsilon, \text{ for any } l, t \right\} \quad (\text{G.15})$$

$$\mathcal{E}_7 = \left\{ \frac{1}{n_k} \sum_{Y^i=k} \|\mathbf{U}^i\|_2^2 \leq (1 + \epsilon)q \right\} \quad (\text{G.16})$$

$$\mathcal{E}_8 = \left\{ \left\| \hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k \right\|_{\max} < \epsilon, \text{ for any } k \right\} \quad (\text{G.17})$$

$$\mathcal{E}_9 = \left\{ \left\| \hat{\boldsymbol{\phi}}_k - \boldsymbol{\phi}_k \right\|_{\max} < \epsilon, \text{ for any } k \right\} \quad (\text{G.18})$$

When $n_k > \frac{\pi_k n}{2}$, by Lemma F.3, $\Pr(\mathcal{E}_4^C) \leq 2p_j^2 \exp(-Cnp_{-j}\epsilon^2)$. By Lemma G.1, $\Pr(\mathcal{E}_5^C) \leq 2q^2 p \exp(-\frac{Cn\epsilon^2}{q^2})$. By standard Chernoff bound, $\Pr(\mathcal{E}_6^C) \leq 2p \exp(-Cn\epsilon^2)$, $\Pr(\mathcal{E}_7^C) \leq 2q \exp(-Cn\epsilon^2)$, $\Pr(\mathcal{E}_8^C) \leq 2p \exp(-Cn\epsilon^2)$, $\Pr(\mathcal{E}_9^C) \leq 2q \exp(-Cn\epsilon^2)$. Hence, all these events happen simultaneously with a probability greater than

$$1 - 2p_j^2 \exp(-Cnp_{-j}\epsilon^2) - 2q^2 p \exp(-\frac{Cn\epsilon^2}{q^2}) - Cp \exp(-Cn\epsilon^2) - 2K \exp(-Cn\pi_{\min}^2)$$

$$\begin{aligned} \hat{\Sigma}_j^{(k)} &= \frac{1}{n_k p_{-j}} \sum_{Y^i=k} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i - \hat{\boldsymbol{\mu}}_k)_{(j)} (\mathbf{X}^i - \hat{\boldsymbol{\alpha}} \bar{\times}_{(M+1)} \mathbf{U}^i - \hat{\boldsymbol{\mu}}_k)_{(j)}^\top \\ &= \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i (\tilde{\mathbf{E}}_{(j)}^i)^\top + \frac{2}{n_k p_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i + (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k))_{(j)}^\top \\ &\quad + \frac{1}{n_k p_{-j}} \sum_{Y^i=k} ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i + (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k))_{(j)} ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i + (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k))_{(j)}^\top \\ &\equiv L_5 + L_6 + L_7 \end{aligned}$$

282

Under events $\mathcal{E}_4\text{--}\mathcal{E}_9$, we have, for L_5 ,

$$\|L_5 - \Sigma_j\|_{\max} \leq \epsilon. \quad (\text{G.19})$$

283

For L_6 ,

$$\|L_6\|_{\max} \leq \left\| \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i]_{(j)} \right\|_{\max} \quad (\text{G.20})$$

$$+ \left\| \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i (\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k)_{(j)}^\top \right\|_{\max} \quad (\text{G.21})$$

284

We analyze these two terms separately. First,

$$\left\| \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i]_{(j)} \right\|_{\max} \quad (\text{G.22})$$

$$= \max_{l,s} \left| \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \sum_{t=1} \tilde{E}_{(j),lt}^i [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i]_{(j),st} \right| \quad (\text{G.23})$$

$$\leq \|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max} \max_{l,s} \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \sum_{t=1} |\tilde{E}_{(j),lt}^i| \cdot \|\mathbf{U}^i\|_{(j),st} \quad (\text{G.24})$$

$$\leq \|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max} \cdot \max_{l,s} \left\{ \frac{1}{n_k p_{-j}} \sum_{Y^i=k} \sum_{t=1}^{p-j} (\tilde{E}_{(j),lt}^i)^2 + \frac{q}{n_k p_{-j}} \sum_{Y^i=k} \sum_{t=1}^{p-j} \|\mathbf{U}^i\|_2^2 \right\} \quad (\text{G.25})$$

$$\leq \|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max} \left\{ \max_{l,t} \frac{1}{n_k} \sum_{Y^i=k} (\tilde{E}_{(j),lt}^i)^2 + \frac{q}{n_k} \sum_{Y^i=k} \|\mathbf{U}^i\|_2^2 \right\} \quad (\text{G.26})$$

$$\leq 2(1 + q^2)\epsilon \quad (\text{G.27})$$

285

Second,

$$\left\| \frac{1}{n_k p-j} \sum_{Y^i=k} \tilde{\mathbf{E}}_{(j)}^i (\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k)_{(j)}^\top \right\|_{\max} \quad (\text{G.28})$$

$$= \max_{l,s} \left| \frac{1}{n_k p-j} \sum_{Y^i=k} \sum_{t=1}^{p-j} \tilde{E}_{(j),lt}^i (\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k)_{(j),st} \right| \quad (\text{G.29})$$

$$\leq \|\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k\|_{\max} \cdot \max_l \frac{1}{n_k p-j} \sum_{Y^i=k} \sum_{t=1}^{p-j} |\tilde{E}_{(j),lt}^i| \quad (\text{G.30})$$

$$\leq \|\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k\|_{\max} \cdot \max_{l,t} \frac{1}{n_k - 1} \sum_{Y^i=k} |\tilde{E}_{(j),lt}^i| \quad (\text{G.31})$$

$$\leq \|\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k\|_{\max} \cdot \max_{l,t} \frac{1}{n_k - 1} \sqrt{\sum_{Y^i=k} (\tilde{E}_{(j),lt}^i)^2 \sum_{Y^i=k} 1^2} \quad (\text{G.32})$$

$$\leq \|\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k\|_{\max} \max_{l,t} \sqrt{\frac{2}{n_k} \sum_{Y^i=k} (\tilde{E}_{(j),lt}^i)^2} \leq 2\epsilon \quad (\text{G.33})$$

286

Therefore,

$$\|L_6\|_{\max} \leq 2(2 + q^2)\epsilon \quad (\text{G.34})$$

287

For L_7 , we further decompose it to be

$$L_7 = \frac{1}{n_k p-j} \sum_{Y^i=k} ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i)_{(j)} ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i)_{(j)}^\top \quad (\text{G.35})$$

$$+ \frac{2}{n_k p-j} \sum_{Y^i=k} ((\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i)_{(j)} (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k)_{(j)}^\top \quad (\text{G.36})$$

$$+ \frac{1}{p-j} \sum_{Y^i=k} (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k)_{(j)} (\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k)_{(j)}^\top \quad (\text{G.37})$$

$$\equiv L_8 + L_9 + L_{10} \quad (\text{G.38})$$

$$\begin{aligned} \|L_8\|_{\max} &= \max_{l,s} \frac{1}{n_k p-j} \sum_{Y^i=k} \sum_{t=1}^{p-j} [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i]_{(j),lt} [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \mathbf{U}^i]_{(j),st} \\ &\leq \frac{\|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max}^2}{n_k p-j} \sum_{Y^i=k} \sum_{t=1}^{p-j} \|\mathbf{U}^i\|_1^2 \leq 2q^2 \epsilon^2 \end{aligned}$$

$$\|L_9\|_{\max} = \max_{l,s} \frac{4}{p_{-j}} \left| \sum_{t=1}^{p-j} [(\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}) \bar{\times}_{(M+1)} \hat{\boldsymbol{\phi}}_k]_{(j),lt} (\boldsymbol{\mu}_k - \hat{\boldsymbol{\mu}}_k)_{(j),st} \right| \quad (\text{G.39})$$

$$\leq \frac{4\|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max}}{p_{-j}} \sum_{t=1}^{p-j} \|\hat{\boldsymbol{\phi}}_k\|_1 \cdot |(\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k)_{(j),st}| \quad (\text{G.40})$$

$$\leq 4\|\boldsymbol{\alpha} - \hat{\boldsymbol{\alpha}}\|_{\max} \cdot (\|\boldsymbol{\phi}_k\|_1 + q\epsilon) \cdot \|\hat{\boldsymbol{\mu}}_k - \boldsymbol{\mu}_k\|_{\max} \quad (\text{G.41})$$

$$\leq 4\epsilon^2(\|\boldsymbol{\phi}_k\|_1 + q\epsilon) \quad (\text{G.42})$$

288 By Lemma F.3, we have that, with a probability greater than $1 - 2p^2 \exp(-Cn^2 p_{-j} \epsilon^2)$,

$$289 \|L_{10}\|_{\max} \leq \epsilon.$$

290 Putting all the results together, we have the desired conclusion.

291 □

292 **Lemma G.4.** *There exists $\epsilon_0 > 0$ such that for $\frac{4}{\pi_k n} < \epsilon < \epsilon_0$, we have*

$$\Pr(\|\hat{\boldsymbol{\Sigma}}_j - \boldsymbol{\Sigma}_j\|_{\max} > Cq^2\epsilon) \leq 2Kp_j^2 \exp(-Cnp_{-j}\epsilon^2) + CKq^2p \exp(-\frac{Cn\epsilon^2}{q^2}) + 2K \exp(-Cn\pi_{\min}^2) \quad (\text{G.43})$$

293 **Proof of Lemma G.4.** Because $\hat{\boldsymbol{\Sigma}}_j = \sum_{k=1}^K \frac{n_k - 1}{n - K} \boldsymbol{\Sigma}_j^{(k)}$, we have

$$\|\hat{\boldsymbol{\Sigma}}_j - \boldsymbol{\Sigma}_j\|_{\max} = \left\| \sum_{k=1}^K \frac{n_k}{n} (\hat{\boldsymbol{\Sigma}}_j^{(k)} - \boldsymbol{\Sigma}_j) \right\|_{\max} \leq \left(\sum_{k=1}^K \frac{n_k}{n} \right) \max_k \|\hat{\boldsymbol{\Sigma}}_j^{(k)} - \boldsymbol{\Sigma}_j\|_{\max} \leq q^2\epsilon$$

294 provided that $\max_k \|\hat{\boldsymbol{\Sigma}}_j^{(k)} - \boldsymbol{\Sigma}_j\|_{\max} \leq q^2\epsilon$. By Lemma G.3 we have the desired conclusion. □

295 **Lemma G.5.** *There exists $\epsilon_0 > 0$ such that for $\frac{4}{\pi_k n} < \epsilon < \epsilon_0$, we have*

$$\Pr(\|\hat{\boldsymbol{\Sigma}} - \boldsymbol{\Sigma}\|_{\max} > \epsilon) \leq 2K \sum_{m=1}^M p_m^2 \exp(-Cnp_{-m}\epsilon^2) + CMKq^2p \exp(-\frac{Cn\epsilon^2}{q^2}) + 2KM \exp(-Cn\pi_{\min}^2) \quad (\text{G.44})$$

Proof of Lemma G.5. Combine Lemma G.4 with Lemma F.4 and we have the desired conclusion.

□

To prove Theorem 2, we define the oracle estimator

$$(\hat{\beta}_{2,\mathcal{D}}^{(\text{oracle})}, \dots, \hat{\beta}_{K,\mathcal{D}}^{(\text{oracle})}) = \arg \min_{\beta_{2,\mathcal{D}}, \dots, \beta_{K,\mathcal{D}}} \sum_{k=2}^K \{ \beta_{k,\mathcal{D}}^\top \hat{\Sigma} \beta_{k,\mathcal{D}} - 2(\hat{\nu}_{k,\mathcal{D}} - \hat{\nu}_{1,\mathcal{D}})^\top \beta_{k,\mathcal{D}} + \lambda \sum_{j \in \mathcal{D}} \|\beta_{\cdot j}\| \} \quad (\text{G.45})$$

We also collect two technical propositions proved in Mai et al. (2017).

Proposition G.1. For a set of real numbers $\{a_1, \dots, a_N\}$, if $\sum_{i=1}^N a_i^2 \leq \kappa^2 < 1$, then $\sum_{i=1}^N (a_i + b)^2$ as long as $b < \frac{1 - \kappa}{\sqrt{N}}$.

Proposition G.2. If $\|\hat{\Sigma} - \Sigma\|_{\max} < \frac{\epsilon}{d}$, $\|(\hat{\nu}_k - \hat{\nu}_1) - (\nu_k - \nu_1)\|_\infty < \epsilon/d$ for $\epsilon < \min\{\frac{1}{2\varphi}, \frac{\lambda}{1 + \varphi\Delta}\}$,

we have that

$$\|\hat{\Sigma}_{\mathcal{D}^c, \mathcal{D}} \hat{\Sigma}_{\mathcal{D}\mathcal{D}}^{-1} - \Sigma_{\mathcal{D}^c, \mathcal{D}} \Sigma_{\mathcal{D}\mathcal{D}}^{-1}\|_\infty < \frac{\varphi\epsilon}{1 - \varphi\epsilon} \quad (\text{G.46})$$

$$\|\hat{\beta}_{k,\mathcal{D}}^0 - \beta_{k,\mathcal{D}}^0\|_1 < \frac{\varphi\epsilon(1 + \varphi\Delta)}{1 - \varphi\epsilon} \quad (\text{G.47})$$

$$\|\hat{\beta}_{k,\mathcal{D}}^{\text{oracle}} - \beta_{k,\mathcal{D}}\|_\infty \leq 4\lambda\varphi \quad (\text{G.48})$$

The following lemmas will be helpful in proving Theorem 2

Lemma G.6. For $(\hat{\beta}_{2,\mathcal{D}}^{(\text{oracle})}, \dots, \hat{\beta}_{K,\mathcal{D}}^{(\text{oracle})})$ defined in (G.45), if

$$\max_{j \in \mathcal{D}^c} \left[\sum_{k=2}^K \{ (\hat{\Sigma}_{\mathcal{D}^c, \mathcal{D}} \hat{\beta}_{k,\mathcal{D}}^{(\text{oracle})})_j - (\hat{\nu}_{kj} - \hat{\nu}_{1j}) \}^2 \right]^{1/2} < \lambda/2 \quad (\text{G.49})$$

then we have that $\hat{\mathcal{D}} = \mathcal{D}$ and $\text{vec}(\hat{\mathbf{B}}_k)_{\mathcal{D}} = \hat{\beta}_{k,\mathcal{D}}^{(\text{oracle})}$.

Proof of Lemma G.6. When (G.49) is true, $\hat{\beta}_k = (\hat{\beta}_{k,\mathcal{D}}^{(\text{oracle})}, 0)$, $k = 2, \dots, K$ satisfies the Karush-Kuhn-Tucker condition. Hence we have the desired conclusion.

□

Proof of Theorem 2. By Lemmas F.4, G.2, & G.5, for any $0 < \epsilon < \min\{\frac{1}{2\varphi}, \frac{\lambda}{1 + \varphi\Delta}\}$, with a probability greater than

$$1 - 2K \sum_{m=1} p_m^2 \exp(-C \frac{np-m\epsilon^2}{d^2}) - CKq^2p \exp\{-\frac{Cn\epsilon^2}{d^2q^2}\} - 2Kp \exp(-Cn\epsilon^2)$$

we have that,

$$\|\widehat{\Sigma} - \Sigma\|_{\max} \leq \frac{\epsilon}{d}, \|\widehat{\mu}_k - \mu_k\|_{\max} \leq \frac{\epsilon}{d}. \quad (\text{G.50})$$

By Lemma G.6 and (G.48), it suffices to show that, under the event in (G.50) with a properly chosen ϵ , (G.49) must be true. To this end, note that

$$\begin{aligned} 0 &= \beta_{\mathcal{D}^C} = (\Sigma^{-1}(\nu_k - \nu_1))_{\mathcal{D}^C} \\ &= -(\Sigma_{\mathcal{D}^C, \mathcal{D}^C} - \Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}^C})^{-1} (\Sigma_{\mathcal{D}^C} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} (\nu_{k, \mathcal{D}} - \nu_{1, \mathcal{D}}) - (\nu_{k, \mathcal{D}^C} - \nu_{1, \mathcal{D}^C})) \end{aligned}$$

which implies that

$$\Sigma_{\mathcal{D}^C} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} (\nu_{k, \mathcal{D}} - \nu_{1, \mathcal{D}}) = \nu_{k, \mathcal{D}^C} - \nu_{1, \mathcal{D}^C} \quad (\text{G.51})$$

Now we assume that (G.50) holds. For any $j \in \mathcal{D}^C$, we have that

$$|(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})})_j - (\widehat{\nu}_k - \widehat{\nu}_1)| \leq |\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})} - \Sigma_{\mathcal{D}^C, \mathcal{D}} \beta_{k, \mathcal{D}}| + |(\widehat{\nu}_{kj} - \widehat{\nu}_{1j}) - (\nu_{kj} - \nu_{1j})|$$

where we make use of (G.51). Under (G.50), $|(\widehat{\nu}_{kj} - \widehat{\nu}_{1j}) - (\nu_{kj} - \nu_{1j})| \leq \epsilon$. Meanwhile, it is easy to show that

$$\widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})} = \widehat{\beta}_{k, \mathcal{D}}^0 - \frac{\lambda}{2} \widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}} \widehat{\mathbf{t}}_{k, \mathcal{D}} \quad (\text{G.52})$$

where $\widehat{\beta}_{k, \mathcal{D}}^0 = \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} (\widehat{\nu}_{k, \mathcal{D}} - \widehat{\nu}_{1, \mathcal{D}})$ and $\widehat{\mathbf{t}}_{k, \mathcal{D}}$ is the sub-gradient of the group lasso penalty at $\widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})}$. Hence,

$$|\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})} - \Sigma_{\mathcal{D}^C, \mathcal{D}} \beta_{k, \mathcal{D}}| \leq |(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^0)_j - (\Sigma_{\mathcal{D}^C, \mathcal{D}} \beta_{k, \mathcal{D}})_j| + \frac{\lambda}{2} |(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} \widehat{\mathbf{t}}_{k, \mathcal{D}})_j|$$

For the first term,

$$\begin{aligned} &|(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^0)_j - (\Sigma_{\mathcal{D}^C, \mathcal{D}} \beta_{k, \mathcal{D}})_j| \\ &\leq \|\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} - \Sigma_{\mathcal{D}^C, \mathcal{D}}\|_{\infty} \|\widehat{\beta}_{k, \mathcal{D}}^0 - \beta_{k, \mathcal{D}}\|_{\infty} + \|\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} - \Sigma_{\mathcal{D}^C, \mathcal{D}}\|_{\infty} \|\beta_{k, \mathcal{D}}\|_{\infty} + \|\Sigma_{\mathcal{D}^C, \mathcal{D}} (\widehat{\beta}_{k, \mathcal{D}}^0 - \beta_{k, \mathcal{D}})\|_{\infty} \\ &\leq C\epsilon \end{aligned}$$

320 For the second term,

$$|(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} \widehat{\mathbf{t}}_{k, \mathcal{D}})_j| \leq |(\Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} \mathbf{t}_{k, \mathcal{D}})_j| + |(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} \widehat{\mathbf{t}}_{k, \mathcal{D}})_j - (\Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} \mathbf{t}_{k, \mathcal{D}})_j|$$

321 Now note that

$$\begin{aligned} & |(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} \widehat{\mathbf{t}}_{k, \mathcal{D}})_j - (\Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} \mathbf{t}_{k, \mathcal{D}})_j| \\ & \leq \|\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} - \Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1}\|_{\infty} \|\widehat{\mathbf{t}}_{k, \mathcal{D}} - \mathbf{t}_{k, \mathcal{D}}\|_{\infty} + \|\Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1}\|_{\infty} \|\widehat{\mathbf{t}}_{k, \mathcal{D}} - \mathbf{t}_{k, \mathcal{D}}\| \\ & \quad + \|\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} - \Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1}\|_{\infty} |(\mathbf{t}_{k, \mathcal{D}})_j| \end{aligned}$$

322 Because, by (G.48),

$$|\widehat{t}_{kj} - t_{kj}| = \left| \frac{\widehat{\beta}_{kj}^{\text{oracle}}}{\|\widehat{\beta}_{\cdot j}^{\text{oracle}}\|} - \frac{\beta_{kj}}{\|\beta_{\cdot j}\|} \right| \leq \frac{2\|\widehat{\beta}_{\cdot j}^{\text{oracle}} - \beta_{\cdot j}\|}{\|\widehat{\beta}_{\cdot j}^{\text{oracle}}\|} \leq C\lambda\varphi$$

323 we have

$$|(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\Sigma}_{\mathcal{D}, \mathcal{D}}^{-1} \widehat{\mathbf{t}}_{k, \mathcal{D}})_j| \leq C\varphi\lambda \quad (\text{G.53})$$

324 Combine the bounds for the two terms and we have

$$|(\widehat{\Sigma}_{\mathcal{D}^C, \mathcal{D}} \widehat{\beta}_{k, \mathcal{D}}^{(\text{oracle})})_j - (\widehat{\nu}_k - \widehat{\nu}_1)| \leq \frac{\lambda}{2} |(\Sigma_{\mathcal{D}^C, \mathcal{D}} \Sigma_{\mathcal{D}, \mathcal{D}}^{-1} \mathbf{t}_{k, \mathcal{D}})_j| + C\lambda^2 \quad (\text{G.54})$$

325 Let $\epsilon = \frac{1 + \varphi\Delta}{4\varphi}\lambda$. Condition (C1) and Proposition G.1 together imply that there exists a
 326 generic constant $M_0 > 0$, such that when $\lambda < M_0(1 - \kappa)$, (G.49) holds. Let $\epsilon = \frac{1 + \varphi\Delta}{4\varphi}\lambda$ and
 327 we have the first conclusion.

328 By the first conclusion, if $\{\frac{d^2 \log(pd)}{n}\}^{1/2} \ll \lambda \ll b_{\min}$, $\lambda \rightarrow 0$, we have that $\widehat{\mathcal{D}} = \mathcal{D}$ and
 329 $\|\text{vec}(\widehat{\mathbf{B}}_k) - \text{vec}(\mathbf{B}_k)\|_{\infty} \rightarrow 0$. $\square$

Proof of Theorem 3. By Lemmas F.4, G.2, & G.5, for any $0 < \epsilon < \min\{\frac{1}{2\varphi}, \frac{\lambda}{1 + \varphi\Delta}\}$, with
 a probability greater than

$$1 - 2K \sum_{m=1} p_m^2 \exp(-C \frac{np_{-m}\epsilon^2}{d^2}) - CKq^2p \exp\{-\frac{Cn\epsilon^2}{d^2q^2}\} - 2Kp \exp(-Cn\epsilon^2)$$

we have that the events in (G.50) have happened. Meanwhile, by the Bernstein inequality, with a probability greater than $1 - 2K \exp(-Cn\epsilon^2)$, we have that

$$|\widehat{\pi}_k - \pi_k| \leq \epsilon \quad (\text{G.55})$$

By the proofs of Theorem 2 in Mai et al. (2017), we have that (G.50) & (G.55) imply $|R_n - R| \leq C\lambda^{1/3}$. Let $\lambda = \frac{1 + \varphi\Delta}{4\varphi}\lambda$ and we have the first conclusion. The second conclusion is a direct consequence of the first conclusion. $\square$

References

- Bertsekas, D. P. (2016). *Nonlinear Programming*. Athena Scientific, 3 edition.
- Bickel, P. and Levina, E. (2004). Some theory for fisher’s linear discriminant function, ‘naive bayes’, and some alternatives when there are many more variables than observations. *Bernoulli*, 10:989–1010.
- Breiman, L. (2001). Random forests. *Machine learning*, 45(1):5–32.
- Clemmensen, L., Hastie, T., Witten, D., and Ersbøll, B. (2011). Sparse discriminant analysis. *Technometrics*, 53(4):406–413.
- Cortes, C. and Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3):273–297.
- Dimitriadou, E., Hornik, K., Leisch, F., Meyer, D., and Weingessel, A. (2009). E1071: Misc functions of the department of statistics (e1071), tu wien.
- Goeman, J., Meijer, R., and Chaturvedi, N. (2012). penalized: L1 (lasso and fused lasso) and l2 (ridge) penalized estimation in glms and in the cox model. *CRAN R package*.
- Lai, Z., Xu, Y., Yang, J., Tang, J., and Zhang, D. (2013). Sparse tensor discriminant analysis. *IEEE Transactions on Image Processing*, 22(10):3904–3915.
- Li, L. and Zhang, X. (2017). Parsimonious tensor response regression. *Journal of the American Statistical Association*, 112(519):1131–1146.
- Li, Q. and Schonfeld, D. (2014). Multilinear discriminant analysis for higher-order tensor data classification. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 36(12):2524–2537.
- Liaw, A. and Wiener, M. (2014). Package ‘randomforest’: Breiman and cutler’s random forests for classification and regression. 4:6–10.
- Mai, Q., Yang, Y., and Zou, H. (2017). Multiclass sparse discriminant analysis. *Statistica Sinica*, In press.

- 359 Tseng, P. (2001). Convergence of a block coordinate descent method for nondifferentiable
360 minimization. *Journal of Optimization Theory and Applications*, 109:475–494.
- 361 Wimalawarne, K., Tomioka, R., and Sugiyama, M. (2016). Theoretical and experimental anal-
362 yses of tensor-based regression and classification. *Neural Computation*, 28(4):686–715.
- 363 Witten, D. M. and Tibshirani, R. (2011). Penalized classification using fisher’s linear discrimi-
364 nant. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 73(5):753–
365 772.
- 366 Zhong, W. and Suslick, K. S. (2015). Matrix discriminant analysis with application to colori-
367 metric sensor array data. *Technometrics*, 57(4):524–534.
- 368 Zhou, H., Li, L., and Zhu, H. (2013). Tensor regression with applications in neuroimaging data
369 analysis. *Journal of the American Statistical Association*, 108(502):540–552.