

Supplementary Materials for “Common Reducing Subspace Model and Network Alternation Analysis”

by Wenjing Wang, Xin Zhang and Lexin Li

A1 Additional Simulations

A1.1 Computation time

We report the computation time of our proposed method in covariance estimation. In addition, we report the running time of two variants of our method, as well as the methods of Franks and Hoff (2016) and Cook and Forzani (2008). We use the simulation examples in Section 5.1. Table S1 reports the average running time, in seconds, over 20 data replications. It is seen that our method runs reasonably fast. It is comparable in terms of running time to the method of Cook and Forzani (2008), and is about 10 to 50 fold faster than that of Franks and Hoff (2016). All simulations have been carried out on a x86 64-redhat-linux-gnu (64-bit) computing platform, with Intel Xeon CPU E5-2670 2.60GHz processors.

Table S1: Running time for covariance estimation. The average running time, in seconds, over 20 data replications for the covariance estimation examples in Section 5.1. CRS: the proposed method with the one-direction-at-a-time (1D) initialization plus the full Grassmann manifold optimization; CRS(1D): the proposed method with the 1D initialization only; CRS(rdm+FG): the proposed method with a random initialization plus the full Grassmann manifold optimization; FH16: the shared subspace model of Franks and Hoff (2016); CF08: the covariance reducing model of Cook and Forzani (2008).

	$p = 100, n_k = 40$	$p = 100, n_k = 100$	$p = 200, n_k = 40$	$p = 200, n_k = 100$	max SE
	Model I				
CRS	38.5	37.1	274	273	1.85
CRS(1D)	32.3	31.2	237	236	1.62
CRS(rdm+FG)	6.02	5.90	36.7	36.8	0.25
FH16	550	559	551	592	3.00
CF08	8.40	42.1	38.8	34.7	13.8
	Model II				
CRS	8.26	8.70	83.2	83.2	1.90
CRS(1D)	8.16	8.59	82.8	82.6	1.90
CRS(rdm+FG)	6.28	5.63	39.9	38.3	0.81
FH16	578	538	571	549	4.18
CF08	3.77	3.28	36.0	36.3	4.52
	Model III				
CRS	7.97	8.61	81.9	86.13	3.42
CRS(1D)	7.90	8.52	81.3	85.7	3.41
CRS(rdm+FG)	4.14	4.24	31.2	27.6	1.78
FH16	594	590	596	531	2.04
CF08	13.2	14.2	143	180	12.9

A1.2 Effect of initialization

Our method has employed the one-direction-at-a-time (1D) approach of Cook and Zhang (2016) for initialization, followed by a full Grassmann manifold optimization. Our initialization step produces a $\sqrt{n}$ -consistent initial estimator, which is important for the subsequent manifold optimization. Based on our empirical experience, we have found that, after our initialization step, the manifold optimization usually converges after only a few iterations. To further evaluate the effect of the initialization, we consider two variants of our method. One is our method with the 1D initialization only, and the other is with a random initialization plus the full Grassmann manifold optimization. Table S1 reports the computation time, and Table S2 reports the corresponding estimation error. It is first observed in Table S1 that the computation time with and without the full Grassmann manifold optimization step after the 1D initialization is close. This suggests that the manifold optimization converges quickly after the 1D initialization. Moreover, it is seen in Table S2 that the method with a random initialization usually results in a much larger estimation error. This suggests the importance of a proper initialization.

A1.3 Number of groups

We further study the effect of the number of groups in the context of sparse precision matrix estimation. We adopt the simulation setup in Section 5.2 of the paper, and add the number of groups $K = 3$ and $K = 5$.

Specifically, when $K = 3$, our model is of the form,

$$\begin{aligned} \text{Model I: } \quad & \Phi_k = B_k + \delta I, \quad \text{where } B_k = \Gamma \Psi_k \Gamma^T, \delta = 0.5. \\ \text{Model II: } \quad & \Phi_k = B_k + \delta I, \quad \text{where } B_k = \sigma_k^2 \Gamma \Gamma^T, \sigma_1 = 1, \sigma_2 = 2, \sigma_3 = 3, \delta = 0.5. \\ \text{Model III: } \quad & \Phi_k = B_k + \delta I, \quad \text{where } B_k = \sigma_k \Gamma \Gamma^T, \sigma_1 = -2, \sigma_2 = -3, \sigma_3 = -4, \delta = 5. \end{aligned}$$

When $K = 5$, our model is of the form,

Table S2: Effect of initialization. The average estimation error over 100 data replications for the covariance estimation examples in Section 5.1. CRS: the proposed method with the one-direction-at-a-time (1D) initialization plus the full Grassmann manifold optimization; CRS(1D): the proposed method with the 1D initialization only; CRS(rdm+FG): the proposed method with a random initialization plus the full Grassmann manifold optimization.

	$p = 100, n_k = 40$	$p = 100, n_k = 100$	$p = 200, n_k = 40$	$p = 200, n_k = 100$	max SE
Model I					
CRS	1.469	0.913	2.388	1.387	0.061
CRS(1D)	1.721	1.049	3.231	1.989	0.070
CRS(rdm+FG)	6.503	6.511	6.802	6.726	0.059
Model II					
CRS	0.004	0.002	0.004	0.003	0.000
CRS(1D)	0.004	0.002	0.004	0.003	0.000
CRS(rdm+FG)	0.755	0.481	1.097	0.700	0.006
Model III					
CRS	0.413	0.263	0.471	0.293	0.006
CRS(1D)	0.415	0.260	0.469	0.295	0.006
CRS(rdm+FG)	3.289	3.320	3.593	3.352	0.100

Table S3: Sparse precision matrix estimation when $K = 3$. Reported are the average estimation error, the sensitivity (%) and specificity (%) of nonzero elements selection, over 100 data replications. CRS: the proposed common reducing subspace model; FH16: the shared subspace model of Franks and Hoff (2016); CF08: the covariance reducing model of Cook and Forzani (2008); DWW14: the fused graphical Lasso estimator of Danaher et al. (2014); LL15: the joint precision estimator of Lee and Liu (2015); Sample: the sample mean estimator. For CRS, FH16, CF08 and Sample estimators, we further combine with a fast sparse estimation method of Liu and Luo (2015) to obtain a sparse precision matrix estimator. The last column, max SE: the maximum standard error. The left panel reports the estimation error. The right panel reports the sensitivity (the first number) and specificity (the second number).

	Estimation error					Selection specificity and sensitivity				
	$p = 100$ $n_k = 40$	$p = 100$ $n_k = 100$	$p = 200$ $n_k = 40$	$p = 200$ $n_k = 100$	max SE	$p = 100$ $n_k = 40$	$p = 100$ $n_k = 100$	$p = 200$ $n_k = 40$	$p = 200$ $n_k = 100$	
Model I						Model I				
CRS	0.647	0.361	1.120	0.625	0.009	90.0, 72.7	94.0, 72.3	75.7, 77.1	87.2, 74.0	
FH16	2.943	2.941	2.955	2.955	0.000	4.76, 97.4	4.81, 97.4	2.48, 98.6	2.38, 98.8	
CF08	3.070	3.001	3.115	3.122	0.022	88.4, 64.1	94.3, 57.1	70.6, 74.7	85.3, 67.3	
DWW14	2.253	1.440	2.519	2.175	0.010	90.3, 49.1	93.8, 48.1	61.2, 76.5	80.8, 36.5	
LL15	2.246	2.166	3.273	2.658	0.001	80.1, 73.6	67.8, 97.6	69.3, 70.3	56.1, 97.5	
Sample	1.509	1.045	2.402	1.791	0.004	82.4, 66.8	90.9, 60.2	49.9, 83.8	77.0, 69.7	
Model II						Model II				
CRS	2.498	1.665	4.202	2.848	0.037	98.2, 52.1	98.9, 51.0	96.6, 57.8	98.0, 55.9	
FH16	35.74	35.74	35.76	35.76	0.001	6.18, 96.0	1.00, 91.2	2.95, 98.3	2.26, 99.0	
CF08	3.059	2.972	3.119	3.087	0.019	97.9, 38.1	99.2, 27.6	94.8, 51.9	97.8, 41.2	
DWW14	9.263	4.919	27.66	7.267	0.027	99.8, 3.05	99.7, 4.44	85.5, 15.3	99.7, 4.17	
LL15	11.48	11.96	33.20	33.44	0.004	92.9, 75.1	92.1, 97.5	66.4, 98.6	58.3, 100.0	
Sample	3.897	2.527	8.631	5.253	0.052	96.7, 42.1	98.1, 39.3	93.4, 50.4	94.6, 44.7	
Model III						Model III				
CRS	1.993	1.188	4.244	2.606	0.033	86.2, 74.3	91.5, 72.1	72.9, 79.2	82.7, 74.4	
FH16	1.040	0.644	1.310	0.919	0.015	97.7, 79.1	98.0, 84.3	95.8, 83.1	96.8, 85.9	
CF08	1.649	1.272	2.114	1.724	0.078	88.3, 70.4	93.0, 67.7	74.3, 76.2	84.8, 71.0	
DWW14	5.753	4.586	17.58	6.343	0.090	84.5, 52.6	83.4, 70.1	88.2, 40.4	73.8, 70.9	
LL15	10.57	6.107	10.10	9.265	0.002	64.4, 73.7	50.4, 97.4	50.5, 70.2	29.5, 97.3	
Sample	11.56	8.104	17.78	13.59	0.028	81.5, 67.9	89.0, 61.3	62.3, 79.4	76.3, 71.2	

Model I: $\Phi_k = B_k + \delta I$, where $B_k = \Gamma \Psi_k \Gamma^T$, $\delta = 0.5$.

Model II: $\Phi_k = B_k + \delta I$, where $B_k = \sigma_k^2 \Gamma \Gamma^T$,
 $\sigma_1 = 0.5, \sigma_2 = 1, \sigma_3 = 1.5, \sigma_4 = 2, \sigma_5 = 3, \delta = 0.5$.

Model III: $\Phi_k = B_k + \delta I$, where $B_k = \sigma_k \Gamma \Gamma^T$,
 $\sigma_1 = -0.5, \sigma_2 = -1, \sigma_3 = -2, \sigma_4 = -3, \sigma_5 = -4, \delta = 5$.

When $K > 2$, we define the estimation error as,

$$\sum_{k=2}^K \sum_{j=1}^{k-1} \|\Phi_k - \Phi_j - (\hat{\Phi}_k - \hat{\Phi}_j)\|_F$$

Let $\Psi_s = \Phi_k - \Phi_j$, $k = 2, \dots, K, j = 1, \dots, k-1, s = 1, \dots, m = \binom{K}{2}$, then we define the specificity and sensitivity as,

Table S4: Sparse precision matrix estimation when $K = 5$. The rest of setup is the same as in Table S3.

	Estimation error					Selection specificity and sensitivity				
	$p = 100$ $n_k = 40$	$p = 100$ $n_k = 100$	$p = 200$ $n_k = 40$	$p = 200$ $n_k = 100$	max SE	$p = 100$ $n_k = 40$	$p = 100$ $n_k = 100$	$p = 200$ $n_k = 40$	$p = 200$ $n_k = 100$	
	Model I					Model I				
CRS	2.172	1.440	3.286	2.094	0.165	93.7, 72.2	96.3, 71.0	85.6, 74.8	91.7, 73.5	
FH16	14.75	14.75	14.79	14.79	0.002	4.93, 96.9	5.40, 96.8	2.53, 98.6	2.50, 98.7	
CF08	10.82	9.394	11.69	10.81	0.234	93.0, 65.4	96.4, 62.5	84.4, 70.7	91.8, 66.2	
DWW14	11.29	11.01	12.91	14.83	0.148	78.0, 56.5	75.6, 78.8	35.9, 78.3	15.8, 93.8	
LL15	10.51	9.686	14.92	14.03	0.075	82.2, 56.6	75.9, 88.2	21.1, 96.9	45.2, 94.8	
Sample	5.958	4.032	10.35	7.062	0.519	84.6, 65.5	91.9, 59.2	56.0, 81.4	80.4, 67.0	
	Model II					Model II				
CRS	6.199	4.087	9.571	6.710	0.089	96.8, 58.2	98.1, 57.2	93.2, 61.7	97.4, 62.3	
FH16	91.51	91.51	91.59	91.60	0.002	6.91, 95.0	8.09, 93.4	3.07, 98.0	3.22, 97.8	
CF08	76.24	70.05	79.51	77.44	0.815	98.1, 39.8	99.0, 38.6	94.3, 50.0	97.4, 42.2	
DWW14	79.12	79.45	85.03	84.92	0.058	82.9, 54.6	78.8, 79.2	71.9, 61.8	69.9, 80.0	
LL15	51.45	39.71	85.53	85.41	0.223	83.2, 94.1	88.6, 95.5	53.2, 96.3	45.7, 99.9	
Sample	11.29	7.347	25.60	15.34	0.116	92.8, 54.6	95.6, 50.7	82.8, 63.7	91.4, 57.6	
	Model III					Model III				
CRS	9.104	5.334	18.21	11.26	0.055	83.4, 76.4	89.5, 74.5	67.9, 81.1	78.6, 78.3	
FH16	5.865	3.210	6.464	3.885	0.040	88.4, 89.1	85.2, 94.2	83.2, 93.3	65.7, 97.8	
CF08	17.67	13.01	25.10	18.92	0.575	76.7, 74.0	78.9, 73.5	62.6, 78.5	60.0, 81.4	
DWW14	31.35	30.82	37.14	37.66	0.159	41.2, 99.5	41.5, 99.6	19.0, 99.7	20.1, 99.8	
LL15	32.32	31.56	42.08	39.78	0.056	44.2, 97.0	42.9, 99.8	272, 97.1	26.0, 99.9	
Sample	32.29	23.05	48.27	36.60	0.058	66.6, 78.4	78.3, 72.5	45.5, 87.1	46.4, 83.6	

$$\text{specificity} = \frac{1}{m} \sum_{s=1}^m \frac{\#(\Psi_s \neq 0, \hat{\Psi}_s \neq 0)}{\#(\Psi_s \neq 0)}$$

$$\text{sensitivity} = \frac{1}{m} \sum_{s=1}^m \frac{\#(\Psi_s = 0, \hat{\Psi}_s = 0)}{\#(\Psi_s = 0)}$$

We report the simulation results over 100 data replications in Tables S3 and S4. It is seen that, for $K > 2$, the performance of our method remains competitive. Moreover, the estimation error of all the methods increases as K increases, since it is defined as the sum of K -choose-2 pairs of differences and a larger K implies more difference terms involved. On the other hand, the selection accuracy in terms of sensitivity and specificity remain about the same as K increases.

A2 Additional Theory

A2.1 Connection with envelope

We present an additional theorem that connects our common reducing subspace model with the notion of the envelope and the central subspace in the covariance reducing model of Cook and Forzani (2008).

Theorem S1. *The following statements are true.*

- (i) *The smallest common reducing subspace that satisfies (1) is $\mathcal{E}_\Delta(\Upsilon)$.*
- (ii) $\mathcal{E}_\Delta(\Upsilon) = \mathcal{E}_\Delta(\mathcal{C}) = \mathcal{E}_\Delta(\mathcal{P})$.

(iii) For all $k = 1, \dots, K$, $\mathcal{E}_\Delta(\mathcal{C}) = \mathcal{E}_{\Delta_k}(\mathcal{C})$ and $\mathcal{E}_\Delta(\mathcal{P}) = \mathcal{E}_{\Delta_k}(\mathcal{P})$.

Proof. (i) First, we show that, if $\text{span}(\Gamma) = \mathcal{E}_\Delta(\Upsilon)$, then it satisfies our common reducing subspace model (1); that is,

$$\Delta_k = \Gamma \Omega_k \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top, \quad k = 1, \dots, K,$$

for some symmetric positive definite matrices Ω_k and Ω_0 . We achieve this by verifying: (a) $P_\Gamma \Delta_k Q_\Gamma = 0$, and (b) $Q_\Gamma \Delta_k Q_\Gamma = Q_\Gamma \Delta Q_\Gamma$. Recall that $P_{\Upsilon(\Delta)} = \Upsilon(\Upsilon^\top \Delta \Upsilon)^{-1} \Upsilon^\top \Delta$ and

$$P_{\Upsilon(\Delta)} Q_\Gamma = \Upsilon(\Upsilon^\top \Delta \Upsilon)^{-1} \Upsilon^\top \Delta Q_\Gamma = \Upsilon(\Upsilon^\top \Delta \Upsilon)^{-1} \Upsilon^\top Q_\Gamma \Delta Q_\Gamma = 0,$$

where the equality $\Delta Q_\Gamma = Q_\Gamma \Delta Q_\Gamma$ is due to the fact that $\text{span}(\Gamma)$ reduces Δ , and the last equality is due to that $\Upsilon^\top Q_\Gamma = 0$ as $\text{span}(\Upsilon) \subseteq \text{span}(\Gamma)$. Then by (5), we have

$$P_\Gamma \Delta_k Q_\Gamma = P_\Gamma \{\Delta + P_{\Upsilon(\Delta)}^\top (\Delta_k - \Delta) P_{\Upsilon(\Delta)}\} Q_\Gamma = P_\Gamma \Delta Q_\Gamma = 0,$$

where the last equality again follows from the fact that $\text{span}(\Gamma)$ is a reducing subspace of Δ . Similarly, we have

$$Q_\Gamma \Delta_k Q_\Gamma = Q_\Gamma \{\Delta + P_{\Upsilon(\Delta)}^\top (\Delta_k - \Delta) P_{\Upsilon(\Delta)}\} Q_\Gamma = Q_\Gamma \Delta Q_\Gamma.$$

Second, we show that $\mathcal{E}_\Delta(\Upsilon)$ is the smallest subspace that satisfies model (1). We prove by contradiction. Suppose there is a subspace $\text{span}(\mathbf{G})$ that satisfies,

$$\Delta_k = \mathbf{G} \Phi_k \mathbf{G}^\top + \mathbf{G}_0 \Phi_0 \mathbf{G}_0^\top, \quad k = 1, \dots, K,$$

for some symmetric positive definite matrices Φ_k and Φ_0 . Then by the definition of $\Delta = \sum_k \pi_k \Delta_k$,

$$\Delta = \mathbf{G} \Phi \mathbf{G}^\top + \mathbf{G}_0 \Phi_0 \mathbf{G}_0^\top, \quad \Phi = \sum_k \pi_k \Phi_k,$$

which implies that $\text{span}(\mathbf{G})$ is a reducing subspace of Δ . Furthermore, we can now write $\text{span}(\Delta_k - \Delta) = \text{span}(\mathbf{G} \Phi_k \mathbf{G}^\top - \mathbf{G} \Phi \mathbf{G}^\top) \subseteq \text{span}(\mathbf{G})$, and $P_\mathbf{G} = P_{\mathbf{G}(\Delta)}$. It is thus straightforward to verify that $\text{span}(\mathbf{G})$ satisfies the sufficient dimension reduction condition (5). Henceforth, $\text{span}(\Upsilon) \subseteq \text{span}(\mathbf{G})$. Together, we have shown that $\text{span}(\mathbf{G})$ is a reducing subspace of Δ that contains $\text{span}(\Upsilon)$. By the definition of the envelope, $\mathcal{E}_\Delta(\Upsilon)$ is the smallest reducing subspace of Δ that contains $\text{span}(\Upsilon)$. This proves the first statement (i).

(ii) First, we show that $\mathcal{E}_\Delta(\mathcal{C}) \subseteq \mathcal{E}_\Delta(\Upsilon)$. Let $\text{span}(\Gamma) = \mathcal{E}_\Delta(\Upsilon)$. Then by Theorem S1 (i), we have $\Delta_k = \Gamma \Omega_k \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top$. Henceforth, $\mathcal{C} \subseteq \text{span}(\Gamma)$, and $\mathcal{E}_\Delta(\mathcal{C}) \subseteq \text{span}(\Gamma)$, because $\text{span}(\Gamma)$ is a reducing subspace of Δ .

Second, we show that $\mathcal{E}_\Delta(\Upsilon) \subseteq \mathcal{E}_\Delta(\mathcal{C})$. Let $\text{span}(\Gamma) = \mathcal{E}_\Delta(\mathcal{C})$. Then we can write $\Delta = \Gamma \Omega \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top$ for some Ω and Ω_0 , because $\text{span}(\Gamma)$ is a reducing subspace of Δ . Since $\mathcal{C} \subseteq \text{span}(\Gamma)$, we have that $\text{span}(\Delta_k - \Delta) \subseteq \mathcal{C} \subseteq \text{span}(\Gamma)$. Therefore, we can write $\Delta_k = \Delta + \Gamma \Phi_k \Gamma^\top$ for some symmetric matrix Φ_k . It follows that $\Delta_k = \Gamma(\Omega + \Phi_k) \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top$, and henceforth $\text{span}(\Gamma)$ is a common reducing subspace. The conclusion follows from the fact that $\mathcal{E}_\Delta(\Upsilon)$ is the smallest common reducing subspace.

Combining the above two steps, we have $\mathcal{E}_\Delta(\Upsilon) = \mathcal{E}_\Delta(\mathcal{C})$. The proof for $\mathcal{E}_\Delta(\Upsilon) = \mathcal{E}_\Delta(\mathcal{P})$ follows a similar argument, and is omitted. This proves the second statement (ii).

(iii) We show that, for all $k = 1, \dots, K$, $\mathcal{E}_\Delta(\mathcal{C}) = \mathcal{E}_{\Delta_k}(\mathcal{C})$. It follows from a similar argument that $\mathcal{E}_\Delta(\mathcal{P}) = \mathcal{E}_{\Delta_k}(\mathcal{P})$.

We prove the equality $\mathcal{E}_\Delta(\mathcal{C}) = \mathcal{E}_{\Delta_k}(\mathcal{C})$, by first assuming $\text{span}(\mathcal{C}) \subseteq \text{span}(\Gamma)$, and showing that: (a) if in addition $\text{span}(\Gamma)$ is a reducing subspace of Δ , then it is also a reducing subspace of Δ_k , and (b) if

$\text{span}(\mathbf{\Gamma})$ is a reducing subspace of $\mathbf{\Delta}$, then it is also a reducing subspace of $\mathbf{\Delta}_k$. Therefore, (a) implies that $\mathcal{E}_{\mathbf{\Delta}}(\mathcal{C}) \subseteq \mathcal{E}_{\mathbf{\Delta}_k}(\mathcal{C})$, and (b) implies that $\mathcal{E}_{\mathbf{\Delta}_k}(\mathcal{C}) \subseteq \mathcal{E}_{\mathbf{\Delta}}(\mathcal{C})$.

If $\text{span}(\mathcal{C}) \subseteq \text{span}(\mathbf{\Gamma})$, then $\mathbf{\Delta}_k - \mathbf{\Delta} = \mathbf{P}_{\mathbf{\Gamma}}^T(\mathbf{\Delta}_k - \mathbf{\Delta})\mathbf{P}_{\mathbf{\Gamma}}$. If $\text{span}(\mathbf{\Gamma})$ is a reducing subspace of $\mathbf{\Delta}$, then we can write

$$\mathbf{\Delta}_k = \mathbf{\Delta} + \mathbf{P}_{\mathbf{\Gamma}}(\mathbf{\Delta}_k - \mathbf{\Delta})\mathbf{P}_{\mathbf{\Gamma}} = \mathbf{P}_{\mathbf{\Gamma}}\mathbf{\Delta}\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\mathbf{\Delta}\mathbf{Q}_{\mathbf{\Gamma}} + \mathbf{P}_{\mathbf{\Gamma}}(\mathbf{\Delta}_k - \mathbf{\Delta})\mathbf{P}_{\mathbf{\Gamma}} = \mathbf{P}_{\mathbf{\Gamma}}\mathbf{\Delta}_k\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\mathbf{\Delta}\mathbf{Q}_{\mathbf{\Gamma}},$$

which implies that $\text{span}(\mathbf{\Gamma})$ is also a reducing subspace of $\mathbf{\Delta}_k$. Similarly, by symmetry, we can show if $\text{span}(\mathbf{\Gamma})$ is a reducing subspace of $\mathbf{\Delta}_k$ then it is also a reducing subspace of $\mathbf{\Delta}$. This proves the third statement (iii), and completes the proof of Theorem S1. $\square$

A2.2 Theory with a fixed sample size

We also study the theoretical properties of our estimator when the sample size $n = \sum_{k=1}^K n_k$ is fixed. For simplicity, we assume that the common reducing subspace $\text{span}(\mathbf{\Gamma})$ is known. We note that it is very difficult to fully characterize the estimation variation in the estimated basis matrix $\hat{\mathbf{\Gamma}}$, and hence difficult to provide a full theoretical analysis under the fixed sample size. Nevertheless, we find that, even the result under the reduced setting with a known $\text{span}(\mathbf{\Gamma})$ offers some useful understanding of our method.

Specifically, we define $\mathbf{I}_{(p,q)}$ as the vec-permutation matrix such that $\text{vec}(\mathbf{A}_{p \times q}) = \mathbf{I}_{(p,q)}\text{vec}(\mathbf{A}_{p \times q}^T)$ for any $p \times q$ matrix $\mathbf{A}_{p \times q}$ (Henderson and Searle, 1979). An important property of $\mathbf{I}_{(p,p)}$ is that $\mathbf{I}_{(p,p)}(\mathbf{B}_{p \times p} \otimes \mathbf{A}_{p \times p}) = (\mathbf{A}_{p \times p} \otimes \mathbf{B}_{p \times p})\mathbf{I}_{(p,p)}$. We have the following result for a fixed sample size n .

Proposition S1. Assume that the observed data matrices $\mathbf{D}_{i,k}$, $i = 1, \dots, n_k$, $k = 1, \dots, K$, are i.i.d. samples from a Wishart distribution $W_p(\mathbf{\Delta}_k/d, d)$ with d degrees of freedom, and that the basis matrix $\mathbf{\Gamma} \in \mathbb{R}^{p \times u}$ is known. Then the variance reduction of the proposed estimator over the sample covariance estimator is $0 \leq \text{cov}\{\text{vec}(\overline{\mathbf{D}}_k)\} - \text{cov}\{\text{vec}(\hat{\mathbf{\Delta}}_k)\} = \mathbf{V}_1 + \mathbf{V}_2 + \mathbf{V}_3$, where

$$\begin{aligned} \mathbf{V}_1 &= \frac{n - n_k}{nn_k d} \{(\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T \otimes (\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T))\} \{\mathbf{I}_{p^2} + \mathbf{I}_{(p,p)}\}, \\ \mathbf{V}_2 &= \frac{1}{n_k d} \{(\mathbf{\Gamma} \mathbf{\Omega}_k \mathbf{\Gamma}^T) \otimes (\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T)\} \{\mathbf{I}_{p^2} + \mathbf{I}_{(p,p)}\}, \\ \mathbf{V}_3 &= \frac{1}{n_k d} \{(\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T) \otimes (\mathbf{\Gamma} \mathbf{\Omega}_k \mathbf{\Gamma}^T)\} \{\mathbf{I}_{p^2} + \mathbf{I}_{(p,p)}\}. \end{aligned}$$

This proposition offers some additional insight on the efficiency gain of our estimator, while it does not require n_k to diverge to infinity. Specifically, our proposed estimator based on the common reducing subspace model is of the form, $\hat{\mathbf{\Delta}}_k = \mathbf{P}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}\mathbf{Q}_{\mathbf{\Gamma}}$. The sample covariance estimator $\overline{\mathbf{D}}_k$ can be written as $\overline{\mathbf{D}}_k = \mathbf{P}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{Q}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{P}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{Q}_{\mathbf{\Gamma}}$. Then the efficiency gain of $\hat{\mathbf{\Delta}}_k$ over $\overline{\mathbf{D}}_k$, under the fixed n_k , includes two parts. The first part, quantified as $\mathbf{V}_1$, comes from pooling the shared information across groups. That is, we estimate $\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T$ by $\mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}\mathbf{Q}_{\mathbf{\Gamma}}$, instead of $\mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{Q}_{\mathbf{\Gamma}}$. Interestingly, if we have $n_1 = \dots = n_K = n/K$, then the factors $(n - n_k)(nn_k d) = (K - 1)(nd)$ increases as the number of groups increases. In other words, we gain more efficiency with more groups. The second part of efficiency gain, quantified as $\mathbf{V}_2 + \mathbf{V}_3$, comes from the reducing subspace structure. Under our model assumption, the proposed estimator identifies and eliminates $\mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{P}_{\mathbf{\Gamma}}$ and $\mathbf{P}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{Q}_{\mathbf{\Gamma}}$ in the estimation. These two terms are essentially zero matrices, and their estimation would only bring extraneous variability. Our method, however, avoids estimating these terms.

Proof. Under a fix n_k and the true $\mathbf{\Gamma}$, we have $\hat{\mathbf{\Delta}}_k = \mathbf{P}_{\mathbf{\Gamma}}\overline{\mathbf{D}}_k\mathbf{P}_{\mathbf{\Gamma}} + \mathbf{Q}_{\mathbf{\Gamma}}\overline{\mathbf{D}}\mathbf{Q}_{\mathbf{\Gamma}}$, and

$$\text{vec}(\hat{\mathbf{\Delta}}_k) = (\mathbf{P}_{\mathbf{\Gamma}} \otimes \mathbf{P}_{\mathbf{\Gamma}})\text{vec}(\overline{\mathbf{D}}_k) + (\mathbf{Q}_{\mathbf{\Gamma}} \otimes \mathbf{Q}_{\mathbf{\Gamma}})\text{vec}(\overline{\mathbf{D}}).$$

Since $\mathbf{D}_{i,k} \sim W_p(\Delta_k/d, d)$ are i.i.d. Wishart samples, $i = 1, \dots, n_k$, $k = 1, \dots, K$, we have $n_k \bar{\mathbf{D}}_k = \sum_{i=1}^{n_k} \mathbf{D}_{i,k} \sim W_p(\Delta_k/d, n_k d)$. Denote $\Sigma_k = \text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\}$, and $\Sigma = \text{cov}\{\text{vec}(\bar{\mathbf{D}})\}$. Then $\Sigma = \sum_{k=1}^K n_k^2/n^2 \Sigma_k$, and $\text{cov}\{\text{vec}(\bar{\mathbf{D}}), \text{vec}(\bar{\mathbf{D}}_k)\} = n_k/n \Sigma_k$. Furthermore, based on Henderson and Searle (1979), we have

$$\Sigma_k = \frac{1}{n_k d} (\Delta_k \otimes \Delta_k) \{I_{p^2} + I_{(p,p)}\},$$

where I_{p^2} is an identity matrix with dimension $p^2 \times p^2$, and $I_{(p,p)}$ is the vec-permutation matrix. Then,

$$\begin{aligned} \text{cov}\{\text{vec}(\hat{\Delta}_k)\} &= (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) + (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \\ &\quad + \frac{n_k}{n} (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) + \frac{n_k}{n} (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \\ &= (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) + (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma), \\ \text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\} &= \{(\mathbf{P}_\Gamma + \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma + \mathbf{Q}_\Gamma)\} \Sigma_k \{(\mathbf{P}_\Gamma + \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma + \mathbf{Q}_\Gamma)\}. \end{aligned}$$

Directly expanding $\text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\}$ leads to 16 additive terms, of which 10 are zero; i.e.,

$$\begin{aligned} (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) &= (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) = 0, \\ (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) &= (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) = 0, \\ (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) &= (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) = 0, \\ (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) &= (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) = 0, \\ (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) &= (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) = 0, \end{aligned}$$

since $\mathbf{P}_\Gamma \Delta_k \mathbf{Q}_\Gamma = \mathbf{Q}_\Gamma \Delta_k \mathbf{P}_\Gamma = 0$. There are six nonzero terms left in $\text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\}$, i.e., $(\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma)$, $(\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma)$, and

$$\begin{aligned} (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) &= \frac{1}{n_k d} (\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma), \\ (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) &= \frac{1}{n_k d} \{(\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma)\} I_{(p,p)}, \\ (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{Q}_\Gamma) &= \frac{1}{n_k d} \{(\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma)\} I_{(p,p)}, \\ (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{P}_\Gamma) &= \frac{1}{n_k d} \{(\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma)\}. \end{aligned}$$

Therefore, we have

$$\begin{aligned} \text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\} &= (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) \Sigma_k (\mathbf{P}_\Gamma \otimes \mathbf{P}_\Gamma) + (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \\ &\quad + \frac{1}{n_k d} \{(\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma)\} \{I_{p^2} + I_{(p,p)}\} \\ &\quad + \frac{1}{n_k d} \{(\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma)\} \{I_{p^2} + I_{(p,p)}\}. \end{aligned}$$

The difference of the variance of the two estimators is,

$$\begin{aligned} \text{cov}\{\text{vec}(\bar{\mathbf{D}}_k)\} - \text{cov}\{\text{vec}(\hat{\Delta}_k)\} &= (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) (\Sigma_k - \Sigma) (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \\ &\quad + \frac{1}{n_k d} \{(\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma)\} \{I_{p^2} + I_{(p,p)}\} \\ &\quad + \frac{1}{n_k d} \{(\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma)\} \{I_{p^2} + I_{(p,p)}\}. \end{aligned}$$

In addition, we note that,

$$\begin{aligned} (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma_k (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) &= \frac{1}{n_k d} \{ (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \} \{ \mathbf{I}_{p^2} + \mathbf{I}_{(p,p)} \}, \\ (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) \Sigma (\mathbf{Q}_\Gamma \otimes \mathbf{Q}_\Gamma) &= \frac{1}{nd} \{ (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \otimes (\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma) \} \{ \mathbf{I}_{p^2} + \mathbf{I}_{(p,p)} \}. \end{aligned}$$

The conclusion then follows by noting that $\mathbf{Q}_\Gamma \Delta_k \mathbf{Q}_\Gamma = \Gamma_0 \Omega_0 \Gamma_0^\top$, and $\mathbf{P}_\Gamma \Delta_k \mathbf{P}_\Gamma = \Gamma \Omega_k \Gamma^\top$. This completes the proof of Proposition S1. $\square$

A3 Proofs

A3.1 Proof of Proposition 1

Proof. The equivalence between (1) and (3), and between (2) and (4), directly follow by noting that $\Delta_k \Delta_k^{-1} = \mathbf{I}_p$. Next we prove the equivalence between (1) and (2).

First we show (1) implies (2). It is sufficient to show that $\mathbf{P}_S \Delta_k \mathbf{Q}_S = 0$. Since $\mathcal{S} = \text{span}(\Gamma)$, we have $\mathbf{P}_S \Delta_k \mathbf{Q}_S = \mathbf{P}_S (\Gamma \Omega_k \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top) \mathbf{Q}_S = 0$, and $\mathbf{Q}_S (\Delta_k - \Delta_j) \mathbf{Q}_S = \mathbf{Q}_S \{ \Gamma (\Omega_k - \Omega_j) \Gamma^\top \} \mathbf{Q}_S = 0$.

Next we show (2) implies (1). Since $\mathcal{S} = \text{span}(\Gamma)$, we rewrite $\mathbf{P}_S = \Gamma \Gamma^\top$, and $\mathbf{Q}_S = \Gamma_0 \Gamma_0^\top$. Then we have $\Delta_k = \Gamma \Gamma^\top \Delta_k \Gamma \Gamma^\top + \Gamma_0 \Gamma_0^\top \Delta_k \Gamma_0 \Gamma_0^\top$. Denote $\Omega_k = \Gamma^\top \Delta_k \Gamma$, and $\Omega_{k0} = \Gamma_0^\top \Delta_k \Gamma_0$. We have $\Gamma_0 (\Omega_{k0} - \Omega_{j0}) \Gamma_0^\top = \mathbf{Q}_S (\Delta_k - \Delta_j) \mathbf{Q}_S = 0$, which holds for all $1 \leq k, j \leq K$. Therefore, for any k, j , $\Omega_{k0} = \Omega_{j0} = \Omega_0$. Then $\Delta_k = \Gamma \Omega_k \Gamma + \Gamma_0 \Omega_0 \Gamma_0^\top$ follows. This completes the proof of Proposition 1. $\square$

A3.2 Proof of Proposition 2

Proof. Proposition 2 is a direct consequence of Theorem S1. $\square$

A3.3 Proof of Proposition 3

Proof. First we describe the joint likelihood-based objective function. For i.i.d. Wishart samples $\mathbf{D}_{i,k} \sim W_p(\Delta_k/d, d)$, $i = 1, \dots, n_k$, $k = 1, \dots, K$, we have $n_k \bar{\mathbf{D}}_k = \sum_{i=1}^{n_k} \mathbf{D}_{i,k} \sim W_p(\Delta_k/d, n_k d)$. Then

$$f(n_1 \bar{\mathbf{D}}_1, \dots, n_K \bar{\mathbf{D}}_K | \Lambda_1, \dots, \Lambda_K, \beta, \Delta_0, d) \propto \prod_{k=1}^K \left| \frac{\Delta_k}{d} \right|^{-\frac{n_k d}{2}} \text{etr} \left\{ - \left(\frac{\Delta_k}{d} \right)^{-1} n_k \bar{\mathbf{D}}_k / 2 \right\}.$$

Therefore we have the negative log-likelihood as the objective function, which is of the form,

$$\ell(\Delta_1, \dots, \Delta_K) = \sum_{k=1}^K \frac{n_k d}{2} \{ \log |\Delta_k| + \text{tr}(\Delta_k^{-1} \bar{\mathbf{D}}_k) \}.$$

For our model, we have $\Delta_k = \Gamma \Omega_k \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top$, and $\Delta_k^{-1} = \Gamma \Omega_k^{-1} \Gamma^\top + \Gamma_0 \Omega_0^{-1} \Gamma_0^\top$, $k = 1, \dots, K$. Then the joint log-likelihood-based objective function is

$$\begin{aligned} \ell(\Gamma, \Omega_0, \Omega_1, \dots, \Omega_K) &= \sum_{k=1}^K \frac{n_k d}{2} \left[\log |\Gamma \Omega_k \Gamma^\top + \Gamma_0 \Omega_0 \Gamma_0^\top| + \text{tr} \{ (\Gamma \Omega_k^{-1} \Gamma^\top + \Gamma_0 \Omega_0^{-1} \Gamma_0^\top) \bar{\mathbf{D}}_k \} \right] \\ &= \sum_{k=1}^K \frac{n_k d}{2} \left[\log |\Omega_k| + \log |\Omega_0| + \text{tr} \{ (\Gamma \Omega_k^{-1} \Gamma^\top + \Gamma_0 \Omega_0^{-1} \Gamma_0^\top) \bar{\mathbf{D}}_k \} \right]. \end{aligned}$$

Furthermore, we have,

$$\begin{aligned}
\frac{\partial \ell}{\partial \Omega_0} &= \sum_{k=1}^K \frac{n_k d}{2} \left\{ \frac{\partial \log |\Omega_0|}{\partial \Omega_0} + \frac{\text{tr}(\Gamma_0 \Omega_0^{-1} \Gamma_0^T \bar{D}_k)}{\partial \Omega_0} \right\}, \\
\frac{\partial \ell}{\partial \Omega_k} &= \sum_{k=1}^K \frac{n_k d}{2} \left\{ \frac{\partial \log |\Omega_k|}{\partial \Omega_k} + \frac{\text{tr}(\Gamma \Omega_k^{-1} \Gamma^T \bar{D}_k)}{\partial \Omega_k} \right\}, \\
\frac{\partial \log |\Omega_0|}{\partial \Omega_0} &= \Omega_0^{-1}, \quad \frac{\partial \log |\Omega_k|}{\partial \Omega_k} = \Omega_k^{-1}, \\
\frac{\partial \text{tr}(\Gamma_0 \Omega_0^{-1} \Gamma_0^T \bar{D}_k)}{\partial \Omega_0} &= \frac{\partial \text{tr}(\Omega_0^{-1} \Gamma_0^T \bar{D}_k \Gamma_0)}{\partial \Omega_0} = -\Omega_0^{-1} \Gamma_0^T \bar{D}_k \Gamma_0 \Omega_0^{-1}, \\
\frac{\partial \text{tr}(\Gamma \Omega_k^{-1} \Gamma^T \bar{D}_k)}{\partial \Omega_k} &= \frac{\partial \text{tr}(\Omega_k^{-1} \Gamma^T \bar{D}_k \Gamma)}{\partial \Omega_k} = -\Omega_k^{-1} \Gamma^T \bar{D}_k \Gamma \Omega_k^{-1}.
\end{aligned}$$

If we set $\partial \ell / \partial \Omega_0 = 0$ and $\partial \ell / \partial \Omega_k = 0$, then the maximum likelihood estimators of Ω_0 , Ω_k and Δ_k satisfy that $\hat{\Omega}_0 = \hat{\Gamma}_0^T \bar{D} \hat{\Gamma}_0$, $\hat{\Omega}_k = \hat{\Gamma}^T \bar{D}_k \hat{\Gamma}$, and $\hat{\Delta}_k = \hat{\Gamma} \hat{\Gamma}^T \bar{D}_k \hat{\Gamma} \hat{\Gamma}^T + \hat{\Gamma}_0 \hat{\Gamma}_0^T \bar{D} \hat{\Gamma}_0 \hat{\Gamma}_0^T$, where $\bar{D} = \sum_{k=1}^K (n_k/n) \bar{D}_k$, and $\hat{\Gamma}$ is the minimizer of the following partially optimized negative log-likelihood,

$$\ell(\Gamma) = \sum_{k=1}^K \frac{n_k d}{2} (\log |\Gamma^T \bar{D}_k \Gamma| + \log |\Gamma_0^T \bar{D} \Gamma_0|),$$

in which (Γ, Γ_0) forms an orthogonal basis matrix of $\mathbb{R}^p$. To further simplify $\ell(\Gamma)$, we apply the identity $|\Gamma_0^T \bar{D} \Gamma_0| = |\bar{D}| |\Gamma^T(\bar{D})^{-1} \Gamma|$ from Proposition A2 of Cook and Forzani (2008). Omitting an additive constant $\log |\bar{D}|$, $\ell(\Gamma)$ can be rewritten as,

$$\ell(\Gamma) = \frac{nd}{2} \log |\Gamma^T(\bar{D})^{-1} \Gamma| + \sum_{k=1}^K \frac{n_k d}{2} \log |\Gamma^T \bar{D}_k \Gamma|.$$

Therefore the likelihood-based estimator $\hat{\Gamma}$ is obtained from minimizing the above objective function $\ell(\Gamma)$ subject to the constraint $\Gamma^T \Gamma = \mathbf{I}_u$. This completes the proof of Proposition 3. $\square$

A3.4 Proof of Theorem 1

Proof. By direct calculation, we can obtain both the asymptotic covariance, $\mathbf{V} = \mathbf{J}_h^{-1}$, of the standard maximum likelihood estimator, as well as the asymptotic covariance, $\mathbf{U} = \mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T$, of our common reducing subspace-based estimator. Specifically, the Fisher information matrix $\mathbf{J}_h$ and the gradient matrix $\mathbf{H}$ is consistent with Cook et al. (2010).

To show $\mathbf{U} \leq \mathbf{V}$, we observe that

$$\mathbf{V} - \mathbf{U} = \mathbf{J}_h^{-1} - \mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T = \mathbf{J}_h^{-\frac{1}{2}} \left\{ \mathbf{I} - \mathbf{J}_h^{\frac{1}{2}} \mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T \mathbf{J}_h^{\frac{1}{2}} \right\} \mathbf{J}_h^{-\frac{1}{2}} = \mathbf{J}_h^{-\frac{1}{2}} \mathbf{Q}_{\mathbf{J}_h^{\frac{1}{2}} \mathbf{H}} \mathbf{J}_h^{-\frac{1}{2}},$$

which is a symmetric positive semi-definite matrix. This completes the proof of Theorem 1. $\square$

A3.5 Proof of Theorem 2

Proof. We first derive an alternative form for $\mathbf{U} = \mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T$, which facilitates the interpretation. Since $\mathbf{U}$ only depends on the column space of $\mathbf{H}$, we can replace $\mathbf{H}$ by $\mathbf{H}_1$ that has the same column space as $\mathbf{H}$, whereas $\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1$ is block-diagonal, and

$$\begin{aligned}
\mathbf{H}_1 &= \begin{pmatrix} 2\mathbf{C}_p(\Gamma \Omega_1 \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) & \mathbf{C}_p(\Gamma \otimes \Gamma) \mathbf{E}_u & \mathbf{0} & \mathbf{C}_p(\Gamma_0 \otimes \Gamma_0) \mathbf{E}_{p-u} \\ 2\mathbf{C}_p(\Gamma \Omega_2 \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) & \mathbf{0} & \mathbf{C}_p(\Gamma \otimes \Gamma) \mathbf{E}_u & \mathbf{C}_p(\Gamma_0 \otimes \Gamma_0) \mathbf{E}_{p-u} \end{pmatrix} \\
&= (\mathbf{H}_{11} \quad \mathbf{H}_{12} \quad \mathbf{H}_{13} \quad \mathbf{H}_{14}).
\end{aligned}$$

Moreover, let

$$\mathbf{H}_2 = \begin{pmatrix} \mathbf{I}_u \otimes \mathbf{\Gamma}_0^T & 0 & 0 & 0 \\ 2\mathbf{C}_u(\mathbf{\Omega}_1 \otimes \mathbf{\Gamma}^T) & \mathbf{I}_{\frac{u(u+1)}{2}} & 0 & 0 \\ 2\mathbf{C}_u(\mathbf{\Omega}_2 \otimes \mathbf{\Gamma}^T) & 0 & \mathbf{I}_{\frac{u(u+1)}{2}} & 0 \\ 0 & 0 & 0 & \mathbf{I}_{\frac{(r-u)(r-u+1)}{2}} \end{pmatrix}.$$

Since $\mathbf{H} = \mathbf{H}_1\mathbf{H}_2$, and $\mathbf{H}_2$ is of the full row rank, we have $\text{span}(\mathbf{H}) = \text{span}(\mathbf{H}_1)$. Furthermore, by straightforward multiplication, we see that $\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1$ is block-diagonal with

$$\begin{aligned} \mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11} &= d \sum_{k=1}^2 \pi_k (\mathbf{\Omega}_k \otimes \mathbf{\Omega}_0^{-1} - 2\mathbf{I}_{u \times (r-u)} + \mathbf{\Omega}_k^{-1} \otimes \mathbf{\Omega}_0), \\ \mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12} &= \frac{d\pi_1}{2} \{ \mathbf{E}_u^T (\mathbf{\Omega}_1^{-1} \otimes \mathbf{\Omega}_1^{-1}) \mathbf{E}_u \}, \\ \mathbf{H}_{13}^T \mathbf{J}_h \mathbf{H}_{13} &= \frac{d\pi_2}{2} \{ \mathbf{E}_u^T (\mathbf{\Omega}_2^{-1} \otimes \mathbf{\Omega}_2^{-1}) \mathbf{E}_u \}, \\ \mathbf{H}_{14}^T \mathbf{J}_h \mathbf{H}_{14} &= \frac{d}{2} \{ \mathbf{E}_{p-u}^T (\mathbf{\Omega}_0^{-1} \otimes \mathbf{\Omega}_0^{-1}) \mathbf{E}_{p-u} \}. \end{aligned}$$

Thus, we can now write $\mathbf{U} = \mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^\dagger \mathbf{H}_1^T = \sum_{j=1}^4 \mathbf{H}_{1j}(\mathbf{H}_{1j}^T \mathbf{J}_h \mathbf{H}_{1j})^\dagger \mathbf{H}_{1j}^T$.

Our interest is $\text{avar}\{\sqrt{n}\text{vech}(\hat{\Delta}_1)\}$, which is the upper left block of $\mathbf{U}$. Since $\mathbf{H}_{13}^T \mathbf{J}_h \mathbf{H}_{13}$ does not contribute to the upper left block, we focus only on the upper left block of $\mathbf{H}_{1j}(\mathbf{H}_{1j}^T \mathbf{J}_h \mathbf{H}_{1j})^\dagger \mathbf{H}_{1j}^T$, for $j = 1, 2, 4$. Let $\mathbf{H}_{1,ij}$ denote the (i, j) th block of the matrix $\mathbf{H}_1$, then the upper left block of $\mathbf{U}$ can be decomposed in the following way:

$$\begin{aligned} \mathbf{H}_{1,11}(\mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11})^\dagger \mathbf{H}_{1,11}^T &= \frac{4}{d} \mathbf{C}_p(\mathbf{\Gamma} \mathbf{\Omega}_1 \otimes \mathbf{\Gamma}_0 - \mathbf{\Gamma} \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0) \\ &\times \left\{ \sum_{k=1}^2 \pi_k (\mathbf{\Omega}_k \otimes \mathbf{\Omega}_0^{-1} - 2\mathbf{I}_{u \times (r-u)} + \mathbf{\Omega}_k^{-1} \otimes \mathbf{\Omega}_0) \right\}^\dagger \\ &\times (\mathbf{\Omega}_1 \mathbf{\Gamma}^T \otimes \mathbf{\Omega}_0^T - \mathbf{\Gamma}^T \otimes \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T) \mathbf{C}_p^T, \\ \mathbf{H}_{1,12}(\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12})^\dagger \mathbf{H}_{1,12}^T &= \mathbf{C}_p(\mathbf{\Gamma} \otimes \mathbf{\Gamma}) \mathbf{E}_u \left\{ \frac{d\pi_1}{2} \mathbf{E}_u^T (\mathbf{\Omega}_1^{-1} \otimes \mathbf{\Omega}_1^{-1}) \mathbf{E}_u \right\}^\dagger \mathbf{E}_u^T (\mathbf{\Gamma}^T \otimes \mathbf{\Gamma}^T) \mathbf{C}_p^T \\ &= \frac{2}{d\pi_1} \mathbf{C}_p(\mathbf{\Gamma} \mathbf{\Omega}_1 \mathbf{\Gamma}^T \otimes \mathbf{\Gamma} \mathbf{\Omega}_1 \mathbf{\Gamma}^T) \mathbf{C}_p^T, \\ \mathbf{H}_{1,14}(\mathbf{H}_{14}^T \mathbf{J}_h \mathbf{H}_{14})^\dagger \mathbf{H}_{1,14}^T &= \mathbf{C}_p(\mathbf{\Gamma}_0 \otimes \mathbf{\Gamma}_0) \mathbf{E}_{p-u} \left\{ \frac{d}{2} \mathbf{E}_{p-u}^T (\mathbf{\Omega}_0^{-1} \otimes \mathbf{\Omega}_0^{-1}) \mathbf{E}_{p-u} \right\}^\dagger \mathbf{E}_{p-u}^T (\mathbf{\Gamma}_0^T \otimes \mathbf{\Gamma}_0^T) \mathbf{C}_p^T \\ &= \frac{2}{d} \mathbf{C}_p(\mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T) \mathbf{C}_p^T. \end{aligned}$$

Henceforth,

$$\text{avar}\{\sqrt{n}\text{vech}(\hat{\Delta}_1)\} = \mathbf{H}_{1,11}(\mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11})^\dagger \mathbf{H}_{1,11}^T + \mathbf{H}_{1,12}(\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12})^\dagger \mathbf{H}_{1,12}^T + \mathbf{H}_{1,14}(\mathbf{H}_{14}^T \mathbf{J}_h \mathbf{H}_{14})^\dagger \mathbf{H}_{1,14}^T.$$

The Fisher information matrix of $\boldsymbol{\theta}$ is simply $\mathbf{H}^T \mathbf{J}_h \mathbf{H}$; that is,

$$\mathbf{J}_\theta = \mathbf{H}^T \mathbf{J}_h \mathbf{H} = \begin{pmatrix} \mathbf{J}_\Gamma & \mathbf{J}_{\Gamma \mathbf{\Omega}_1} & \mathbf{J}_{\Gamma \mathbf{\Omega}_2} & 0 \\ \mathbf{J}_{\mathbf{\Omega}_1 \Gamma} & \mathbf{J}_{\mathbf{\Omega}_1} & 0 & 0 \\ \mathbf{J}_{\mathbf{\Omega}_2 \Gamma} & 0 & \mathbf{J}_{\mathbf{\Omega}_2} & 0 \\ 0 & 0 & 0 & \mathbf{J}_{\mathbf{\Omega}_0} \end{pmatrix}.$$

Each nonzero block is of the form,

$$\begin{aligned}
J_{\Gamma} &= 2d \sum_{k=1}^2 \pi_k (\Omega_k^T \Gamma^T \otimes I_p - \Gamma^T \otimes \Gamma_0 \Omega_0 \Gamma_0^T) C_p^T E_p^T (\Delta_k^{-1} \otimes \Delta_k^{-1}) E_p C_p (\Gamma \Omega_k \otimes I_p - \Gamma \otimes \Gamma_0 \Omega_0 \Gamma_0^T), \\
J_{\Gamma \Omega_1} &= J_{\Omega_1 \Gamma}^T = d\pi_1 (I_u \otimes \Delta_1^{-1} \Gamma) E_u, \\
J_{\Gamma \Omega_2} &= J_{\Omega_2 \Gamma}^T = d\pi_2 (I_u \otimes \Delta_2^{-1} \Gamma) E_u, \\
J_{\Omega_1} &= \frac{d\pi_1}{2} \{E_u^T (\Omega_1^{-1} \otimes \Omega_1^{-1}) E_u\}, \\
J_{\Omega_2} &= \frac{d\pi_2}{2} \{E_u^T (\Omega_2^{-1} \otimes \Omega_2^{-1}) E_u\}, \\
J_{\Omega_0} &= \frac{d}{2} \{E_{p-u}^T (\Omega_0^{-1} \otimes \Omega_0^{-1}) E_{p-u}\}.
\end{aligned}$$

Considering $\text{avar}\{\sqrt{n}\text{vec}(\hat{\Gamma})\}$, the asymptotic variance is the upper left block of J_{θ}^{-1} :

$$\left\{ J_{\Gamma} - \begin{pmatrix} J_{\Gamma \Omega_1} & J_{\Gamma \Omega_2} \end{pmatrix} \begin{pmatrix} J_{\Omega_1}^{-1} & 0 \\ 0 & J_{\Omega_2}^{-1} \end{pmatrix} \begin{pmatrix} J_{\Omega_1 \Gamma} \\ J_{\Omega_2 \Gamma} \end{pmatrix} \right\}^{-1}.$$

We have

$$\begin{aligned}
J_{\Gamma} &= 2d \sum_{k=1}^2 \pi_k (\Omega_k^T \Gamma^T \otimes I_p - \Gamma^T \otimes \Gamma_0 \Omega_0 \Gamma_0^T) C_p^T E_p^T (\Delta_k^{-1} \otimes \Delta_k^{-1}) E_p C_p (\Gamma \Omega_k \otimes I_p - \Gamma \otimes \Gamma_0 \Omega_0 \Gamma_0^T) \\
&= d \sum_{k=1}^2 \pi_k \{(\Omega_k \otimes \Delta_k^{-1}) + (\Omega_k \Gamma^T \Delta_k^{-1} \otimes \Delta_k^{-1} \Gamma \Omega_k) K_{ru} - 2(I_u \otimes \Gamma_0 \Gamma_0^T) + (\Omega_k^{-1} \otimes \Gamma_0 \Omega_0 \Gamma_0^T)\}.
\end{aligned}$$

Through calculation, we have

$$\begin{pmatrix} J_{\Gamma \Omega_1} & J_{\Gamma \Omega_2} \end{pmatrix} \begin{pmatrix} J_{\Omega_1}^{-1} & 0 \\ 0 & J_{\Omega_2}^{-1} \end{pmatrix} \begin{pmatrix} J_{\Omega_1 \Gamma} \\ J_{\Omega_2 \Gamma} \end{pmatrix} = d \sum_{k=1}^2 \pi_k \{ \Omega_k \otimes \Gamma \Omega_k^{-1} \Gamma^T + (\Omega_k \Gamma^T \Delta_k^{-1} \otimes \Delta_k^{-1} \Gamma \Omega_k) K_{ru} \},$$

where $K_{pu} \in \mathbb{R}^{pu \times pu}$ is the commutation matrix, such that for $M \in \mathbb{R}^{p \times u}$, $\text{vec}(M^T) = K_{pu} \text{vec}(M)$. Then we have

$$\left[\text{avar} \left\{ \sqrt{n} \text{vec}(\hat{\Gamma}) \right\} \right]^{-1} = d \sum_{k=1}^2 \pi_k \{ (\Omega_k \otimes \Gamma_0 \Omega_0^{-1} \Gamma_0^T) - 2(I_u \otimes \Gamma_0 \Gamma_0^T) + (\Omega_k^{-1} \otimes \Gamma_0 \Omega_0 \Gamma_0^T) \}.$$

In addition, we have $\text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Omega}_{1,\Gamma}) \right\} = J_{\Omega_1}^{-1}$, $\text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Omega}_{0,\Gamma}) \right\} = J_{\Omega_0}^{-1}$. By comparing these with $\text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Delta}_1) \right\}$, we see that

$$\begin{aligned}
H_{1,12} (H_{12}^T J_h H_{12})^\dagger H_{1,12}^T &= C_p (\Gamma \otimes \Gamma) E_u \cdot \text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Omega}_{1,\Gamma}) \right\} \cdot E_u^T (\Gamma^T \otimes \Gamma^T) C_p^T \\
&= \text{avar} \left[\sqrt{n} \{ C_p (\Gamma \otimes \Gamma) E_u \} \text{vech}(\hat{\Omega}_{1,\Gamma}) \right] \\
&= \text{avar} \left\{ \sqrt{n} \text{vech}(\Gamma \hat{\Omega}_{1,\Gamma} \Gamma^T) \right\}.
\end{aligned}$$

Similarly, we can show that $H_{1,14} (H_{14}^T J_h H_{14})^\dagger H_{1,14}^T = \text{avar} \left\{ \sqrt{n} \text{vech}(\Gamma_0 \hat{\Omega}_{0,\Gamma} \Gamma_0^T) \right\}$. Therefore we have $H_{1,12} (H_{12}^T J_h H_{12})^\dagger H_{1,12}^T + H_{1,14} (H_{14}^T J_h H_{14})^\dagger H_{1,14}^T = \text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Delta}_{1,\Gamma}) \right\}$. Also, comparing $H_{11}^T J_h H_{11}$ with $\text{avar} \left\{ \sqrt{n} \text{vec}(\hat{\Gamma}) \right\}$, we see that

$$H_{11}^T J_h H_{11} = (I_u \otimes \Gamma_0^T) \text{avar} \left\{ \sqrt{n} \text{vec}(\hat{\Gamma}) \right\}^{-1} (I_u \otimes \Gamma_0).$$

Therefore,

$$\begin{aligned} \mathbf{H}_{1,11}(\mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11})^\dagger \mathbf{H}_{1,11}^T &= 4\mathbf{C}_p(\mathbf{\Gamma}\mathbf{\Omega}_1 \otimes \mathbf{I}_p - \mathbf{\Gamma} \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T)(\mathbf{I}_u \otimes \mathbf{\Gamma}_0) \\ &\times \left[(\mathbf{I}_u \otimes \mathbf{\Gamma}_0^T) \text{avar} \left\{ \sqrt{n} \text{vec}(\widehat{\mathbf{\Gamma}}) \right\}^{-1} (\mathbf{I}_u \otimes \mathbf{\Gamma}_0) \right]^\dagger \\ &\times (\mathbf{I}_u \otimes \mathbf{\Gamma}_0^T)(\mathbf{\Omega}_1 \mathbf{\Gamma}^T \otimes \mathbf{I}_p - \mathbf{\Gamma}^T \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T). \end{aligned}$$

From Cook et al. (2010), since $\mathbf{P}_G = \mathbf{I}_u \otimes \mathbf{\Gamma}_0 \mathbf{\Gamma}_0^T$, and $\left[\text{avar} \left\{ \sqrt{n} \text{vec}(\widehat{\mathbf{\Gamma}}) \right\} \right]^{-1}$ commutes, we have

$$\begin{aligned} \mathbf{H}_{1,11}(\mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11})^\dagger \mathbf{H}_{1,11}^T &= 4\mathbf{C}_p(\mathbf{\Gamma}\mathbf{\Omega}_1 \otimes \mathbf{I}_p - \mathbf{\Gamma} \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T)(\mathbf{I}_u \otimes \mathbf{Q}_\Gamma) \times \text{avar} \left\{ \sqrt{n} \text{vec}(\widehat{\mathbf{\Gamma}}) \right\} \\ &\times (\mathbf{I}_u \otimes \mathbf{Q}_\Gamma)(\mathbf{\Omega}_1 \mathbf{\Gamma}^T \otimes \mathbf{I}_p - \mathbf{\Gamma}^T \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T) \mathbf{C}_p^T \\ &= \mathbf{A} \cdot \text{avar} \left\{ \sqrt{n} \text{vec}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}}) \right\} \cdot \mathbf{A}^T. \end{aligned}$$

We see that

$$\begin{aligned} \mathbf{A} &= 2\mathbf{C}_p(\mathbf{\Gamma}\mathbf{\Omega}_1 \otimes \mathbf{I}_p - \mathbf{\Gamma} \otimes \mathbf{\Gamma}_0 \mathbf{\Omega}_0 \mathbf{\Gamma}_0^T) = \frac{\partial \text{vech}(\widehat{\mathbf{\Delta}}_1)}{\partial \text{vec}^T(\mathbf{\Gamma})} \\ &= \left[\frac{\partial \text{vech} \left\{ \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0}(\widehat{\mathbf{\Gamma}}) \right\}}{\partial \text{vec}^T(\widehat{\mathbf{\Gamma}})} \right]_{\widehat{\mathbf{\Gamma}}=\mathbf{\Gamma}} = \left[\frac{\partial \text{vech} \left\{ \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}}) \right\}}{\partial \text{vec}^T(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}})} \right]_{\widehat{\mathbf{\Gamma}}=\mathbf{\Gamma}}, \end{aligned}$$

where

$$\widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}}) = \mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}} \mathbf{\Omega}_1 \widehat{\mathbf{\Gamma}}^T \mathbf{Q}_\Gamma + \mathbf{P}_\Gamma \widehat{\mathbf{\Gamma}}_0 \mathbf{\Omega}_0 \widehat{\mathbf{\Gamma}}_0^T \mathbf{P}_\Gamma = \mathbf{Q}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{Q}_\Gamma + \mathbf{P}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{P}_\Gamma.$$

By the delta method, we have that

$$\mathbf{A} \cdot \text{avar} \left\{ \sqrt{n} \text{vec}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}}) \right\} \cdot \mathbf{A}^T = \text{avar} \left\{ \sqrt{n} \text{vech}(\widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Gamma}})) \right\}.$$

Henceforth,

$$\begin{aligned} \mathbf{H}_{1,11}(\mathbf{H}_{11}^T \mathbf{J}_h \mathbf{H}_{11})^\dagger \mathbf{H}_{1,11}^T &= \text{avar} \left\{ \sqrt{n} \text{vech}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{Q}_\Gamma + \mathbf{P}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{P}_\Gamma) \right\} \\ &= \text{avar} \left\{ \sqrt{n} \text{vech}(\mathbf{Q}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{Q}_\Gamma) \right\} + \text{avar} \left\{ \sqrt{n} \text{vech}(\mathbf{P}_\Gamma \widehat{\mathbf{\Delta}}_{1,\Omega_1,\Omega_0} \mathbf{P}_\Gamma) \right\}. \end{aligned}$$

Then the result in (8) of Theorem 2 follows.

Next we consider

$$\widehat{\mathbf{g}}_{\text{crs}} = \text{vech}(\widehat{\mathbf{\Delta}}_1) - \text{vech}(\widehat{\mathbf{\Delta}}_2) = \mathbf{G} \begin{pmatrix} \text{vech} \widehat{\mathbf{\Delta}}_1 \\ \text{vech} \widehat{\mathbf{\Delta}}_2 \end{pmatrix},$$

where $\mathbf{G} = (\mathbf{I}_{p(p+1)/2} - \mathbf{I}_{p(p+1)/2})$. Then, we have $\text{avar}(\sqrt{n} \widehat{\mathbf{g}}_{\text{crs}}) = \mathbf{G} \mathbf{U} \mathbf{G}^T$, and

$$\mathbf{U} = \mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^\dagger \mathbf{H}_1^T = \begin{pmatrix} \mathbf{U}_1 & \mathbf{U}_2 \\ \mathbf{U}_2^T & \mathbf{U}_3 \end{pmatrix},$$

following the proof of Theorem 1. Since $\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1$ is block-diagonal, we can see that $(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^\dagger$ is a block-diagonal matrix with blocks $\mathbf{A}_j^\dagger = (\mathbf{H}_{1j}^T \mathbf{J}_h \mathbf{H}_{1j})^\dagger$, $j = 1, 2, 3, 4$. Then by straightforward calculation, we have

$$\begin{aligned} \mathbf{U}_1 &= \mathbf{H}_{1,11} \mathbf{A}_1^\dagger \mathbf{H}_{1,11}^T + \mathbf{H}_{1,12} \mathbf{A}_2^\dagger \mathbf{H}_{1,12}^T + \mathbf{H}_{1,14} \mathbf{A}_4^\dagger \mathbf{H}_{1,14}^T, \\ \mathbf{U}_2 &= \mathbf{H}_{1,11} \mathbf{A}_1^\dagger \mathbf{H}_{1,21}^T + \mathbf{H}_{1,14} \mathbf{A}_4^\dagger \mathbf{H}_{1,14}^T, \\ \mathbf{U}_3 &= \mathbf{H}_{1,21} \mathbf{A}_1^\dagger \mathbf{H}_{1,21}^T + \mathbf{H}_{1,23} \mathbf{A}_3^\dagger \mathbf{H}_{1,23}^T + \mathbf{H}_{1,24} \mathbf{A}_4^\dagger \mathbf{H}_{1,24}^T. \end{aligned}$$

$$\begin{aligned}
GUG^T &= H_{1,11}A_1^\dagger H_{1,11}^T + H_{1,12}A_2^\dagger H_{1,12}^T - H_{1,21}A_1^\dagger H_{1,11}^T - H_{1,11}A_1^\dagger H_{1,21}^T \\
&+ H_{1,21}A_1^\dagger H_{1,21}^T + H_{1,23}A_3^\dagger H_{1,23}^T.
\end{aligned}$$

Furthermore,

$$\begin{aligned}
H_{1,11}A_1^\dagger H_{1,11}^T &= 4C_p(\Gamma\Omega_1 \otimes \Gamma_0 - \Gamma \otimes \Gamma_0\Omega_0)(H_{11}^T J_h H_{11})^\dagger (\Omega_1\Gamma^T \otimes \Gamma_0^T - \Gamma^T \otimes \Omega_0\Gamma_0^T)C_p^T, \\
H_{1,12}A_2^\dagger H_{1,12}^T &= C_p(\Gamma \otimes \Gamma)E_u \left\{ \frac{d\pi_1}{2}(E_u^T(\Omega_1^{-1} \otimes \Omega_1^{-1})E_u) \right\}^\dagger E_u^T(\Gamma^T \otimes \Gamma^T)C_p^T, \\
H_{1,23}A_3^\dagger H_{1,23}^T &= C_p(\Gamma \otimes \Gamma)E_u \left\{ \frac{d\pi_2}{2}(E_u^T(\Omega_2^{-1} \otimes \Omega_2^{-1})E_u) \right\}^\dagger E_u^T(\Gamma^T \otimes \Gamma^T)C_p^T, \\
H_{1,21}A_1^\dagger H_{1,11}^T &= 4C_p(\Gamma\Omega_2 \otimes \Gamma_0 - \Gamma \otimes \Gamma_0\Omega_0)[H_{11}^T J_h H_{11}]^\dagger (\Omega_1\Gamma^T \otimes \Gamma_0^T - \Gamma^T \otimes \Omega_0\Gamma_0^T)C_p^T, \\
H_{1,21}A_1^\dagger H_{1,21}^T &= 4C_p(\Gamma\Omega_2 \otimes \Gamma_0 - \Gamma \otimes \Gamma_0\Omega_0)[H_{11}^T J_h H_{11}]^\dagger (\Omega_2\Gamma^T \otimes \Gamma_0^T - \Gamma^T \otimes \Omega_0\Gamma_0^T)C_p^T.
\end{aligned}$$

As shown earlier, we have

$$\begin{aligned}
H_{1,12}A_2^\dagger H_{1,12}^T + H_{1,23}A_3^\dagger H_{1,23}^T &= \text{avar} \left\{ \sqrt{n} \text{vech}(\Gamma\hat{\Omega}_{1,\Gamma}\Gamma^T) \right\} + \text{avar} \left\{ \sqrt{n} \text{vech}(\Gamma\hat{\Omega}_{2,\Gamma}\Gamma^T) \right\} \\
&= \text{avar}(\sqrt{n} \text{vech} \left\{ (\hat{\Delta}_{1,\Gamma} - \hat{\Delta}_{2,\Gamma}) \right\}).
\end{aligned}$$

In addition,

$$\begin{aligned}
&H_{1,11}A_1^\dagger H_{1,11}^T - H_{1,11}A_1^\dagger H_{1,21}^T + H_{1,21}A_1^\dagger H_{1,11}^T - H_{1,21}A_1^\dagger H_{1,21}^T \\
&= 4C_p \{ \Gamma(\Omega_1 - \Omega_2) \otimes I_p \} (I_u \otimes Q_\Gamma) \times \text{avar} \left\{ \sqrt{n} \text{vec}(\hat{\Gamma}) \right\} \times (I_u \otimes Q_\Gamma) \{ (\Omega_1 - \Omega_2)\Gamma^T \otimes I_p \} \\
&= B \cdot \text{avar} \left\{ \sqrt{n} \text{vec}(Q_\Gamma \hat{\Gamma}) \right\} \cdot B^T,
\end{aligned}$$

where $B = 2C_p [\Gamma(\Omega_1 - \Omega_2) \otimes I_p]$. Define $\Delta_{\text{diff}} = \Delta_1 - \Delta_2$, we have

$$B = \frac{\partial \text{vech}(\Delta_{\text{diff}})}{\partial \text{vec}^T(\Gamma)} = \left[\frac{\partial \text{vech} \left\{ \hat{\Delta}_{\text{diff},\Omega_1,\Omega_2}(\hat{\Gamma}) \right\}}{\partial \text{vec}^T(\hat{\Gamma})} \right]_{\hat{\Gamma}=\Gamma} = \left[\frac{\partial \text{vech} \left\{ \hat{\Delta}_{\text{diff},\Omega_1,\Omega_2}(Q_\Gamma \hat{\Gamma}) \right\}}{\partial \text{vec}^T(Q_\Gamma \hat{\Gamma})} \right]_{\hat{\Gamma}=\Gamma},$$

where $\hat{\Delta}_{\text{diff},\Omega_1,\Omega_2}(Q_\Gamma \hat{\Gamma}) = Q_\Gamma \hat{\Gamma}(\Omega_1 - \Omega_2)\hat{\Gamma}^T Q_\Gamma = Q_\Gamma(\hat{\Delta}_{1,\Omega_1,\Omega_0} - \hat{\Delta}_{2,\Omega_2,\Omega_0})Q_\Gamma$. Again by the delta method, we have that

$$\begin{aligned}
B \cdot \text{avar}(\sqrt{n} \text{vec}(Q_\Gamma \hat{\Gamma})) \cdot B^T &= \text{avar} \left\{ \sqrt{n} \text{vech}(\hat{\Delta}_{\text{diff},\Omega_1,\Omega_2}(Q_\Gamma \hat{\Gamma})) \right\} \\
&= \text{avar} \left\{ \sqrt{n} \text{vech}(Q_\Gamma(\hat{\Delta}_{1,\Omega_1,\Omega_0} - \hat{\Delta}_{2,\Omega_2,\Omega_0})Q_\Gamma) \right\}
\end{aligned}$$

Henceforth, the result in (9) of Theorem 2 follows. This completes the proof of Theorem 2. $\square$

A3.6 Proof of Theorem 3

Proof. Recall that

$$\ell(\Delta_1, \Delta_2) = - \sum_{k=1}^2 \frac{n_k d}{2} \{ \log |\Delta_k| + \text{tr}(\Delta_k^{-1} \overline{D}_k) \}.$$

We define a minimum discrepancy function,

$$\begin{aligned}
\mathbf{F}(\mathbf{h}, \hat{\mathbf{h}}) &= -\frac{2}{n} [\ell(\boldsymbol{\Delta}_1, \boldsymbol{\Delta}_2) - \ell(\overline{\mathbf{D}}_1, \overline{\mathbf{D}}_2)] \\
&= \sum_{k=1}^2 \pi_k d \left\{ \log |\boldsymbol{\Delta}_k| + \text{tr}(\boldsymbol{\Delta}_k^{-1} \overline{\mathbf{D}}_k) - \log |\overline{\mathbf{D}}_k| - \text{tr}(\overline{\mathbf{D}}_k^{-1} \overline{\mathbf{D}}_k) \right\},
\end{aligned}$$

with $\mathbf{h} = \{\boldsymbol{\Delta}_1(\boldsymbol{\theta}), \boldsymbol{\Delta}_2(\boldsymbol{\theta})\}$, and $\hat{\mathbf{h}} = (\overline{\mathbf{D}}_1, \overline{\mathbf{D}}_2)$.

Then the Fisher information is calculated as

$$\mathbf{J}_h = \left\{ \frac{1}{2} \frac{\partial \mathbf{F}(\mathbf{h}, \hat{\mathbf{h}})}{\partial \mathbf{h} \partial \mathbf{h}^T} \right\}_{\hat{\mathbf{h}}=\mathbf{h}=\mathbf{h}_{\text{true}}}.$$

Also, we have that: (a) $\mathbf{F}(\mathbf{h}, \hat{\mathbf{h}}) \geq 0$ for all $\mathbf{h}$ and $\hat{\mathbf{h}}$, (b) $\mathbf{F}(\mathbf{h}, \hat{\mathbf{h}}) = 0$ if and only if $\hat{\mathbf{h}} = \mathbf{h}$, and (c) $\mathbf{F}(\mathbf{h}, \hat{\mathbf{h}})$ is twice continuously differentiable in $\mathbf{h}$ and $\hat{\mathbf{h}}$. Since $\overline{\mathbf{D}}_k$ is the sample covariance matrix, by the Central Limit Theorem, we have $\sqrt{n} \text{vech}(\hat{\mathbf{h}}) \rightarrow \mathbf{W}$, for some positive definite covariance matrix $\mathbf{W}$. Then by Proposition 4.1 of Shapiro (1986), we have the consistency result of $\hat{\mathbf{h}}_{\text{crs}}$.

Moreover, we have

$$\begin{aligned}
\mathbf{W} - \mathbf{K} &= \mathbf{W} - \mathbf{J}_h^{-\frac{1}{2}} \mathbf{P}_{\mathbf{J}_h^{\frac{1}{2}} \mathbf{H}} \mathbf{J}_h^{\frac{1}{2}} \mathbf{W} \mathbf{J}_h^{\frac{1}{2}} \mathbf{P}_{\mathbf{J}_h^{\frac{1}{2}} \mathbf{H}} \mathbf{J}_h^{-\frac{1}{2}} \\
&= \mathbf{J}_h^{-\frac{1}{2}} \left[\mathbf{J}_h^{\frac{1}{2}} \mathbf{W} \mathbf{J}_h^{\frac{1}{2}} - \mathbf{P}_{\mathbf{J}_h^{\frac{1}{2}} \mathbf{H}} \mathbf{J}_h^{\frac{1}{2}} \mathbf{W} \mathbf{J}_h^{\frac{1}{2}} \mathbf{P}_{\mathbf{J}_h^{\frac{1}{2}} \mathbf{H}} \right] \mathbf{J}_h^{-\frac{1}{2}}.
\end{aligned}$$

This completes the proof of Theorem 3. □

A4 Computer Code

The computer code associated with this paper is available at the Biometrics website on Wiley Online Library.

References

- Cook, R. D. and Forzani, L. (2008). Covariance reducing models: An alternative to spectral modelling of covariance matrices. *Biometrika*, 95(4):799–812.
- Cook, R. D., Li, B., and Chiaromonte, F. (2010). Envelope models for parsimonious and efficient multivariate linear regression. *Statist. Sinica*, 20(3):927–960.
- Cook, R. D. and Zhang, X. (2016). Algorithms for envelope estimation. *Journal of Computational and Graphical Statistics*, 25(1):284–300.
- Danaher, P., Wang, P., and Witten, D. M. (2014). The joint graphical lasso for inverse covariance estimation across multiple classes. *Journal of the Royal Statistical Society: Series B.*, 76(2):373–397.
- Franks, A. and Hoff, P. (2016). Shared subspace models for multi-group covariance estimation. *arXiv preprint arXiv:1607.03045*.
- Henderson, H. V. and Searle, S. (1979). Vec and vech operators for matrices, with some uses in jacobians and multivariate statistics. *Canadian Journal of Statistics*, 7(1):65–81.

- Lee, W. and Liu, Y. (2015). Joint estimation of multiple precision matrices with common structures. *The Journal of Machine Learning Research*, 16(1):1035–1062.
- Liu, W. and Luo, X. (2015). Fast and adaptive sparse precision matrix estimation in high dimensions. *Journal of Multivariate Analysis*, 135:153–162.
- Shapiro, A. (1986). Asymptotic theory of overparameterized structural models. *Journal of the American Statistical Association*, 81(393):142–149.