

Supplementary Materials for “Efficient integration of sufficient dimension reduction and prediction in discriminant analysis”

Xin Zhang and Qing Mai

S1 Additional numerical results

S1.1 Sensitivity of λ

As suggested in Section 5, for the ENDS estimator based on (5.1), we use $\lambda = 1$ if the LDA classifier is applied and select λ based on cross-validation from the set of $\{0, 0.1, \dots, 1\}$ if the QDA classifier is applied. In our experience, the ENDS estimator is not sensitive to the choice of the tuning parameter λ . We use model (Q2) in the simulation (Section 8.1) of the paper for illustration.

Recall from Section 8.1, in this model, we considered the QDA model with $K = 4$ classes and equal number of observations per class, $n_1 = \dots = n_K = 75$. Basis matrix for the envelope $\mathbf{\Gamma} \in \mathbb{R}^{p \times u}$ is first generated randomly and orthogonalized, then $\mathbf{\Gamma}_0 \in \mathbb{R}^{p \times (p-u)}$ is obtained correspondingly so that $(\mathbf{\Gamma}, \mathbf{\Gamma}_0)$ is an orthogonal matrix. The envelope dimension $u = 2$ and the original predictor dimension $p = 15$. Symmetric and positive definite matrices $\mathbf{\Omega}_k \in \mathbb{R}^{u \times u}$, $k = 1, \dots, K$, and $\mathbf{\Omega}_0 \in \mathbb{R}^{(p-u) \times (p-u)}$ are generated as $\mathbf{A}\mathbf{A}^T / \|\mathbf{A}\mathbf{A}^T\|_F$ with $\mathbf{A}$ being a random matrix with the same dimension as $\mathbf{\Omega}_k$ or $\mathbf{\Omega}_0$, and have independent Uniform(0, 1) elements. Within each class k , we generate data from the QDA model, $\mathbf{X} \mid (Y = k) \sim N(\boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)$, $k = 1, \dots, K$, where $\boldsymbol{\mu}_k = \mathbf{\Gamma}\boldsymbol{\eta}_k$ for some randomly generated $\boldsymbol{\eta}_k \in \mathbb{R}^{u \times 1}$ with $N(0, 1)$ elements. We then let $\boldsymbol{\Sigma}_k^* = \exp(-k) \cdot \mathbf{\Gamma}\mathbf{\Omega}_k\mathbf{\Gamma}^T + \mathbf{\Gamma}_0\mathbf{\Omega}_0\mathbf{\Gamma}_0^T$ and standardize it to get $\boldsymbol{\Sigma}_k = 25 \cdot \boldsymbol{\Sigma}_k^* / \|\boldsymbol{\Sigma}_k^*\|_F$, where the scaling constant 25 makes the Bayes' classification error about 0.22.

We varied λ from 0 to 1, and calculated the classification error rate and the principal angle between the estimated and the true subspace spanned by $\hat{\mathbf{\Gamma}}$ and $\mathbf{\Gamma}$, respectively. For each λ value, we repeated 20 times for independent training and testing sets, and also included the results of

$p = 15, K = 4$	Bayes	LDA	ENDS-LDA	QDA	ENDS-QDA	S.E. $\leq$
$n_k = 10$	22.0	48.8	46.7	54.8	38.4	0.5
$n_k = 15$	22.0	46.8	43.8	55.9	34.6	0.4

Table 1: Small n_k simulation setting. The QDA uses pseudoinverse of sample covariance due to insufficient sample size. The last column is the largest standard error among all four estimators.

SIR, SAVE and DR as references. From Figure A1, we confirm that the ENDS estimator is very insensitive to the tuning parameter λ for either the basis estimation or the final classification.

S1.2 ENDS-QDA with small sample size

Using the same data generating procedure as in the previous section, we now consider the scenarios where $15 = n_k = p$ and $10 = n_k < p$. Table 1 summarizes the classification results from LDA, QDA, ENDS-LDA and ENDS-QDA. Recall that, for ENDS-LDA, we use $\lambda = 1$ and the group-specific covariance estimators all equal to the pooled sample covariance: $\mathbf{S}_k(\lambda) = \mathbf{S}$. For ENDS-QDA, by letting $\lambda = 0.1$, we impose a small amount of regularization just to avoid singularity in $\mathbf{S}_k(\lambda) = 0.9\mathbf{S}_k + 0.1\mathbf{S}$. As we discussed at the end of Section 5, if we do not impose such regularization, the symmetric matrix $\mathbf{\Gamma}^T \mathbf{S}_k \mathbf{\Gamma}$ from (5.1) can be singular during iterations because the rank of sample covariance $\mathbf{S}_k$ is at most $(n_k - 1)$. From Table 1, we see that QDA is worse than LDA in terms of classification accuracy because of $n_k \leq p$, but ENDS-QDA with a small amount of regularization is substantially better than LDA and ENDS-LDA.

S2 A technical lemma and its proof

Lemma 1. Under the QDA model (4.1), $\mathcal{S}_{D(Y|\mathbf{X})} = \mathcal{S}_{Y|\mathbf{X}} = \mathcal{L} \cup \mathcal{Q}$ and, moreover, $\mathcal{L} \cup \mathcal{Q} = \mathcal{B} \cup \mathcal{Q}$.

Proof. First, we show that $\mathcal{B} \cup \mathcal{Q}$ is the central subspace $\mathcal{S}_{Y|\mathbf{X}}$. From Cook and Forzani (2009, Theorem 1), the central subspace under the QDA model is the smallest subspace $\mathcal{S}$ such that (i) $\Sigma^{-1}(\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}}) \in \mathcal{S}$ for all $k = 1, \dots, K$; and (ii) $\mathbf{Q}_{\mathcal{S}} \Sigma_1^{-1} = \dots = \mathbf{Q}_{\mathcal{S}} \Sigma_K^{-1}$. Then condition (i) is equivalent to $\mathcal{B} \subseteq \mathcal{S}$ and condition (ii) is equivalent to $\mathcal{Q} \subseteq \mathcal{S}$. Therefore, the smallest such $\mathcal{S}$ is simply $\mathcal{B} \cup \mathcal{Q}$ and is the central subspace $\mathcal{S}_{Y|\mathbf{X}}$.

Next, we show that $\mathcal{L} \cup \mathcal{Q}$ is the central discriminant subspace $\mathcal{S}_{D(Y|\mathbf{X})}$. By combining the QDA Bayes' rule 4.3 with the definitions of $\mathcal{L}$ and $\mathcal{Q}$, it is clear that $\phi(\mathbf{X}) = \phi(\mathbf{P}_{\mathcal{L} \cup \mathcal{Q}} \mathbf{X})$. Hence

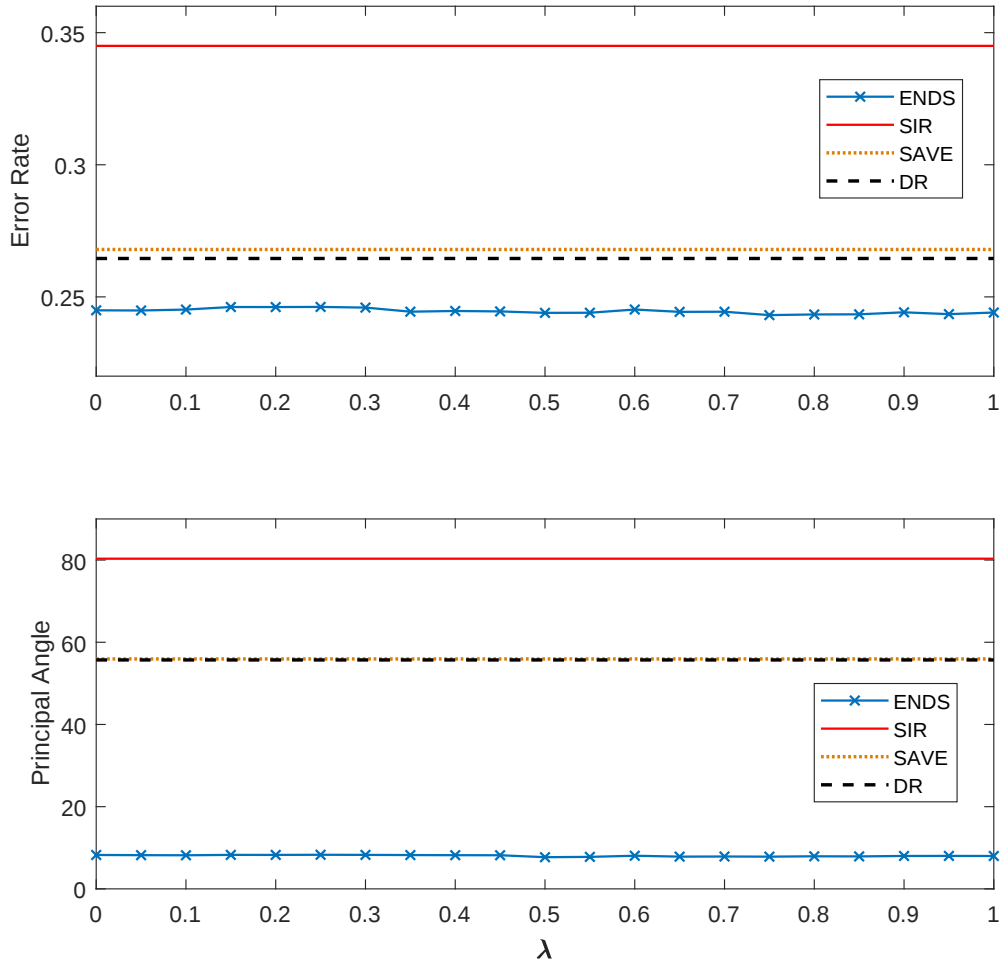


Figure A1: ENDS estimator with different tuning parameter λ . The averaged classification error rate (top plot) and the principal angle between the estimated and the true subspaces (bottom plot).

48 $\mathcal{S}_{D(Y|X)} \subseteq \mathcal{L} \cup \mathcal{Q}$. Then we prove by contradiction for the equality $\mathcal{S}_{D(Y|X)} = \mathcal{L} \cup \mathcal{Q}$. Suppose that
 49 $\mathcal{S}_{D(Y|X)} \subset \mathcal{L} \cup \mathcal{Q}$. Then there exists some vector $\mathbf{v} \notin \mathcal{S}_{D(Y|X)}$ and $\mathbf{v} \in \mathcal{L} \cup \mathcal{Q}$. We can vary the
 50 value of the variable $\mathbf{v}_k^T \mathbf{Q}_{\mathcal{S}_{D(Y|X)}} \mathbf{X}$ without any effect on $\phi(\mathbf{P}_{\mathcal{S}_{D(Y|X)}} \mathbf{X}) = \phi(\mathbf{X})$. We can re-write
 51 the QDA Bayes' rule in the quadratic form of $\phi(\mathbf{X}) = \arg \max_k \{a_k + \mathbf{b}_k^T \mathbf{X} + \mathbf{X}^T \mathbf{C}_k \mathbf{X}\}$ for some
 52 $\mathbf{b}_k \in \mathbb{R}^p$ and $\mathbf{C}_k \in \mathbb{R}^{p \times p}$ such that $\text{span}(\mathbf{b}_k) \subseteq \mathcal{L} \cup \mathcal{Q}$ and $\text{span}(\mathbf{C}_k) \subseteq \mathcal{L} \cup \mathcal{Q}$. Then we notice
 53 that $\mathbf{v}^T \mathbf{Q}_{\mathcal{S}_{D(Y|X)}} \mathbf{X}$ is stochastic (because $\mathbf{v} \notin \mathcal{S}_{D(Y|X)}$) and could still change $\phi(\mathbf{X})$ since the value
 54 of $a_k + \mathbf{b}_k^T \mathbf{X} + \mathbf{X}^T \mathbf{C}_k \mathbf{X}$ can vary from $-\infty$ to $+\infty$.

55 Finally, we show that $\mathcal{L} \cup \mathcal{Q} = \mathcal{B} \cup \mathcal{Q}$. By definition, the union of two subspace of $\mathbb{R}^p$ is
 56 $\mathcal{S}_1 \cup \mathcal{S}_2 = \{\mathbf{a} + \mathbf{b} \mid \mathbf{a} \in \mathcal{S}_1, \mathbf{b} \in \mathcal{S}_2\}$. For vectors in $\mathcal{L}$, we re-write $\Sigma_k^{-1} \boldsymbol{\mu}_k - \Sigma_1^{-1} \boldsymbol{\mu}_1 =$
 57 $(\Sigma_k^{-1} - \Sigma_1^{-1}) \boldsymbol{\mu}_k + \Sigma_1^{-1} (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)$, where the first term is in $\mathcal{Q}$. Therefore,

$$\mathcal{L} \cup \mathcal{Q} = \text{span}\{\Sigma_1^{-1} (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) \mid k = 2, \dots, K\} \cup \mathcal{Q}.$$

58 Similarly, we re-write $\Sigma_1^{-1} (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) = \Sigma^{-1} (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) + (\Sigma_1^{-1} - \Sigma^{-1}) (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)$, where the
 59 first term on the right hand side of the equation is in $\mathcal{B} = \text{span}\{\Sigma^{-1} (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) \mid k = 2, \dots, K\}$.
 60 Therefore, to prove $\mathcal{L} \cup \mathcal{Q} = \mathcal{B} \cup \mathcal{Q}$, it suffices to show that $\text{span}(\Sigma_1^{-1} - \Sigma^{-1}) \subseteq \mathcal{Q}$. From Cook
 61 and Forzani (2009, Proposition 1), we have the following equivalence for any subspace $\mathcal{S} \subseteq \mathbb{R}^p$,

$$\mathbf{Q}_{\mathcal{S}} \Sigma_1^{-1} = \dots = \mathbf{Q}_{\mathcal{S}} \Sigma_K^{-1} \iff \mathbf{Q}_{\mathcal{S}} (\Sigma_k^{-1} - \Sigma^{-1}) = 0, \forall k.$$

62 This implies that, for $\mathcal{Q}$, we have $\mathbf{Q}_{\mathcal{Q}} (\Sigma_k^{-1} - \Sigma^{-1}) = 0$ and hence $\text{span}(\Sigma_1^{-1} - \Sigma^{-1}) \subseteq \mathcal{Q}$.

63 □

64 S3 Proof for Proposition 1

65 By definition, ENDS will contain $\mathcal{S}_{D(Y|X)}$ because of the first condition in (2.1). The second
 66 condition in (2.1) is equivalent to stating that $\mathcal{S}$ is a reducing subspace of $\Sigma_{\mathbf{X}}$. Together, we have
 67 that ENDS is the intersection of all reducing subspace of $\Sigma_{\mathbf{X}}$ that contains $\mathcal{S}_{D(Y|X)}$, which is the
 68 $\Sigma_{\mathbf{X}}$ -envelope of the CDS, $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{S}_{D(Y|X)})$ (cf. Definition 1).

69 The existence and uniqueness of ENDS is thus guaranteed by the existence of CDS. When
 70 CDS exists, the ENDS $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{S}_{D(Y|X)})$ is the intersection of all reducing subspace of $\Sigma_{\mathbf{X}}$ that con-
 71 tains $\mathcal{S}_{D(Y|X)}$ by Definition 1. Trivially, $\mathbb{R}^p$ is a reducing subspace of $\Sigma_{\mathbf{X}}$ that contains $\mathcal{S}_{D(Y|X)}$.
 72 Moreover, if any two subspaces $\mathcal{S}_1 \subseteq \mathbb{R}^p$ and $\mathcal{S}_2 \subseteq \mathbb{R}^p$ both are reducing subspace of $\Sigma_{\mathbf{X}}$ that

contains $\mathcal{S}_{D(Y|\mathbf{X})}$, then we need to show that their intersection is still a reducing subspace of $\Sigma_{\mathbf{X}}$ that contains $\mathcal{S}_{D(Y|\mathbf{X})}$. This implies the uniqueness and existence of the ENDS $\mathcal{E}_{\Sigma_{\mathbf{X}}}(\mathcal{S}_{D(Y|\mathbf{X})})$.

First, it is easy to see that the intersection $\mathcal{S}_1 \cap \mathcal{S}_2$ still contains $\mathcal{S}_{D(Y|\mathbf{X})}$. To see that it is still a reducing subspace of $\Sigma_{\mathbf{X}}$, we note that $\mathcal{S}_1$ is a reducing subspace of $\Sigma_{\mathbf{X}}$ implies that $\mathbf{P}_{\mathcal{S}_1} \Sigma_{\mathbf{X}} \mathbf{Q}_{\mathcal{S}_1} = 0$. Similarly, we have $\mathbf{P}_{\mathcal{S}_2} \Sigma_{\mathbf{X}} \mathbf{Q}_{\mathcal{S}_2} = 0$. Therefore,

$$\begin{aligned}
\mathbf{P}_{\mathcal{S}_1 \cap \mathcal{S}_2} \Sigma_{\mathbf{X}} \mathbf{Q}_{\mathcal{S}_1 \cap \mathcal{S}_2} &= \mathbf{P}_{\mathcal{S}_1} \mathbf{P}_{\mathcal{S}_2} \Sigma_{\mathbf{X}} (\mathbf{I} - \mathbf{P}_{\mathcal{S}_1} \mathbf{P}_{\mathcal{S}_2}) \\
&= \mathbf{P}_{\mathcal{S}_1} \mathbf{P}_{\mathcal{S}_2} \Sigma_{\mathbf{X}} (\mathbf{Q}_{\mathcal{S}_1} + \mathbf{P}_{\mathcal{S}_1} - \mathbf{P}_{\mathcal{S}_1} \mathbf{P}_{\mathcal{S}_2}) \\
&= \mathbf{P}_{\mathcal{S}_1} \mathbf{P}_{\mathcal{S}_2} \Sigma_{\mathbf{X}} (\mathbf{Q}_{\mathcal{S}_1} + \mathbf{P}_{\mathcal{S}_1} \mathbf{Q}_{\mathcal{S}_2}) \\
&= \mathbf{P}_{\mathcal{S}_2} (\mathbf{P}_{\mathcal{S}_1} \Sigma_{\mathbf{X}} \mathbf{Q}_{\mathcal{S}_1}) + \mathbf{P}_{\mathcal{S}_1} (\mathbf{P}_{\mathcal{S}_2} \Sigma_{\mathbf{X}} \mathbf{Q}_{\mathcal{S}_2}) \mathbf{P}_{\mathcal{S}_1} \\
&= 0,
\end{aligned}$$

which implies that $\mathcal{S}_1 \cap \mathcal{S}_2$ is a reducing subspace of $\Sigma_{\mathbf{X}}$. This completes the proof of existence and uniqueness of the ENDS (provided that the CDS exists).

S4 Proof for Proposition 2 and ENDS-LDA parameterization

Under the LDA model assumption (3.1), we prove Proposition 2 by showing in the following that, (2.2) $\implies$ (2.1); (2.1) $\implies$ (3.4); and (3.4) $\implies$ (2.2).

(1) The strong envelope pursuit (2.2) implies the weak envelope pursuit (2.1).

By the definition of the Bayes' rule, the first statement in (2.2), $\Pr(Y = k \mid \mathbf{X}) = \Pr(Y = k \mid \mathbf{P}_{\mathcal{S}} \mathbf{X})$ for all $k = 1, \dots, K$, implies that $\phi(\mathbf{X}) = \phi(\mathbf{P}_{\mathcal{S}} \mathbf{X})$. The second statement in (2.2), $\mathbf{P}_{\mathcal{S}} \mathbf{X} \perp \mathbf{Q}_{\mathcal{S}} \mathbf{X} \mid (Y = k)$ for all $k = 1, \dots, K$, implies that $\mathbf{P}_{\mathcal{S}} \mathbf{X} \perp \mathbf{Q}_{\mathcal{S}} \mathbf{X}$, which further implies $\text{cov}(\mathbf{P}_{\mathcal{S}} \mathbf{X}, \mathbf{Q}_{\mathcal{S}} \mathbf{X}) = 0$.

(2) The weak envelope pursuit (2.1) implies the envelope parameterization (3.4).

First, we show that (2.1) implies $\mathcal{B} \subseteq \mathcal{S}$ in (3.4). Recall from (3.3) that the Bayes' rule under the LDA model (3.1) is $\phi(\mathbf{X}) = \arg \max_k [\log(\pi_k/\pi_1) + (\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)^T \Sigma^{-1} \{\mathbf{X} - (\boldsymbol{\mu}_k + \boldsymbol{\mu}_1)/2\}]$. From this LDA Bayes' rule, we can see that $\phi(\mathbf{X}) = \phi(\mathbf{P}_{\mathcal{B}} \mathbf{X})$ by definitions of $\mathcal{B} = \text{span}(\boldsymbol{\beta}_2, \dots, \boldsymbol{\beta}_K)$ and $\boldsymbol{\beta}_k = \Sigma^{-1}(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)$, $k = 2, \dots, K$. Moreover, $\phi(\mathbf{X}) = \phi(\mathbf{P}_{\mathcal{S}} \mathbf{X})$ implies that $\mathcal{B} \subseteq \mathcal{S}$. To see this, we prove by contradiction. Suppose that for some $\boldsymbol{\beta}_k \notin \mathcal{S}$, then we can vary the value of the univariate variable $\boldsymbol{\beta}_k^T \mathbf{Q}_{\mathcal{S}} \mathbf{X}$ without any effect on $\phi(\mathbf{P}_{\mathcal{S}} \mathbf{X})$. But, by re-writing the Bayes' rule

in the form of $\phi(\mathbf{X}) = \arg \max_k \{\alpha_k + \beta_k^T \mathbf{X}\}$, we notice that $\beta_k^T \mathbf{Q}_S \mathbf{X}$ could still change $\phi(\mathbf{X})$ since the value of $\alpha_k + \beta_k^T \mathbf{X}$ can vary from $-\infty$ to $+\infty$. Therefore, $\phi(\mathbf{X}) = \phi(\mathbf{P}_S \mathbf{X})$ implies that $\mathcal{B} \subseteq \mathcal{S}$.

Next, we need to show $\Sigma = \mathbf{P}_S \Sigma \mathbf{P}_S + \mathbf{Q}_S \Sigma \mathbf{Q}_S$. More specifically, we want to derive $\mathbf{P}_S \Sigma \mathbf{Q}_S = 0$ from $\text{cov}(\mathbf{P}_S \mathbf{X}, \mathbf{Q}_S \mathbf{X}) = \mathbf{P}_S \Sigma_{\mathbf{X}} \mathbf{Q}_S = 0$. Notice that,

$$\begin{aligned} \Sigma_{\mathbf{X}} &= \text{cov}(\mathbf{X}) = \text{cov}\{\mathbb{E}(\mathbf{X} | Y)\} + \mathbb{E}\{\text{cov}(\mathbf{X} | Y)\} \\ &= \sum_k \pi_k (\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}})(\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}})^T + \Sigma \\ &= \Sigma \mathbf{P}_{\mathcal{B}} \mathbf{M}_1 \mathbf{P}_{\mathcal{B}} \Sigma + \Sigma, \end{aligned}$$

where the last equality follows from introducing $\mathbf{M}_1 = \Sigma^{-1} \sum_k \pi_k (\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}})(\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}})^T \Sigma^{-1} \geq 0$ and hence $\text{span}(\mathbf{M}_1) = \mathcal{B}$ by the definition of $\mathcal{B}$. Since we have shown that (2.1) implies $\mathcal{B} \subseteq \mathcal{S}$. We can further write $\Sigma_{\mathbf{X}} = \Sigma \mathbf{P}_S \mathbf{M}_1 \mathbf{P}_S \Sigma + \Sigma$, and thus

$$\begin{aligned} \mathbf{P}_S \Sigma_{\mathbf{X}} \mathbf{Q}_S &= \mathbf{P}_S \Sigma \mathbf{P}_S \mathbf{M}_1 \mathbf{P}_S \Sigma \mathbf{Q}_S + \mathbf{P}_S \Sigma \mathbf{Q}_S \\ &= (\mathbf{P}_S \Sigma \mathbf{P}_S \mathbf{M}_1 + \mathbf{I}_p) \mathbf{P}_S \Sigma \mathbf{Q}_S. \end{aligned}$$

Finally, because $\mathbf{P}_S \Sigma \mathbf{P}_S \mathbf{M}_1 + \mathbf{I}_p > 0$ in the above equation, we have $\mathbf{P}_S \Sigma_{\mathbf{X}} \mathbf{Q}_S = 0$ implies $\mathbf{P}_S \Sigma \mathbf{Q}_S = 0$.

(3) The envelope parameterization (3.4) implies the strong envelope pursuit (2.2).

From the LDA Bayes' rule (3.3), we know that $\Pr(Y = k | \mathbf{X}) = \Pr(Y = k | \mathbf{P}_{\mathcal{B}} \mathbf{X})$, $k = 1, \dots, K$. Therefore, $\mathcal{B} \subseteq \mathcal{S}$ from (3.4) immediately implies the first statement in (2.2). The second statement in (2.2) follows directly from the normal distribution and common within class covariance.

Finally, since we have $\text{span}(\Gamma) = \mathcal{E}_{\Sigma}(\mathcal{B})$. Therefore, $\text{span}(\beta_k) \subseteq \mathcal{B} \subseteq \text{span}(\Gamma)$ and we have (3.5). Next, we show $\text{span}(\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}}) \subseteq \text{span}(\Gamma)$ for (3.6). By definition of $\beta_k = \Sigma^{-1}(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)$, $k = 2, \dots, K$, we can see that $\text{span}(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) = \text{span}(\Sigma \beta_k) \subseteq \text{span}(\Gamma)$, where the last inequality can be derived directly from plugging in (3.5). Notice that $\sum_{j=1}^K \pi_j = 1$, we can write $\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}} = \sum_{j=1}^K \pi_j (\boldsymbol{\mu}_k - \boldsymbol{\mu}_j) = \sum_{j=1}^K \pi_j \{(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1) - (\boldsymbol{\mu}_j - \boldsymbol{\mu}_1)\}$ and have $\text{span}(\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}}) \subseteq \text{span}(\Gamma)$ and $\boldsymbol{\mu}_k - \bar{\boldsymbol{\mu}} = \Gamma \boldsymbol{\alpha}_k$. Finally, the intrinsic constraint $\sum_{j=1}^K \pi_j \boldsymbol{\alpha}_j = 0$ is from $\sum_{j=1}^K \pi_j (\boldsymbol{\mu}_j - \bar{\boldsymbol{\mu}}) = 0$.

S5 Derivations of MLEs for ENDS-LDA

In the LDA model with envelope parameterization (3.6), the conditional log-likelihood of $\mathbf{X}$ given $\mathbf{Y}$ can be written conveniently as

$$\begin{aligned} & L_n(\bar{\boldsymbol{\mu}}, \{\boldsymbol{\alpha}_k\}_{k=1}^K, \boldsymbol{\Gamma}, \boldsymbol{\Omega}, \boldsymbol{\Omega}_0) \\ &= -\frac{n}{2} (\log |\boldsymbol{\Omega}| + \log |\boldsymbol{\Omega}_0|) \\ & - \frac{1}{2} \left\{ \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\boldsymbol{\mu}})^T (\boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\boldsymbol{\mu}}) \right\}, \quad (\text{A1}) \end{aligned}$$

which is to be maximized over the intrinsic mean constraint that $\sum_{j=1}^K \hat{\pi}_j \boldsymbol{\alpha}_j = 0$ where $\hat{\pi}_j = n_j/n$ is the MLE for π_j . The log-determinant terms are derived from the fact that $\log |\boldsymbol{\Sigma}| = \log |\boldsymbol{\Gamma} \boldsymbol{\Omega} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T| = \log |\boldsymbol{\Omega}| + \log |\boldsymbol{\Omega}_0|$.

MLE for $\bar{\boldsymbol{\mu}}$. By taking derivative with respect to $\bar{\boldsymbol{\mu}}$ in (A1) and set it to zero, we have the following equality,

$$0 = \sum_{k=1}^K \sum_{i=1}^{n_k} (\boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) \{ \bar{\boldsymbol{\mu}} - (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k) \},$$

which lead to the MLE for $\bar{\boldsymbol{\mu}}$:

$$\hat{\bar{\boldsymbol{\mu}}} = n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \hat{\boldsymbol{\Gamma}} \hat{\boldsymbol{\alpha}}_k) = n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} \mathbf{X}_{i(k)} - \hat{\boldsymbol{\Gamma}} \sum_{k=1}^K \hat{\pi}_k \hat{\boldsymbol{\alpha}}_k = \bar{\mathbf{X}},$$

where the last equality holds because $\sum_{k=1}^K \hat{\pi}_k \hat{\boldsymbol{\alpha}}_k = 0$ from the intrinsic mean constraint.

MLE for $\boldsymbol{\alpha}_k$. Replacing $\bar{\boldsymbol{\mu}}$ by $\bar{\mathbf{X}}$, the second part of the conditional log-likelihood in (A1) is now partially maximized and can be simplified as follows,

$$\begin{aligned} & -\frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\mathbf{X}})^T (\boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\mathbf{X}}) \\ &= -\frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \}^T \boldsymbol{\Omega}^{-1} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \} \\ & - \frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}})^T \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}). \quad (\text{A2}) \end{aligned}$$

Then because $\boldsymbol{\alpha}_k$ only appears in the above quadratic form, we now need to maximize $F_k(\boldsymbol{\alpha}_k) \equiv -\frac{1}{2} \sum_{i=1}^{n_k} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \}^T \boldsymbol{\Omega}^{-1} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \}$ for $k = 1, \dots, K$, under the constraint

130 that $\sum_{j=1}^K \hat{\pi}_j \alpha_j = 0$. If we ignore the constraint, then the maximizer for each $F_k(\alpha_k)$ will be
 131 $\Gamma^T(\bar{\mathbf{X}}_k - \bar{\mathbf{X}})$, which automatically satisfies the constraint. Therefore, the maximum likelihood
 132 estimators for α_k under the constraint that $\sum_{j=1}^K \hat{\pi}_j \alpha_j = 0$ is simply $\hat{\alpha}_k = \hat{\Gamma}^T(\bar{\mathbf{X}}_k - \bar{\mathbf{X}})$, where $\hat{\Gamma}$
 133 is the MLE for Γ .

134 **MLE for Ω and Ω_0 .** Substitute MLEs for $\bar{\mu}$ and α_k 's into the log-likelihood function, the
 135 partially maximized log-likelihood is $L_n(\Gamma, \Omega, \Omega_0)$ as shown in the following,

$$\begin{aligned}
 -\frac{2}{n}L_n(\Gamma, \Omega, \Omega_0) &\simeq \log |\Omega| + n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)\}^T \Omega^{-1} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)\} \\
 &+ \log |\Omega_0| + n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}})^T \Gamma_0 \Omega_0^{-1} \Gamma_0^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) \\
 &= \log |\Omega| + \text{tr} \left[\Omega^{-1} \Gamma^T \left\{ n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)^T \right\} \Gamma \right] \\
 &+ \log |\Omega_0| + \text{tr} \left[\Omega_0^{-1} \Gamma_0^T \left\{ n^{-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T \right\} \Gamma_0 \right] \\
 &= \log |\Omega| + \text{tr}(\Omega^{-1} \Gamma^T \mathbf{S} \Gamma) + \log |\Omega_0| + \text{tr}(\Omega_0^{-1} \Gamma_0^T \mathbf{S}_X \Gamma_0),
 \end{aligned}$$

136 where $\mathbf{S} = n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)^T$ and $\mathbf{S}_X = n^{-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T$
 137 are the sample covariance matrices for Σ and $\Sigma_X = \text{cov}(\mathbf{X})$. Then the MLEs for Ω and Ω_0
 138 are $\hat{\Omega} = \hat{\Gamma}^T \mathbf{S} \hat{\Gamma}$ and $\hat{\Omega}_0 = \hat{\Gamma}_0^T \mathbf{S}_X \hat{\Gamma}_0$ by noticing that $\mathbf{B} = \arg \min_{\mathbf{A}} \{\log |\mathbf{A}| + \text{tr}(\mathbf{A}^{-1} \mathbf{B})\}$ for
 139 optimization over positive definite matrices.

140 **MLE for Γ .** Finally, the partially maximized log-likelihood satisfies

$$\begin{aligned}
 -\frac{2}{n}L_n(\Gamma, \Omega, \Omega_0) &\simeq \log |\Gamma^T \mathbf{S} \Gamma| + \log |\Gamma_0^T \mathbf{S}_X \Gamma_0| \\
 &\simeq \log |\Gamma^T \mathbf{S} \Gamma| + \log |\Gamma^T \mathbf{S}_X^{-1} \Gamma|,
 \end{aligned}$$

141 which implies that the MLE for Γ is obtained by minimizing $\log |\Gamma^T \mathbf{S} \Gamma| + \log |\Gamma^T \mathbf{S}_X^{-1} \Gamma|$ over
 142 the semi-orthogonal constraint that $\Gamma^T \Gamma = \mathbf{I}_u$, so that the optimization is essentially over the
 143 Grassmann manifold $\mathcal{G}(p, u)$.

144 **MLE for β_k .** From the definition, $\beta_k = \Gamma \eta_k$, we have $\eta_k = \Gamma^T \beta_k = \Gamma^T \Sigma^{-1}(\mu_k - \mu_1)$. By
 145 plugging in the MLEs of Γ , Σ and $(\mu_k - \mu_1)$ to the parameterizations in Proposition 2, we have
 146 $\hat{\eta}_k = (\hat{\Gamma}^T \mathbf{S} \hat{\Gamma})^{-1} \hat{\Gamma}^T(\bar{\mathbf{X}}_k - \bar{\mathbf{X}}_1)$ and hence $\hat{\beta}_k = \hat{\Gamma} \hat{\eta}_k = \mathbf{P}_{\hat{\Gamma}(\mathbf{S})} \mathbf{S}^{-1}(\bar{\mathbf{X}}_k - \bar{\mathbf{X}}_1)$, $k = 2, \dots, K$.

S6 Proof for Proposition 3 and ENDS-QDA parameterization

Similar to the proof for Proposition 2, now under the QDA model (4.1), we prove this Proposition by showing in the following that, (2.2) $\implies$ (2.1); (2.1) $\implies$ (4.5); (4.5) $\implies$ (4.6); and (4.6) $\implies$ (2.2).

(1) The strong envelope pursuit (2.2) implies the weak envelope pursuit (2.1).

By the definition of the Bayes' rule, the first statement in (2.2), $\Pr(Y = k \mid \mathbf{X}) = \Pr(Y = k \mid \mathbf{P}_S \mathbf{X})$ for all $k = 1, \dots, K$, implies that $\phi(\mathbf{X}) = \phi(\mathbf{P}_S \mathbf{X})$. The second statement in (2.2), $\mathbf{P}_S \mathbf{X} \perp\!\!\!\perp \mathbf{Q}_S \mathbf{X} \mid (Y = k)$ for all $k = 1, \dots, K$, implies that $\mathbf{P}_S \mathbf{X} \perp\!\!\!\perp \mathbf{Q}_S \mathbf{X}$, which further implies $\text{cov}(\mathbf{P}_S \mathbf{X}, \mathbf{Q}_S \mathbf{X}) = 0$.

(2) The weak envelope pursuit (2.1) implies the envelope parameterization (4.5).

First, we show that $\phi(\mathbf{X}) = \phi(\mathbf{P}_S \mathbf{X})$ in (2.1) implies $\mathcal{B} \cup \mathcal{Q} \subseteq \mathcal{S}$ in (4.5). From the definition of the central discriminant subspace, $\phi(\mathbf{X}) = \phi(\mathbf{P}_S \mathbf{X})$ implies that $\mathcal{S}_{D(Y|\mathbf{X})} \subseteq \mathcal{S}$. From Lemma 1, we have $\mathcal{B} \cup \mathcal{Q} = \mathcal{S}_{D(Y|\mathbf{X})}$ under the QDA model. Therefore, $\mathcal{B} \cup \mathcal{Q} \subseteq \mathcal{S}$.

Next, it is easy to have $\Sigma = \mathbf{P}_S \Sigma \mathbf{P}_S + \mathbf{Q}_S \Sigma \mathbf{Q}_S$ from $\text{cov}(\mathbf{P}_S \mathbf{X}, \mathbf{Q}_S \mathbf{X}) = 0$ and $\mathcal{B} \cup \mathcal{Q} \subseteq \mathcal{S}$. From the proof of Proposition 2, we know that $\text{cov}(\mathbf{P}_S \mathbf{X}, \mathbf{Q}_S \mathbf{X}) = 0$ and $\mathcal{B} \subseteq \mathcal{S}$ implies $\Sigma = \mathbf{P}_S \Sigma \mathbf{P}_S + \mathbf{Q}_S \Sigma \mathbf{Q}_S$.

(3) The envelope parameterization (4.5) implies the envelope parameterization (4.6).

Assuming the QDA model and $\mathcal{B} \cup \mathcal{Q} \subseteq \mathcal{S}$, we need to show that $\Sigma = \mathbf{P}_S \Sigma \mathbf{P}_S + \mathbf{Q}_S \Sigma \mathbf{Q}_S$ implies $\Sigma_k = \mathbf{P}_S \Sigma_k \mathbf{P}_S + \mathbf{Q}_S \Sigma_k \mathbf{Q}_S$ for all $k = 1, \dots, K$. Let $(\Gamma, \Gamma_0) \in \mathbb{R}^{p \times p}$ be an orthogonal basis matrix such that $\text{span}(\Gamma) = \mathcal{S}$. From Cook and Forzani (2009, Proposition 1), we have the following equality because of $\mathcal{Q} \subseteq \mathcal{S}$,

$$\Sigma_k = \Sigma + \mathbf{P}_{\Gamma(\Sigma)}^T (\Sigma_k - \Sigma) \mathbf{P}_{\Gamma(\Sigma)}, \text{ for all } k = 1, \dots, K, \quad (\text{A3})$$

where $\mathbf{P}_{\Gamma(\Sigma)} = \Gamma(\Gamma^T \Sigma \Gamma)^{-1} \Gamma^T \Sigma$ is the projection matrix onto $\text{span}(\Gamma)$ with Σ inner product. Because $\Sigma = \mathbf{P}_S \Sigma \mathbf{P}_S + \mathbf{Q}_S \Sigma \mathbf{Q}_S$, we can write $\Sigma = \Gamma \Gamma^T \Sigma \Gamma \Gamma^T + \Gamma_0 \Gamma_0^T \Sigma \Gamma_0 \Gamma_0^T$ and $\mathbf{P}_{\Gamma(\Sigma)} = \mathbf{P}_\Gamma = \mathbf{P}_S$. Therefore, (A3) become $\Sigma_k = \Sigma + \mathbf{P}_S (\Sigma_k - \Sigma) \mathbf{P}_S$, which implies that $\mathbf{Q}_S \Sigma_k = \mathbf{Q}_S \Sigma$ for all k . Hence, $\mathbf{Q}_S \Sigma_k \mathbf{P}_S = \mathbf{Q}_S \Sigma \mathbf{P}_S = 0$ since $\mathcal{S}$ is a reducing subspace for Σ . The reducing subspace statement $\Sigma_k = \mathbf{P}_S \Sigma_k \mathbf{P}_S + \mathbf{Q}_S \Sigma_k \mathbf{Q}_S$ is equivalent to $\mathbf{Q}_S \Sigma_k \mathbf{P}_S = 0$.

(4) The envelope parameterization (4.6) implies the strong envelope pursuit (2.2).

The first statement, $\Pr(Y = k \mid \mathbf{X}) = \Pr(Y = k \mid \mathbf{P}_S \mathbf{X})$ in (2.2) is because $\mathcal{S}_{Y|\mathbf{X}} = \mathcal{B} \cup \mathcal{Q} \subseteq \mathcal{S}$ (from Lemma (1)) implies that $Y \perp\!\!\!\perp \mathbf{X} \mid \mathbf{P}_S \mathbf{X}$. The second statement, $\mathbf{P}_S \mathbf{X} \perp\!\!\!\perp \mathbf{Q}_S \mathbf{X} \mid (Y = k)$ is

equivalent to $\Sigma_k = \mathbf{P}_S \Sigma_k \mathbf{P}_S + \mathbf{Q}_S \Sigma_k \mathbf{Q}_S$ under the QDA model assumption.

S7 Derivations of MLEs for ENDS-QDA

In the QDA model with envelope parameterization (4.7), the conditional log-likelihood of $\mathbf{X}$ given Y can be written conveniently as

$$\begin{aligned} & L_n(\bar{\boldsymbol{\mu}}, \{\boldsymbol{\alpha}_k\}_{k=1}^K, \boldsymbol{\Gamma}, \{\boldsymbol{\Omega}_k\}_{k=1}^K, \boldsymbol{\Omega}_0) \\ &= \sum_{k=1}^K \left(-\frac{n_k}{2}\right) (\log |\boldsymbol{\Omega}_k| + \log |\boldsymbol{\Omega}_0|) \\ & - \frac{1}{2} \left\{ \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\boldsymbol{\mu}})^T (\boldsymbol{\Gamma} \boldsymbol{\Omega}_k^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\boldsymbol{\mu}}) \right\}, \quad (\text{A4}) \end{aligned}$$

which is to be maximized over the intrinsic mean constraint that $\sum_{j=1}^K \hat{\pi}_j \boldsymbol{\alpha}_j = 0$ where $\hat{\pi}_j = n_j/n$ is the MLE for π_j . The log-determinant terms are derived from the fact that $\log |\Sigma_k| = \log |\boldsymbol{\Gamma} \boldsymbol{\Omega}_k \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T| = \log |\boldsymbol{\Omega}_k| + \log |\boldsymbol{\Omega}_0|$.

MLE for $\bar{\boldsymbol{\mu}}$. By taking derivative with respect to $\bar{\boldsymbol{\mu}}$ in (A4) and set it to zero, we have the following equality,

$$0 = \sum_{k=1}^K \sum_{i=1}^{n_k} (\boldsymbol{\Gamma} \boldsymbol{\Omega}_k^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) \{ \bar{\boldsymbol{\mu}} - (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k) \},$$

which lead to the MLE for $\bar{\boldsymbol{\mu}}$:

$$\hat{\bar{\boldsymbol{\mu}}} = n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \hat{\boldsymbol{\Gamma}} \hat{\boldsymbol{\alpha}}_k) = n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} \mathbf{X}_{i(k)} - \hat{\boldsymbol{\Gamma}} \sum_{k=1}^K \hat{\pi}_k \hat{\boldsymbol{\alpha}}_k = \bar{\mathbf{X}},$$

where the last equality holds because $\sum_{k=1}^K \hat{\pi}_k \hat{\boldsymbol{\alpha}}_k = 0$ from the intrinsic mean constraint.

MLE for $\boldsymbol{\alpha}_k$. Replacing $\bar{\boldsymbol{\mu}}$ by $\bar{\mathbf{X}}$, the second part of the conditional log-likelihood in (A4) is now partially maximized and can be simplified as follows,

$$\begin{aligned} & -\frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\mathbf{X}})^T (\boldsymbol{\Gamma} \boldsymbol{\Omega}_k^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T) (\mathbf{X}_{i(k)} - \boldsymbol{\Gamma} \boldsymbol{\alpha}_k - \bar{\mathbf{X}}) \\ &= -\frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \}^T \boldsymbol{\Omega}_k^{-1} \{ \boldsymbol{\Gamma}^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \boldsymbol{\alpha}_k \} \\ & - \frac{1}{2} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}})^T \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} \boldsymbol{\Gamma}_0^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}). \quad (\text{A5}) \end{aligned}$$

Then because α_k only appears in the above quadratic form, we now need to maximize $F_k(\alpha_k) \equiv -\frac{1}{2} \sum_{i=1}^{n_k} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \alpha_k\}^T \Omega_k^{-1} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) - \alpha_k\}$ for $k = 1, \dots, K$, under the constraint that $\sum_{j=1}^K \hat{\pi}_j \alpha_j = 0$. If we ignore the constraint, then the maximizer for each $F_k(\alpha_k)$ will be $\Gamma^T(\bar{\mathbf{X}}_k - \bar{\mathbf{X}})$, which automatically satisfies the constraint. Therefore, the maximum likelihood estimators for α_k under the constraint that $\sum_{j=1}^K \hat{\pi}_j \alpha_j = 0$ is simply $\hat{\alpha}_k = \hat{\Gamma}^T(\bar{\mathbf{X}}_k - \bar{\mathbf{X}})$, where $\hat{\Gamma}$ is the MLE for Γ .

MLE for Ω_k and Ω_0 . Substitute MLEs for $\bar{\mu}$ and α_k 's into the log-likelihood function, the partially maximized log-likelihood is $L_n(\Gamma, \Omega, \Omega_0)$ as shown in the following,

$$\begin{aligned} -\frac{2}{n} L_n(\Gamma, \{\Omega_k\}_{k=1}^K, \Omega_0) &\simeq n^{-1} \sum_{k=1}^K \left[n_k \log |\Omega_k| + \sum_{i=1}^{n_k} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)\}^T \Omega_k^{-1} \{\Gamma^T(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)\} \right] \\ &+ \log |\Omega_0| + n^{-1} \sum_{k=1}^K \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}})^T \Gamma_0 \Omega_0^{-1} \Gamma_0^T (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}) \\ &= n^{-1} \sum_{k=1}^K \left(n_k \log |\Omega_k| + \text{tr} \left[\Omega_k^{-1} \Gamma^T \left\{ \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)^T \right\} \Gamma \right] \right) \\ &+ \log |\Omega_0| + \text{tr} \left[\Omega_0^{-1} \Gamma_0^T \left\{ n^{-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T \right\} \Gamma_0 \right] \\ &= \sum_{k=1}^K \frac{n_k}{n} \{ \log |\Omega_k| + \text{tr}(\Omega_k^{-1} \Gamma^T \mathbf{S}_k \Gamma) \} + \log |\Omega_0| + \text{tr}(\Omega_0^{-1} \Gamma_0^T \mathbf{S}_X \Gamma_0), \end{aligned}$$

where $\mathbf{S}_k = n_k^{-1} \sum_{i=1}^{n_k} (\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)(\mathbf{X}_{i(k)} - \bar{\mathbf{X}}_k)^T$ and $\mathbf{S}_X = n^{-1} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T$ are the sample covariance matrices for Σ_k and $\Sigma_X = \text{cov}(\mathbf{X})$. Then the MLEs for Ω_k and Ω_0 are $\hat{\Omega}_k = \hat{\Gamma}^T \mathbf{S}_k \hat{\Gamma}$ and $\hat{\Omega}_0 = \hat{\Gamma}_0^T \mathbf{S}_X \hat{\Gamma}_0$ by noticing that $\mathbf{B} = \arg \min_{\mathbf{A}} \{\log |\mathbf{A}| + \text{tr}(\mathbf{A}^{-1} \mathbf{B})\}$ for optimization over positive definite matrices.

MLE for Γ . Finally, the partially maximized log-likelihood satisfies

$$\begin{aligned} -\frac{2}{n} L_n(\Gamma) &\simeq \sum_{k=1}^K \frac{n_k}{n} \log |\Gamma^T \mathbf{S}_k \Gamma| + \log |\Gamma_0^T \mathbf{S}_X \Gamma_0| \\ &\simeq \sum_{k=1}^K \frac{n_k}{n} \log |\Gamma^T \mathbf{S}_k \Gamma| + \log |\Gamma^T \mathbf{S}_X^{-1} \Gamma|, \end{aligned}$$

which implies that the MLE for Γ is obtained by minimizing $\sum_{k=1}^K \frac{n_k}{n} \log |\Gamma^T \mathbf{S}_k \Gamma| + \log |\Gamma^T \mathbf{S}_X^{-1} \Gamma|$ over the semi-orthogonal constraint that $\Gamma^T \Gamma = \mathbf{I}_u$, so that the optimization is essentially over the Grassmann manifold $\mathcal{G}(p, u)$.

MLE for β_k, Σ_k and Σ . The conclusion follows by plugging in the MLEs of Γ, Ω_k and Ω_0 to the parameterizations in Proposition 3.

S8 Proof for Proposition 4 and Equation (7.3)

The parameters of the original LDA model are $\mu_1, \dots, \mu_K$ and Σ . We can also re-parameterize them as $\bar{\mu}, \beta_2, \dots, \beta_K$ and Σ . For all methods, the estimator of $\bar{\mu}$ is simply the sample mean of $\mathbf{X}$, which will be asymptotically independent of all other estimators. We thus focus on the asymptotic properties of the discriminant directions $\beta_k, k = 2, \dots, K$.

The parameters involved in the coordinate representation of the ENDS-LDA model are Γ, Ω, Ω_0 and $\eta_k, k = 2, \dots, K$. We focus on the asymptotic properties of the estimable functions $\beta_k = \Sigma^{-1}(\mu_k - \mu_1) = \Gamma\eta_k, k = 2, \dots, K$, and the covariance $\Sigma = \Gamma\Omega\Gamma^T + \Gamma_0\Omega_0\Gamma_0^T$. Specifically, we study the asymptotic covariances of the following parameters ϕ and estimable functions $\mathbf{h}$,

$$\phi = \begin{pmatrix} \text{vec}(\Gamma) \\ \eta_2 \\ \vdots \\ \eta_K \\ \text{vech}(\Omega) \\ \text{vech}(\Omega_0) \end{pmatrix}, \quad \mathbf{h} = \begin{pmatrix} \beta_2 \\ \vdots \\ \beta_K \\ \text{vech}(\Sigma) \end{pmatrix}. \quad (\text{A6})$$

Since $\mathbf{h}$ is a vector of length $(K-1)p + p(p+1)/2$ and ϕ is a vector of length $pu - u^2 + (K-1)u + u(u+1)/2 + (p-u)(p-u+1)/2 = (K-1)u + p(p+1)/2$, the ENDS reduces the total number of parameters by $(K-1)(p-u)$.

We let $\tilde{\mathbf{h}} = (\tilde{\beta}_2^T, \dots, \tilde{\beta}_K^T, \text{vech}^T(\mathbf{S}))^T$ be the standard MLE of the LDA model parameters, where recall that $\tilde{\beta}_k = \mathbf{S}^{-1}(\bar{\mathbf{X}}_k - \bar{\mathbf{X}}_1)$.

The asymptotic covariance matrix $\mathbf{C}_1$ of $\tilde{\psi}_1 = (\bar{\mathbf{X}}_1^T, \dots, \bar{\mathbf{X}}_K^T, \text{vech}^T(\mathbf{S}))^T$, the MLE of $\psi_1 = (\mu_1^T, \dots, \mu_K^T, \text{vech}^T(\Sigma))$, can be derived straightforwardly as

$$\mathbf{C}_1 = \mathbf{J}_1^{-1}, \quad \mathbf{J}_1 = \text{diag} \left\{ \pi_1 \Sigma^{-1}, \dots, \pi_K \Sigma^{-1}, \frac{1}{2} \mathbf{E}_p^T (\Sigma^{-1} \otimes \Sigma^{-1}) \mathbf{E}_p \right\}, \quad (\text{A7})$$

where $\mathbf{J}_1$ is the Fisher information matrix of parameter vector ψ_1 .

Next, the asymptotic covariance matrix $\mathbf{C}_2$ of $\tilde{\psi}_2 \equiv (\bar{\mathbf{X}}_2^T - \bar{\mathbf{X}}_1^T, \dots, \bar{\mathbf{X}}_K^T - \bar{\mathbf{X}}_1^T, \text{vech}^T(\mathbf{S}))^T$,

the MLE of $\psi_2 = (\boldsymbol{\mu}_2^T - \boldsymbol{\mu}_1^T \dots, \boldsymbol{\mu}_K^T - \boldsymbol{\mu}_1^T, \text{vech}^T(\boldsymbol{\Sigma}))^T$ is thus

$$\mathbf{C}_2 = \mathbf{L}_1 \mathbf{C}_1 \mathbf{L}_1^T = \begin{pmatrix} (\pi_2^{-1} + \pi_1^{-1})\boldsymbol{\Sigma} & \cdots & \pi_1^{-1}\boldsymbol{\Sigma} & 0 \\ \vdots & \ddots & \vdots & 0 \\ \pi_1^{-1}\boldsymbol{\Sigma} & \cdots & (\pi_K^{-1} + \pi_1^{-1})\boldsymbol{\Sigma} & 0 \\ 0 & 0 & 0 & \{\frac{1}{2}\mathbf{E}_p^T(\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1})\mathbf{E}_p\}^{-1} \end{pmatrix}, \quad (\text{A8})$$

where the matrix $\mathbf{L}_1$ is the linear map between the estimators: $(\bar{\mathbf{X}}_2^T - \bar{\mathbf{X}}_1^T, \dots, \bar{\mathbf{X}}_K^T - \bar{\mathbf{X}}_1^T, \text{vech}^T(\mathbf{S}))^T = \mathbf{L}_1 \cdot (\bar{\mathbf{X}}_1^T, \dots, \bar{\mathbf{X}}_K^T, \text{vech}^T(\mathbf{S}))^T$. Specifically,

$$\mathbf{L}_1 = \text{diag} \left\{ (-\mathbf{1}_{(K-1)}, \mathbf{I}_{(K-1)}) \otimes \mathbf{I}_p, \mathbf{I}_{p(p+1)/2} \right\}, \quad (\text{A9})$$

We can write $\mathbf{C}_2$ more compactly as

$$\mathbf{C}_2 = \text{diag} \left[\mathbf{A} \otimes \boldsymbol{\Sigma}, \left\{ \frac{1}{2}\mathbf{E}_p^T(\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1})\mathbf{E}_p \right\}^{-1} \right], \quad (\text{A10})$$

where $\mathbf{A} \in \mathbb{R}^{(K-1) \times (K-1)}$ is a symmetric matrix has diagonal elements $(\pi_k^{-1} + \pi_1^{-1})$ and off-diagonal π_1^{-1} .

An alternative way of reaching these results is by calculating the Fisher information matrix of $\psi_3 = (\bar{\boldsymbol{\mu}}^T, \boldsymbol{\mu}_2^T - \boldsymbol{\mu}_1^T \dots, \boldsymbol{\mu}_K^T - \boldsymbol{\mu}_1^T, \text{vech}^T(\boldsymbol{\Sigma}))^T = (\bar{\boldsymbol{\mu}}^T, \boldsymbol{\psi}_2^T)^T$. It is easy to see that there is a full-rank square matrix to transform between ψ_1 and ψ_3 . Let $\psi_1 = \mathbf{L}_3 \psi_3$, then $\mathbf{L}_3 = \partial \psi_1 / \partial \psi_3$ and $\mathbf{J}_3 = \mathbf{L}_3^T \mathbf{J}_1 \mathbf{L}_3$. Straightforward calculation leads to block diagonal structure as $\mathbf{J}_3 = \text{diag} \left\{ \boldsymbol{\Sigma}^{-1}, \mathbf{A}^{-1} \otimes \boldsymbol{\Sigma}^{-1}, \frac{1}{2}\mathbf{E}_p^T(\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1})\mathbf{E}_p \right\}$, which also verifies that $\bar{\mathbf{X}}$ is asymptotically independent of $(\bar{\mathbf{X}}_2^T - \bar{\mathbf{X}}_1^T, \dots, \bar{\mathbf{X}}_K^T - \bar{\mathbf{X}}_1^T, \text{vech}^T(\mathbf{S}))^T$. Then write $\mathbf{J}_3 = \text{diag}\{\boldsymbol{\Sigma}^{-1}, \mathbf{J}_2\}$, we have the Fisher information matrix for ψ_2 ,

$$\mathbf{J}_2 = \mathbf{C}_2^{-1} = \text{diag} \left\{ \mathbf{A}^{-1} \otimes \boldsymbol{\Sigma}^{-1}, \frac{1}{2}\mathbf{E}_p^T(\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1})\mathbf{E}_p \right\}. \quad (\text{A11})$$

We can write $\mathbf{A} = \text{diag}(\pi_2^{-1}, \dots, \pi_K^{-1}) + \pi_1^{-1} \mathbf{1}_{K-1} \mathbf{1}_{K-1}^T$, then apply the Woodbury matrix identity, we have

$$\begin{aligned} \mathbf{A}^{-1} &= \text{diag}(\pi_2, \dots, \pi_K) \\ &- \text{diag}(\pi_2, \dots, \pi_K) \mathbf{1}_{K-1} (\pi_1 + \mathbf{1}_{K-1}^T \text{diag}(\pi_2, \dots, \pi_K) \mathbf{1}_{K-1})^{-1} \mathbf{1}_{K-1}^T \text{diag}(\pi_2, \dots, \pi_K) \\ &= \text{diag}(\pi_2, \dots, \pi_K) - \text{diag}(\pi_2, \dots, \pi_K) \mathbf{1}_{K-1} \mathbf{1}_{K-1}^T \text{diag}(\pi_2, \dots, \pi_K) \end{aligned}$$

that is, $[\mathbf{A}^{-1}]_{ij}$ is $\pi_j(1 - \pi_j)$ for $i = j$ and is $-\pi_i \pi_j$ for $i \neq j$.

240 Finally, the asymptotic covariance matrix $\mathbf{C}_h$ of $\tilde{\mathbf{h}}$ can be calculated from the one-to-one
 241 transformation between ψ_2 and $\mathbf{h}$. Specifically, we noticed the simple linear transformation that
 242 $\boldsymbol{\mu}_k - \boldsymbol{\mu}_1 = \boldsymbol{\Sigma}\boldsymbol{\beta}_k$, and thus,

$$\frac{\partial(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)}{\partial\boldsymbol{\beta}_k} = \boldsymbol{\Sigma} \quad (\text{A12})$$

$$\frac{\partial(\boldsymbol{\mu}_k - \boldsymbol{\mu}_1)}{\partial\text{vech}(\boldsymbol{\Sigma})} = (\boldsymbol{\beta}_k^T \otimes \mathbf{I}_p) \frac{\partial\text{vec}(\boldsymbol{\Sigma})}{\partial\text{vech}(\boldsymbol{\Sigma})} = (\boldsymbol{\beta}_k^T \otimes \mathbf{I}_p) \mathbf{E}_p. \quad (\text{A13})$$

244 where $\mathbf{E}_p \in \mathbb{R}^{p^2 \times p(p+1)/2}$ is the expansion matrix such that $\text{vec}(\boldsymbol{\Sigma}) = \mathbf{E}_p \text{vech}(\boldsymbol{\Sigma})$. Therefore,

$$\mathbf{C}_h = \mathbf{J}_h^{-1} = (\mathbf{G}^T \mathbf{J}_2 \mathbf{G})^{-1}, \quad \text{where } \mathbf{G} = \frac{\partial\psi_2}{\partial\mathbf{h}} = \begin{pmatrix} \mathbf{I}_{K-1} \otimes \boldsymbol{\Sigma} & (\mathbf{B}^T \otimes \mathbf{I}_p) \mathbf{E}_p \\ 0 & \mathbf{I}_{p(p+1)/2} \end{pmatrix}, \quad (\text{A14})$$

245 where $\mathbf{B} = (\boldsymbol{\beta}_2, \dots, \boldsymbol{\beta}_K) \in \mathbb{R}^{p \times (K-1)}$.

246 Combining (A11) and (A14), we have

$$\begin{aligned} \mathbf{J}_h &= \mathbf{G}^T \mathbf{J}_2 \mathbf{G} \\ &= \begin{pmatrix} \mathbf{I}_{K-1} \otimes \boldsymbol{\Sigma} & 0 \\ \mathbf{E}_p^T (\mathbf{B} \otimes \mathbf{I}_p) & \mathbf{I}_{p(p+1)/2} \end{pmatrix} \text{diag} \left\{ \mathbf{A}^{-1} \otimes \boldsymbol{\Sigma}^{-1}, \frac{1}{2} \mathbf{E}_p^T (\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \right\} \begin{pmatrix} \mathbf{I}_{K-1} \otimes \boldsymbol{\Sigma} & (\mathbf{B}^T \otimes \mathbf{I}_p) \mathbf{E}_p \\ 0 & \mathbf{I}_{p(p+1)/2} \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{I}_{K-1} \otimes \boldsymbol{\Sigma} & 0 \\ \mathbf{E}_p^T (\mathbf{B} \otimes \mathbf{I}_p) & \mathbf{I}_{p(p+1)/2} \end{pmatrix} \begin{pmatrix} \mathbf{A}^{-1} \otimes \mathbf{I}_p & (\mathbf{A}^{-1} \mathbf{B}^T \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \\ 0 & \frac{1}{2} \mathbf{E}_p^T (\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \end{pmatrix} \\ &= \begin{pmatrix} \mathbf{A}^{-1} \otimes \boldsymbol{\Sigma} & (\mathbf{A}^{-1} \mathbf{B}^T \otimes \mathbf{I}_p) \mathbf{E}_p \\ \mathbf{E}_p^T (\mathbf{B} \mathbf{A}^{-1} \otimes \mathbf{I}_p) & \mathbf{E}_p^T (\mathbf{B} \mathbf{A}^{-1} \mathbf{B}^T \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p + \frac{1}{2} \mathbf{E}_p^T (\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \end{pmatrix}, \end{aligned}$$

247 whose inverse can be obtained using the block matrix inversion formula,

$$\mathbf{C}_h = \mathbf{J}_h^{-1} = \begin{pmatrix} \mathbf{C}_B & \mathbf{C}_{B\Sigma} \\ \mathbf{C}_{B\Sigma}^T & \mathbf{C}_\Sigma \end{pmatrix}, \quad (\text{A15})$$

248 where the top left block is of primary interests. We calculate the asymptotic covariance of $\text{vec}(\mathbf{B}) =$
 249 $(\boldsymbol{\beta}_2^T, \dots, \boldsymbol{\beta}_K^T)^T$, i.e. $\mathbf{C}_B \in \mathbb{R}^{p(K-1) \times p(K-1)}$, in the following,

$$\begin{aligned} \mathbf{C}_B &= \mathbf{A} \otimes \boldsymbol{\Sigma}^{-1} + (\mathbf{B}^T \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \left\{ \frac{1}{2} \mathbf{E}_p^T (\boldsymbol{\Sigma}^{-1} \otimes \boldsymbol{\Sigma}^{-1}) \mathbf{E}_p \right\}^{-1} \mathbf{E}_p^T (\mathbf{B} \otimes \boldsymbol{\Sigma}^{-1}) \\ &= \mathbf{A} \otimes \boldsymbol{\Sigma}^{-1} + (\mathbf{B}^T \otimes \boldsymbol{\Sigma}^{-1}) 2\mathbf{P}_{\mathbf{E}_p} (\boldsymbol{\Sigma} \otimes \boldsymbol{\Sigma}) \mathbf{P}_{\mathbf{E}_p} (\mathbf{B} \otimes \boldsymbol{\Sigma}^{-1}) \\ &= \mathbf{A} \otimes \boldsymbol{\Sigma}^{-1} + 2(\mathbf{B}^T \boldsymbol{\Sigma} \otimes \mathbf{I}_p) \mathbf{P}_{\mathbf{E}_p} (\mathbf{B} \otimes \boldsymbol{\Sigma}^{-1}) \\ &= \mathbf{A} \otimes \boldsymbol{\Sigma}^{-1} + (\mathbf{B}^T \boldsymbol{\Sigma} \mathbf{B} \otimes \boldsymbol{\Sigma}^{-1}) + (\mathbf{B}^T \otimes \mathbf{B}) \mathbf{K}_{p,(k-1)}. \end{aligned}$$

250 Finally, for binary case, $\mathbf{B} = \boldsymbol{\beta} \in \mathbb{R}^p$ is a vector itself. We have $\mathbf{K}_{p,1} = \mathbf{I}_p$ since $\text{vec}(\boldsymbol{\beta}) =$
 251 $\text{vec}(\boldsymbol{\beta}^T)$. The conclusion follows by plugging in $\boldsymbol{\beta} = (\Delta, 0, \dots, 0)^T$ and $\boldsymbol{\Sigma} = \mathbf{I}_p$ and noticing $\mathbf{C}_B$
 252 will coincide with Efron's (1975) results.

S9 Proof for Proposition 5

Recall that, $\beta = \Sigma^{-1}(\mu_2 - \mu_1) = \Gamma\eta$ under the envelope LDA model, therefore $\mu_{Diff} \equiv (\mu_2 - \mu_1) = \Sigma\beta = \Gamma\Omega\eta$ under the envelope parameterization. We define the following parameters and estimable functions,

$$\phi = \begin{pmatrix} \phi_1 \\ \phi_2 \\ \phi_3 \\ \phi_4 \end{pmatrix} = \begin{pmatrix} \eta \\ \text{vec}(\Gamma) \\ \text{vech}(\Omega) \\ \text{vech}(\Omega_0) \end{pmatrix}, \mathbf{h} = \mathbf{h}(\phi) = \begin{pmatrix} \beta \\ \text{vech}(\Sigma) \end{pmatrix}, \mathbf{g} = \mathbf{g}(\phi) = \begin{pmatrix} \mu_{Diff} \\ \text{vech}(\Sigma) \end{pmatrix}. \quad (\text{A16})$$

In the proof of Proposition 4, we have already obtained the Fisher information of $\mathbf{h}$ and $\mathbf{g}$ under the classical LDA model,

$$\mathbf{J}_g = \begin{pmatrix} \pi_1\pi_2\Sigma^{-1} & 0 \\ 0 & \frac{1}{2}\mathbf{E}_p^T(\Sigma^{-1} \otimes \Sigma^{-1})\mathbf{E}_p \end{pmatrix}; \quad (\text{A17})$$

$$\begin{aligned} \mathbf{J}_h &= \begin{pmatrix} \pi_1\pi_2\Sigma & (\pi_1\pi_2\beta^T \otimes \mathbf{I}_p)\mathbf{E}_p \\ \mathbf{E}_p^T(\pi_1\pi_2\beta \otimes \mathbf{I}_p) & \mathbf{E}_p^T(\pi_1\pi_2\beta\beta^T \otimes \Sigma^{-1})\mathbf{E}_p + \frac{1}{2}\mathbf{E}_p^T(\Sigma^{-1} \otimes \Sigma^{-1})\mathbf{E}_p \end{pmatrix} \\ &= \pi_1\pi_2 \begin{pmatrix} \Sigma & (\beta^T \otimes \mathbf{I}_p)\mathbf{E}_p \\ \mathbf{E}_p^T(\beta \otimes \mathbf{I}_p) & \mathbf{E}_p^T\{(\beta\beta^T + (2\pi_1\pi_2)^{-1}\Sigma^{-1}) \otimes \Sigma^{-1}\}\mathbf{E}_p \end{pmatrix} \end{aligned}$$

The Fisher information of ϕ is hence, $\mathbf{J}_\phi = \mathbf{G}^T \mathbf{J}_g \mathbf{G}$, where $\mathbf{G}$ is

$$\mathbf{G} = \frac{\partial \mathbf{g}}{\partial \phi} = \begin{pmatrix} \Gamma\Omega & \eta^T\Omega \otimes \mathbf{I}_p & (\eta^T \otimes \Gamma)\mathbf{E}_u & 0 \\ 0 & 2\mathbf{C}_p(\Gamma\Omega \otimes \mathbf{I}_p - \Gamma \otimes \Gamma_0\Omega_0\Gamma_0^T) & \mathbf{C}_p(\Gamma \otimes \Gamma)\mathbf{E}_u & \mathbf{C}_p(\Gamma_0 \otimes \Gamma_0)\mathbf{E}_{p-u} \end{pmatrix}. \quad (\text{A18})$$

The asymptotic covariance of envelope estimator $\hat{\mathbf{h}}$ is

$$\mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T = \mathbf{H} \mathbf{J}_\phi^\dagger \mathbf{H}^T, \quad (\text{A19})$$

where

$$\mathbf{H} = \frac{\partial \mathbf{h}}{\partial \phi} = \begin{pmatrix} \Gamma & \eta^T \otimes \mathbf{I}_p & 0 & 0 \\ 0 & 2\mathbf{C}_p(\Gamma\Omega \otimes \mathbf{I}_p - \Gamma \otimes \Gamma_0\Omega_0\Gamma_0^T) & \mathbf{C}_p(\Gamma \otimes \Gamma)\mathbf{E}_u & \mathbf{C}_p(\Gamma_0 \otimes \Gamma_0)\mathbf{E}_{p-u} \end{pmatrix}. \quad (\text{A20})$$

Following Cook et al. (2010), we have the following transformation,

$$\mathbf{H}(\mathbf{H}^T \mathbf{J}_h \mathbf{H})^\dagger \mathbf{H}^T = \mathbf{H}_1(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^\dagger \mathbf{H}_1^T, \quad (\text{A21})$$

264 where $\mathbf{H}_1$ satisfies $\text{span}(\mathbf{H}_1) = \text{span}(\mathbf{H})$ and is given in the following,

$$\begin{aligned}\mathbf{H}_1 &= \begin{pmatrix} \mathbf{\Gamma} & \boldsymbol{\eta}^T \otimes \mathbf{\Gamma}_0 & 0 & 0 \\ 0 & 2\mathbf{C}_p(\mathbf{\Gamma}\mathbf{\Omega} \otimes \mathbf{\Gamma}_0 - \mathbf{\Gamma} \otimes \mathbf{\Gamma}_0\mathbf{\Omega}_0) & \mathbf{C}_p(\mathbf{\Gamma} \otimes \mathbf{\Gamma})\mathbf{E}_u & \mathbf{C}_p(\mathbf{\Gamma}_0 \otimes \mathbf{\Gamma}_0)\mathbf{E}_{p-u} \end{pmatrix} \\ &\equiv \begin{pmatrix} \mathbf{H}_{11} & \mathbf{H}_{12} & \mathbf{H}_{13} & \mathbf{H}_{14} \end{pmatrix}.\end{aligned}$$

265 We next calculate each block in $\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1$ as follows,

$$\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1 = \pi_1 \pi_2 \begin{pmatrix} \mathbf{\Omega} & 0 & (\boldsymbol{\eta}^T \otimes \mathbf{I}_u)\mathbf{E}_u & 0 \\ 0 & \mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12} & 0 & 0 \\ \mathbf{E}_u^T(\boldsymbol{\eta} \otimes \mathbf{I}_u) & 0 & \mathbf{E}_u^T\{(\boldsymbol{\eta}\boldsymbol{\eta}^T + (2\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1}) \otimes \mathbf{\Omega}^{-1}\}\mathbf{E}_u & 0 \\ 0 & 0 & 0 & \mathbf{H}_{14}^T \mathbf{J}_h \mathbf{H}_{14} \end{pmatrix}. \quad (\text{A22})$$

266 Since the top blocks of $\mathbf{H}_{13}$ and $\mathbf{H}_{14}$ are both zero, they have no contribution to the asymptotic
267 covariance of $\hat{\boldsymbol{\beta}}$ (top block of $\mathbf{h}$). We hence have,

$$\mathbf{V}_\beta = \mathbf{\Gamma}\{(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1}\}_{11}\mathbf{\Gamma}^T + (\boldsymbol{\eta}^T \otimes \mathbf{I}_p)(\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12})^{-1}(\boldsymbol{\eta} \otimes \mathbf{I}_p), \quad (\text{A23})$$

268 where $\{(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1}\}_{11}$ is the top left $p \times p$ block of $(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1}$.

269 The first term in (A23) can be calculated as,

$$\begin{aligned}\{(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1}\}_{11} &= (\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1} \\ &+ (\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1}(\boldsymbol{\eta}^T \otimes \mathbf{I}_u)\mathbf{E}_u\{(2\pi_1\pi_2)^{-1}\mathbf{E}_u^T(\mathbf{\Omega}^{-1} \otimes \mathbf{\Omega}^{-1})\mathbf{E}_u\}^{-1}\mathbf{E}_u^T(\boldsymbol{\eta} \otimes \mathbf{I}_u)\mathbf{\Omega}^{-1} \\ &= (\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1} \\ &+ (\pi_1\pi_2)^{-1}(\boldsymbol{\eta}^T \otimes \mathbf{\Omega}^{-1})\mathbf{E}_u\{(2\pi_1\pi_2)^{-1}\mathbf{E}_u^T(\mathbf{\Omega}^{-1} \otimes \mathbf{\Omega}^{-1})\mathbf{E}_u\}^{-1}\mathbf{E}_u^T(\boldsymbol{\eta} \otimes \mathbf{\Omega}^{-1}) \\ &= (\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1} + (\boldsymbol{\eta}^T \otimes \mathbf{\Omega}^{-1})2\mathbf{P}_{\mathbf{E}_u}(\mathbf{\Omega} \otimes \mathbf{\Omega})(\boldsymbol{\eta} \otimes \mathbf{\Omega}^{-1}) \\ &= (\pi_1\pi_2)^{-1}\mathbf{\Omega}^{-1} + (\boldsymbol{\eta}^T \mathbf{\Omega} \boldsymbol{\eta})\mathbf{\Omega}^{-1} + (\boldsymbol{\eta}^T \otimes \boldsymbol{\eta}),\end{aligned}$$

270 therefore,

$$\begin{aligned}\mathbf{\Gamma}\{(\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1}\}_{11}\mathbf{\Gamma}^T &= (\pi_1\pi_2)^{-1}\mathbf{\Gamma}\mathbf{\Omega}^{-1}\mathbf{\Gamma}^T + (\boldsymbol{\eta}^T \mathbf{\Omega} \boldsymbol{\eta})\mathbf{\Gamma}\mathbf{\Omega}^{-1}\mathbf{\Gamma}^T + \mathbf{\Gamma}(\boldsymbol{\eta}^T \otimes \boldsymbol{\eta})\mathbf{\Gamma}^T \\ &= (\pi_1\pi_2)^{-1}\mathbf{\Gamma}\mathbf{\Omega}^{-1}\mathbf{\Gamma}^T + (\boldsymbol{\beta}^T \boldsymbol{\Sigma} \boldsymbol{\beta})\mathbf{\Gamma}\mathbf{\Omega}^{-1}\mathbf{\Gamma}^T + \boldsymbol{\beta}^T \otimes \boldsymbol{\beta},\end{aligned}$$

271 where the last equality is because of $\boldsymbol{\beta}^T \boldsymbol{\Sigma} \boldsymbol{\beta} = \boldsymbol{\eta}^T \mathbf{\Gamma}^T (\mathbf{\Gamma}\mathbf{\Omega}\mathbf{\Gamma}^T + \mathbf{\Gamma}_0\mathbf{\Omega}_0\mathbf{\Gamma}_0^T)\mathbf{\Gamma}\boldsymbol{\eta} = \boldsymbol{\eta}^T \mathbf{\Omega} \boldsymbol{\eta}$.

272

The second term in (A23) $\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12}$ can be calculated as,

$$\begin{aligned}
& \mathbf{J}_h \mathbf{H}_{12} \\
= & \pi_1 \pi_2 \begin{pmatrix} \Sigma & (\beta^T \otimes \mathbf{I}_p) \mathbf{E}_p \\ \mathbf{E}_p^T (\beta \otimes \mathbf{I}_p) & \mathbf{E}_p^T \{ (\beta \beta^T + (2\pi_1 \pi_2)^{-1} \Sigma^{-1}) \otimes \Sigma^{-1} \} \mathbf{E}_p \end{pmatrix} \times \\
& \times \begin{pmatrix} \eta^T \otimes \Gamma_0 \\ 2\mathbf{C}_p(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \end{pmatrix} \\
= & \pi_1 \pi_2 \begin{pmatrix} \eta^T \otimes \Gamma_0 \Omega_0 + \eta^T \Omega \otimes \Gamma_0 - \eta^T \otimes \Gamma_0 \Omega_0 \\ \mathbf{E}_p^T (\beta \eta^T \otimes \Gamma_0) + \mathbf{E}_p^T \{ (\beta \beta^T + (2\pi_1 \pi_2)^{-1} \Sigma^{-1}) \otimes \Sigma^{-1} \} 2\mathbf{E}_p \mathbf{C}_p(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \end{pmatrix} \\
= & \pi_1 \pi_2 \begin{pmatrix} \eta^T \Omega \otimes \Gamma_0 \\ \mathbf{E}_p^T (\beta \eta^T \otimes \Gamma_0) + \mathbf{E}_p^T \{ (\beta \beta^T + (2\pi_1 \pi_2)^{-1} \Sigma^{-1}) \otimes \Sigma^{-1} \} 2\mathbf{P}_{\mathbf{E}_p}(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \end{pmatrix},
\end{aligned}$$

273

where the bottom block of the matrix can be simplified as follows. Notice that (Corollary D1 and

274

D2 in Cook et al. (2010)) $2(\mathbf{A} \otimes \mathbf{B})\mathbf{P}_{\mathbf{E}}(\mathbf{C} \otimes \mathbf{D}) = \mathbf{AC} \otimes \mathbf{BD} + (\mathbf{AD} \otimes \mathbf{BC})\mathbf{K}$ where $\mathbf{E}$ and

275

$\mathbf{K}$ are the expansion matrix and the communication matrix with correct dimensions, and also that

276

$\mathbf{E}^T(\mathbf{A} \otimes \mathbf{A})\mathbf{P}_{\mathbf{E}} = \mathbf{E}^T(\mathbf{A} \otimes \mathbf{A})$. We have the following derivations,

$$\begin{aligned}
\mathbf{E}_p^T (\beta \beta^T \otimes \Sigma^{-1}) 2\mathbf{P}_{\mathbf{E}_p}(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) &= \mathbf{E}_p^T (\beta \beta^T \otimes \Sigma^{-1})(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \\
&= \mathbf{E}_p^T (\Gamma \eta \eta^T \Omega \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \eta \eta^T \otimes \Gamma_0) \\
\mathbf{E}_p^T (\Sigma^{-1} \otimes \Sigma^{-1}) 2\mathbf{P}_{\mathbf{E}_p}(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) &= 2\mathbf{E}_p^T (\Sigma^{-1} \otimes \Sigma^{-1})(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \\
&= 2\mathbf{E}_p^T (\Gamma \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \Omega^{-1} \otimes \Gamma_0)
\end{aligned}$$

277

Therefore, we continue the calculation for $\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12}$ as

$$\begin{aligned}
& \mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12} \\
= & \mathbf{H}_{12}^T \pi_1 \pi_2 \begin{pmatrix} \eta^T \Omega \otimes \Gamma_0 \\ \mathbf{E}_p^T (\Gamma \eta \eta^T \otimes \Gamma_0) + \mathbf{E}_p^T (\Gamma \eta \eta^T \Omega \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \eta \eta^T \otimes \Gamma_0) \\ \quad + (2\pi_1 \pi_2)^{-1} 2\mathbf{E}_p^T (\Gamma \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \Omega^{-1} \otimes \Gamma_0) \end{pmatrix} \\
= & \mathbf{H}_{12}^T \begin{pmatrix} \eta^T \Omega \otimes \Gamma_0 \\ \mathbf{E}_p^T (\Gamma \eta \eta^T \Omega \otimes \Gamma_0 \Omega_0^{-1}) + (\pi_1 \pi_2)^{-1} \mathbf{E}_p^T (\Gamma \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \Omega^{-1} \otimes \Gamma_0) \end{pmatrix} \\
= & \pi_1 \pi_2 \begin{pmatrix} \eta^T \otimes \Gamma_0 \\ 2\mathbf{C}_p(\Gamma \Omega \otimes \Gamma_0 - \Gamma \otimes \Gamma_0 \Omega_0) \end{pmatrix}^T \times \\
& \times \begin{pmatrix} \eta^T \Omega \otimes \Gamma_0 \\ \mathbf{E}_p^T (\Gamma \eta \eta^T \Omega \otimes \Gamma_0 \Omega_0^{-1}) + (\pi_1 \pi_2)^{-1} \mathbf{E}_p^T (\Gamma \otimes \Gamma_0 \Omega_0^{-1} - \Gamma \Omega^{-1} \otimes \Gamma_0) \end{pmatrix}.
\end{aligned}$$

278 Direct calculation shows,

$$\begin{aligned}
\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12} &= \pi_1 \pi_2 (\boldsymbol{\eta} \otimes \boldsymbol{\Gamma}_0^T) (\boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Gamma}_0) \\
&+ \pi_1 \pi_2 2 (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) \mathbf{P}_{\mathbf{E}_p} \times \\
&\times \{ (\boldsymbol{\Gamma} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1}) + (\pi_1 \pi_2)^{-1} (\boldsymbol{\Gamma} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0) \} \\
&= \pi_1 \pi_2 (\boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \mathbf{I}_{p-u}) + \pi_1 \pi_2 2 (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) \mathbf{P}_{\mathbf{E}_p} (\boldsymbol{\Gamma} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1}) \\
&+ 2 (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) \mathbf{P}_{\mathbf{E}_p} (\boldsymbol{\Gamma} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0) \\
&= \pi_1 \pi_2 (\boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \mathbf{I}_{p-u}) + \pi_1 \pi_2 (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) (\boldsymbol{\Gamma} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1}) \\
&+ (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) (\boldsymbol{\Gamma} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0) \\
&= \pi_1 \pi_2 (\boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \mathbf{I}_{p-u}) + \pi_1 \pi_2 (\boldsymbol{\Omega} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \mathbf{I}_{p-u}) \\
&+ (\boldsymbol{\Omega} \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma}_0^T - \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Omega}_0 \boldsymbol{\Gamma}_0^T) (\boldsymbol{\Gamma} \otimes \boldsymbol{\Gamma}_0 \boldsymbol{\Omega}_0^{-1} - \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Gamma}_0) \\
&= \pi_1 \pi_2 (\boldsymbol{\Omega} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1}) + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - \mathbf{I}_u \otimes \mathbf{I}_{p-u} - \mathbf{I}_u \otimes \mathbf{I}_{p-u} + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0 \\
&= \pi_1 \pi_2 (\boldsymbol{\Omega} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1}) + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - 2(\mathbf{I}_u \otimes \mathbf{I}_{p-u}) + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0.
\end{aligned}$$

279 Finally, we have the asymptotic covariance matrix of $\hat{\boldsymbol{\beta}}$ as

$$\begin{aligned}
\mathbf{V}_\beta &= \boldsymbol{\Gamma} \{ (\mathbf{H}_1^T \mathbf{J}_h \mathbf{H}_1)^{-1} \}_{11} \boldsymbol{\Gamma}^T + (\boldsymbol{\eta}^T \otimes \mathbf{I}_p) (\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12})^{-1} (\boldsymbol{\eta} \otimes \mathbf{I}_p), \\
&= (\pi_1 \pi_2)^{-1} \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + (\boldsymbol{\beta}^T \boldsymbol{\Sigma} \boldsymbol{\beta}) \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\beta}^T \otimes \boldsymbol{\beta} \\
&+ (\boldsymbol{\eta}^T \otimes \mathbf{I}_p) (\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12})^\dagger (\boldsymbol{\eta} \otimes \mathbf{I}_p),
\end{aligned}$$

280 where $\mathbf{H}_{12}^T \mathbf{J}_h \mathbf{H}_{12} = \pi_1 \pi_2 (\boldsymbol{\Omega} \boldsymbol{\eta} \boldsymbol{\eta}^T \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1}) + \boldsymbol{\Omega} \otimes \boldsymbol{\Omega}_0^{-1} - 2(\mathbf{I}_u \otimes \mathbf{I}_{p-u}) + \boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}_0$.

281 Finally, we identify the first part is indeed $\text{avar}(\sqrt{n} \hat{\boldsymbol{\beta}}_\Gamma) = (\pi_1 \pi_2)^{-1} \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + (\boldsymbol{\beta}^T \boldsymbol{\Sigma} \boldsymbol{\beta}) \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T +$
282 $\boldsymbol{\beta}^T \otimes \boldsymbol{\beta}$. Since $\text{avar}(\sqrt{n} \hat{\boldsymbol{\beta}}_\Gamma) = \text{avar}(\sqrt{n} \boldsymbol{\Gamma} \hat{\boldsymbol{\eta}}_\Gamma) = \boldsymbol{\Gamma} \text{avar}(\sqrt{n} \hat{\boldsymbol{\eta}}_\Gamma) \boldsymbol{\Gamma}^T$, we need to find the asymp-
283 totic distribution of $\hat{\boldsymbol{\eta}}_\Gamma$, which is the LDA direction of the reduced variables $\boldsymbol{\Gamma}^T \mathbf{X}$. Notice that,
284 given $\boldsymbol{\Gamma}$, we can let $\mathbf{Z} = \boldsymbol{\Gamma}^T \mathbf{X} \in \mathbb{R}^u$ and have the LDA model of $\mathbf{Z} \mid (Y = k) \sim N(\boldsymbol{\nu}_k, \boldsymbol{\Omega})$,
285 where the discriminant directions are exactly $\boldsymbol{\eta}_k = \boldsymbol{\Omega}^{-1}(\boldsymbol{\nu}_k - \boldsymbol{\nu}_1)$. Hence, the asymptotic prop-
286 erty of $\hat{\boldsymbol{\eta}}_\Gamma$ can be easily obtained by replacing $\boldsymbol{\Sigma} \rightarrow \boldsymbol{\Omega}$, $\boldsymbol{\beta} \rightarrow \boldsymbol{\eta}$ in the asymptotic distribu-
287 tion of $\tilde{\boldsymbol{\beta}}$. Specifically, $\text{avar}(\sqrt{n} \hat{\boldsymbol{\eta}}_\Gamma) = (\pi_1 \pi_2)^{-1} \boldsymbol{\Omega}^{-1} + (\boldsymbol{\eta}^T \boldsymbol{\Omega} \boldsymbol{\eta}) \boldsymbol{\Omega}^{-1} + \boldsymbol{\eta}^T \otimes \boldsymbol{\eta}$. Therefore,
288 $\text{avar}(\sqrt{n} \hat{\boldsymbol{\beta}}_\Gamma) = \boldsymbol{\Gamma} \text{avar}(\sqrt{n} \hat{\boldsymbol{\eta}}_\Gamma) \boldsymbol{\Gamma}^T = (\pi_1 \pi_2)^{-1} \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + (\boldsymbol{\eta}^T \boldsymbol{\Omega} \boldsymbol{\eta}) \boldsymbol{\Gamma} \boldsymbol{\Omega}^{-1} \boldsymbol{\Gamma}^T + \boldsymbol{\eta}^T \boldsymbol{\Gamma}^T \otimes \boldsymbol{\Gamma} \boldsymbol{\eta}$. The
289 conclusion follows by noticing that $\boldsymbol{\beta} = \boldsymbol{\Gamma} \boldsymbol{\eta}$ and $\boldsymbol{\beta}^T \boldsymbol{\Sigma} \boldsymbol{\beta} = \boldsymbol{\eta}^T \boldsymbol{\Omega} \boldsymbol{\eta}$.

S10 Proof for Propositions 6, 7 and 8

First, for LDA model, the $\sqrt{n}$ -consistency of the envelope estimator from maximizing the likelihood-based objective function $L_n(\{\boldsymbol{\mu}_k\}_{k=1}^K, \boldsymbol{\Sigma})$ relies on Shapiro's (1986; Proposition 4.1) results on the asymptotics of over-parameterized structural models. The proof is parallel to the proof of Proposition 4 in Cook and Zhang (2015) for simultaneous envelope and is thus omitted. Similarly, for QDA model, the likelihood-based objective function $L_n(\{\boldsymbol{\mu}_k\}_{k=1}^K, \{\boldsymbol{\Sigma}_k\}_{k=1}^K)$ satisfies conditions to apply Shapiro's results. Thus we have the similar results. Under normality, the estimators becomes MLE, and thus $\text{avar}(\sqrt{n}\tilde{\mathbf{h}}) = \boldsymbol{\Xi} = \mathbf{J}_{\mathbf{h}}^{-1}$ and $\text{avar}(\sqrt{n}\hat{\mathbf{h}}) = \mathbf{H}(\mathbf{H}^T \mathbf{J}_{\mathbf{h}} \mathbf{H})^+ \mathbf{H}^T$ is always less than or equals to $\text{avar}(\sqrt{n}\tilde{\mathbf{h}})$ because any subspace, including $\text{span}(\mathbf{J}_{\mathbf{h}}^{1/2} \mathbf{H})$, is a reducing subspace of $\mathbf{J}_{\mathbf{h}}^{1/2} \boldsymbol{\Xi} \mathbf{J}_{\mathbf{h}}^{1/2} = \mathbf{I}_p$.

S11 Implementation Details

Recall that the ENDS objective function is

$$F(\boldsymbol{\Gamma}) = \log |\boldsymbol{\Gamma}^T \mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Gamma}| + \sum_{k=1}^K \frac{n_k}{n} \log |\boldsymbol{\Gamma}^T \mathbf{S}_k(\lambda) \boldsymbol{\Gamma}|, \quad (\text{A24})$$

which is to be minimized subject to constraint $\boldsymbol{\Gamma}^T \boldsymbol{\Gamma} = \mathbf{I}_u$, and includes the ENDS-LDA and ENDS-QDA objective functions as its special cases. We hence describe the implementation of algorithms for solving this optimization problem of $F(\boldsymbol{\Gamma})$. The derivative can be calculated straightforwardly as,

$$\frac{dF(\boldsymbol{\Gamma})}{d\boldsymbol{\Gamma}} = 2\mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Gamma} (\boldsymbol{\Gamma}^T \mathbf{S}_{\mathbf{X}}^{-1} \boldsymbol{\Gamma})^{-1} + 2 \sum_{k=1}^K \frac{n_k}{n} \mathbf{S}_k(\lambda) \boldsymbol{\Gamma} \{\boldsymbol{\Gamma}^T \mathbf{S}_k(\lambda) \boldsymbol{\Gamma}\}^{-1}. \quad (\text{A25})$$

Then based on the explicit forms of the objective function and its derivatives, we can use any off-the-shelf methods to find the optimizer $\hat{\boldsymbol{\Gamma}}$. To preserve the geometric constraint, $\boldsymbol{\Gamma}^T \boldsymbol{\Gamma} = \mathbf{I}_u$, we choose to implement the estimation using the *sg_min* package in Matlab that is a generic computing package for Stiefel and Grassmann manifolds optimization (see Edelman et al. (1998) for more backgrounds). Other more recent methods for optimization with orthogonality constraints (for example, Wen and Yin (2013)) can also be straightforwardly incorporated into our implementation.

When u is bigger than one, the optimization can be expensive and requires good initial value to avoid local minima. Therefore, we adopt the sequential 1D algorithm (Cook and Zhang, 2016) to

Algorithm 1 ENDS optimization based on the 1D algorithm from Cook and Zhang (2016).

Let $\mathbf{g}_j \in \mathbb{R}^p$, $j = 1, \dots, u$, be the sequential directions obtained. Let $\mathbf{G}_j = (\mathbf{g}_1, \dots, \mathbf{g}_j)$, let $(\mathbf{G}_j, \mathbf{G}_{0j})$ be an orthogonal basis for $\mathbb{R}^p$ and set initial value $\mathbf{g}_0 = \mathbf{G}_0 = \mathbf{0}$.

For $j = 0, \dots, u - 1$, repeat Step 1 and 2 in the following.

1. Define $\mathbf{A}_j = \mathbf{G}_{0j}^T \mathbf{S}_{\mathbf{X}} \mathbf{G}_{0j}$, $\mathbf{B}_{jk} = \mathbf{G}_{0j}^T \mathbf{S}_k(\lambda) \mathbf{G}_{0j}$, and the objective function for $\mathbf{w} \in \mathbb{R}^{p-j+1}$

$$\phi_j(\mathbf{w}) = \log |\mathbf{w}^T \mathbf{A}_j^{-1} \mathbf{w}| + \sum_{k=1}^K \frac{n_k}{n} \log |\mathbf{w}^T \mathbf{B}_{jk} \mathbf{w}|. \quad (\text{A26})$$

2. Solve $\mathbf{w}_{j+1} = \arg \min \phi_j(\mathbf{w})$ subject to $\mathbf{w}^T \mathbf{w} = 1$, then the $(j + 1)$ -th envelope direction is $\mathbf{g}_{j+1} = \mathbf{G}_{0j} \mathbf{w}_{j+1}$.

The resulting ENDS basis estimator is $\hat{\mathbf{\Gamma}} = \mathbf{G}_u$.

speed up the computation. The 1D algorithm and its variations (Cook and Zhang, 2017) provides a fast and reliable way to obtain envelope subspaces one direction at a time. It breaks down the full optimization of $F(\mathbf{\Gamma})$ into u sequential optimizations as summarized in Algorithm 1.

References

- Cook, R. D. and Forzani, L. (2009). Likelihood-based sufficient dimension reduction. *Journal of the American Statistical Association*, 104(485):197–208. S2, S6
- Cook, R. D. and Zhang, X. (2016). Algorithms for envelope estimation. *Journal of Computational and Graphical Statistics*, 25(1):284–300. S11, 1
- Cook, R. D. and Zhang, X. (2017). Fast envelope algorithms. *Statistica Sinica (just-accepted)*. S11
- Edelman, A., Arias, T. A., and Smith, S. T. (1998). The geometry of algorithms with orthogonality constraints. *SIAM journal on Matrix Analysis and Applications*, 20(2):303–353. S11
- Wen, Z. and Yin, W. (2013). A feasible method for optimization with orthogonality constraints. *Mathematical Programming*, 142(1-2):397–434. S11