

Tackling Small Eigen-Gaps: Fine-Grained Eigenvector Estimation and Inference Under Heteroscedastic Noise

Chen Cheng, Yuting Wei^{ID}, and Yuxin Chen^{ID}, *Member, IEEE*

Abstract—This paper aims to address two fundamental challenges arising in eigenvector estimation and inference for a low-rank matrix from noisy observations: 1) how to estimate an unknown eigenvector when the eigen-gap (i.e. the spacing between the associated eigenvalue and the rest of the spectrum) is particularly small; 2) how to perform estimation and inference on linear functionals of an eigenvector—a sort of “fine-grained” statistical reasoning that goes far beyond the usual ℓ_2 analysis. We investigate how to address these challenges in a setting where the unknown $n \times n$ matrix is symmetric and the additive noise matrix contains independent (and non-symmetric) entries. Based on eigen-decomposition of the asymmetric data matrix, we propose estimation and uncertainty quantification procedures for an unknown eigenvector, which further allow us to reason about linear functionals of an unknown eigenvector. The proposed procedures and the accompanying theory enjoy several important features: 1) distribution-free (i.e. prior knowledge about the noise distributions is not needed); 2) adaptive to heteroscedastic noise; 3) minimax optimal under Gaussian noise. Along the way, we establish valid procedures to construct confidence intervals for the unknown eigenvalues. All this is guaranteed even in the presence of a small eigen-gap (up to $O(\sqrt{n}/\text{poly} \log(n))$ times smaller than the requirement in prior theory), which goes significantly beyond what generic matrix perturbation theory has to offer.

Index Terms—Eigen-gap, linear form of eigenvectors, confidence interval, uncertainty quantification, heteroscedasticity.

I. INTRODUCTION

A VARIETY of science and engineering applications ask for spectral analysis of low-rank matrices in high

Manuscript received November 22, 2020; revised June 13, 2021; accepted September 7, 2021. Date of publication September 10, 2021; date of current version October 20, 2021. The work of Chen Cheng was supported by the William R. Hewlett Stanford Graduate Fellowship. The work of Yuting Wei was supported in part by NSF under Grant DMS-2015447/2147546 and Grant CCF-2007911. The work of Yuxin Chen was supported in part by Air Force Office of Scientific Research (AFOSR) Young Investigator Program (YIP) under Award FA9550-19-1-0030; in part by Office of Naval Research (ONR) under Grant N00014-19-1-2120; in part by Army Research Office (ARO) under Grant W911NF-20-1-0097 and Grant W911NF-18-1-0303; and in part by NSF under Grant CCF-1907661, Grant IIS-1900140, Grant IIS-2100158, and Grant DMS-2014279. (Corresponding author: Yuting Wei.)

Chen Cheng is with the Department of Statistics, Stanford University, Stanford, CA 94305 USA (e-mail: chencheng@stanford.edu).

Yuting Wei is with the Department of Statistics and Data Science, The Wharton School, University of Pennsylvania, Philadelphia, PA 19104 USA (e-mail: ytwei@wharton.upenn.edu).

Yuxin Chen is with the Department of Electrical and Computer Engineering, Princeton University, Princeton, NJ 08544 USA (e-mail: yuxin.chen@princeton.edu).

Communicated by M. Davenport, Associate Editor for Signal Processing.

Color versions of one or more figures in this article are available at <https://doi.org/10.1109/TIT.2021.3111828>.

Digital Object Identifier 10.1109/TIT.2021.3111828

dimension [1]. Setting the stage, imagine that we are interested in a large low-rank matrix

$$M^* = \sum_{l=1}^r \lambda_l^* \mathbf{u}_l^* \mathbf{u}_l^{*\top} \in \mathbb{R}^{n \times n}, \quad (1)$$

where r ($r \ll n$) represents the rank, and λ_l^* stands for the l th largest eigenvalue of M^* , with $\mathbf{u}_l^* \in \mathbb{R}^n$ the associated eigenvector. Suppose, however, that we do not have access to perfect measurements about the entries of this matrix; rather, the observations we have available, represented by a data matrix $M = M^* + H$, are contaminated by a substantial amount of random noise (reflected by the noise matrix H). The aim is to perform reliable estimation and inference on the unseen eigenvectors of M^* on the basis of noisy data. Motivated by the abundance of applications (e.g. collaborative filtering, harmonic retrieval, sensor network localization, joint shape matching [2]–[6]), research on eigenspace estimation in this context has flourished in the past several years, typically built upon proper exploitation of low-rank structures. We have now been equipped with a rich suite of modern statistical theory that delivers statistical performance guarantees for a number of spectral estimators (e.g. [7]–[16], [16]–[22]); see [23] for a contemporary overview of spectral methods.

A. Motivation and Challenges

Despite a large body of work tackling the above problem, there are several fundamental yet unaddressed challenges that deserve further attention.

1) *Stringent Requirements on Eigen-Gaps*: A crucial identifiability issue stands out when estimating individual eigenvectors. In general, one cannot possibly disentangle the eigenvectors $\mathbf{u}_1^*, \dots, \mathbf{u}_r^*$ unless there is sufficient spacing between adjacent eigenvalues. After all, even in the noiseless setting, one can only hope to recover the subspace spanned by $\{\mathbf{u}_i^*\}$, rather than individual eigenvectors, if all non-zero eigenvalues are identical.

In principle, a “sufficient” eigen-gap criterion in the presence of noise is dictated by the noise levels, or more precisely, the signal-to-noise ratios (SNRs). However, generic linear algebra theory typically imposes fairly stringent, and hence pessimistic, eigen-gap requirements for both eigenvector and eigenvalue estimation. More concretely, imagine we wish to estimate the l th eigenvector $\mathbf{u}_l^*$: generic matrix perturbation theory (e.g. the Davis-Kahan $\sin \Theta$ theorem or the Wedin theorem [12], [24], [25]) typically cannot guarantee meaningful

estimation of $\mathbf{u}_l^*$ unless

$$(\text{classical theory}) \quad \underbrace{\min_{k:k \neq l} |\lambda_l^* - \lambda_k^*|}_{\text{eigen-gap}} > \|\mathbf{H}\|. \quad (2)$$

This eigen-gap requirement can be problematic and hard to satisfy as the noise size grows, casting doubts on our ability to perform informative inference on individual eigenvectors.

2) *Fine-Grained Estimation and Inference (Beyond ℓ_2 Guarantees)*: In many applications, it is often the case that the ultimate goal is not to characterize the ℓ_2 or “bulk” behavior (e.g. the mean squared estimation error) of an eigenvector estimator, but rather to reason about the eigenvectors along a few preconceived yet important directions. Take community recovery for instance: the eigenvector of a certain adjacency matrix encodes community membership of a set of users [26]; if we wish to infer the community memberships of a few important users, or the similarities between a few pairs of users, it boils down to assessing the entrywise behavior or certain pairwise linear functional of an eigenvector estimator. Another example is concerned with harmonic retrieval: the leading eigenvector of a properly arranged Toeplitz matrix represents the time-domain response of the sinusoidal signals of interest [27]; thus, retrieving the underlying frequency involves inferring the Fourier coefficients of this eigenvector at some given frequencies. These problems can be formulated as estimation and inference for *linear functionals of an eigenvector*, namely, quantities of the form $\mathbf{a}^\top \mathbf{u}_l^*$ ($1 \leq l \leq r$) with $\mathbf{a}$ some prescribed vector. In principle, this task can be viewed as a sort of *fine-grained* statistical analysis, given that it pursues highly “local” and “delicate” information of interest.

Towards estimating $\mathbf{a}^\top \mathbf{u}_l^*$, a natural starting point is a “plug-in” estimator, which computes a reasonable estimate $\hat{\mathbf{u}}_l$ of $\mathbf{u}_l^*$ and outputs $\mathbf{a}^\top \hat{\mathbf{u}}_l$. There are several critical challenges that we shall bear in mind. To begin with, a dominant fraction of prior theory focuses on ℓ_2 risk analysis of an eigenvector estimator, which is often too coarse to deliver tight uncertainty assessment for the plug-in estimator. In fact, the ℓ_2 risk bounds alone often lead to highly conservative estimates for the risk under consideration, making it hard to assess the performance of the plug-in estimator. To further complicate matters, there is often a severe bias issue surrounding the plug-in estimator. Even when an estimator $\hat{\mathbf{u}}_l$ is nearly unbiased in a strong entrywise sense (meaning that the estimation bias is dominated by the variability in every single coordinate), this property alone does not preclude the possibility of bias accumulation along the direction $\mathbf{a}$. Addressing these issues calls for refined risk analysis as well as careful examination into the bias-variance trade-off.

B. A Glimpse of Our Approach and Our Contributions

1) *Eigen-Decomposition Meets Statistical Asymmetry*: The current paper makes progress in a setting where the noise matrix $\mathbf{H}$ consists of independent (but possibly heterogeneous and heteroscedastic) zero-mean components. Our approach is inspired by the findings of [17]. Consider, for example, the case when $\mathbf{H}$ is a random and asymmetric matrix and when $\mathbf{M}^*$ is a rank-1 symmetric matrix. The results in [17]

reveal that an eigen-decomposition approach applied to $\mathbf{M}$ (without symmetrization) achieves appealing statistical accuracy when estimating the leading eigenvalue of $\mathbf{M}^*$. The key enabler is an implicit bias reduction feature when computing vanilla eigen-decomposition of an asymmetric data matrix. While [17] only provides highly partial results when going beyond the rank-1 case, it hints at the potential benefits of eigen-decomposition in super-resolving the spectrum.

2) *Our Contributions*: The main contributions of this paper are summarized below, all of which are built upon an eigen-decomposition approach applied to the asymmetric data matrix $\mathbf{M}$ (without symmetrization). As an important advantage, the proposed procedures and their accompanying theory are distribution-free, meaning that they do not require prior knowledge about the distribution of the noise and thus are fully adaptive to heteroscedasticity of data.

- We demonstrate that the l th eigenvector $\mathbf{u}_l^*$ and eigenvalue λ_l^* can be estimated with near-optimal accuracy even when the eigen-gap (2) is extremely small. More concretely, for various noise distributions our results only require

$$(\text{our theory}) \quad \min_{k:k \neq l} |\lambda_l^* - \lambda_k^*| \gtrsim \|\mathbf{H}\| \frac{\text{poly log}(n)}{\sqrt{n}}, \quad (3)$$

which is about $O(\sqrt{n})$ times less stringent (up to log factor) than generic linear algebra theory (cf. (2)).

- We propose a new estimator for the linear functional $\mathbf{a}^\top \mathbf{u}_l^*$ (obtained via proper de-biasing of certain plug-in estimators) that achieves minimax-optimal statistical accuracy. In addition, we demonstrate how to construct valid confidence intervals for $\mathbf{a}^\top \mathbf{u}_l^*$.
- Additionally, we demonstrate how to compute valid confidence intervals for the eigenvalues of interest. Interestingly, de-biasing is not needed at all for performing inference on eigenvalues, as the eigen-decomposition approach implicitly alleviates the estimation bias.

Our findings unveil new insights into the capability of spectral methods in a statistical context, going far beyond what conventional matrix perturbation theory has to offer.

C. Notation

For any vector $\mathbf{x}$, denote by $\|\mathbf{x}\|_2$ and $\|\mathbf{x}\|_\infty$ the ℓ_2 norm and the ℓ_∞ norm of $\mathbf{x}$, respectively. For any matrix $\mathbf{M}$, we let $\|\mathbf{M}\|$, $\|\mathbf{M}\|_F$, $\|\mathbf{M}\|_\infty$ and $\|\mathbf{M}\|_{2,\infty}$ represent the spectral norm, the Frobenius norm, the entrywise ℓ_∞ norm (i.e. $\|\mathbf{M}\|_\infty := \max_{i,j} |M_{ij}|$), and the two-to-infinity norm (i.e. $\|\mathbf{M}\|_{2,\infty} := \sup_{\|\mathbf{x}\|_2=1} \|\mathbf{M}\mathbf{x}\|_\infty$) of $\mathbf{M}$, respectively. The notation $f(n) = O(g(n))$ or $f(n) \lesssim g(n)$ means that there is a universal constant $c > 0$ such that $|f(n)| \leq c|g(n)|$, $f(n) \gtrsim g(n)$ means that there is a universal constant $c > 0$ such that $|f(n)| \geq c|g(n)|$, and $f(n) \asymp g(n)$ means that there exist constants $c_1, c_2 > 0$ such that $c_1|g(n)| \leq |f(n)| \leq c_2|g(n)|$. The notation $f(n) \gg g(n)$ (resp. $f(n) \ll g(n)$) means that there exists a sufficiently large (resp. small) constant $c > 0$ such that $f(n) \geq cg(n)$ (resp. $f(n) \leq cg(n)$).

In addition, we denote by $u_{l,j}$, $w_{l,j}$ and $u_{l,j}^*$ the j th entry of $\mathbf{u}_l$, $\mathbf{w}_l$ and $\mathbf{u}_l^*$, respectively. Let $\mathbf{e}_1, \dots, \mathbf{e}_n$ represent the standard basis vectors in $\mathbb{R}^n$, $\mathbf{1}_n \in \mathbb{R}^n$ the all-one vector, and $\mathbf{I}_n \in \mathbb{R}^{n \times n}$ the identity matrix. We shall also abbreviate the interval $[b - c, b + c]$ to $[b \pm c]$, and denote $\min |b \pm c| = \min\{|b + c|, |b - c|\}$, $\max |b \pm c| = \max\{|b + c|, |b - c|\}$, and $\min \|\mathbf{b} \pm \mathbf{c}\| = \min\{\|\mathbf{b} - \mathbf{c}\|, \|\mathbf{b} + \mathbf{c}\|\}$ for any norm $\|\cdot\|$. Moreover, denote by $\Phi(\cdot)$ the cumulative distribution function (CDF) of a standard Gaussian random variable.

II. PROBLEM FORMULATION

A. Model

Imagine we seek to estimate an unknown rank- r matrix

$$\mathbf{M}^* = \sum_{l=1}^r \lambda_l^* \mathbf{u}_l^* \mathbf{u}_l^{*\top} =: \mathbf{U}^* \mathbf{\Sigma}^* \mathbf{U}^{*\top} \in \mathbb{R}^{n \times n}, \quad (4)$$

where $\lambda_1^* \geq \lambda_2^* \geq \dots \geq \lambda_r^*$ denote the r nonzero eigenvalues of $\mathbf{M}^*$, and $\mathbf{u}_1^*, \dots, \mathbf{u}_r^*$ represent the associated (normalized) eigenvectors. Here, for notational convenience we let

$$\begin{aligned} \mathbf{U}^* &:= [\mathbf{u}_1^*, \dots, \mathbf{u}_r^*] \in \mathbb{R}^{n \times r}, \\ \mathbf{\Sigma}^* &:= \begin{bmatrix} \lambda_1^* & & \\ & \ddots & \\ & & \lambda_r^* \end{bmatrix} \in \mathbb{R}^{r \times r}. \end{aligned} \quad (5)$$

We shall also define

$$\lambda_{\max}^* := \max_{1 \leq l \leq r} |\lambda_l^*|, \quad \lambda_{\min}^* := \min_{1 \leq l \leq r} |\lambda_l^*| \quad \text{and} \quad \kappa := \lambda_{\max}^* / \lambda_{\min}^*. \quad (6)$$

What we have observed is a corrupted copy of $\mathbf{M}^*$, namely,

$$\mathbf{M} = \mathbf{M}^* + \mathbf{H}, \quad (7)$$

where $\mathbf{H} \in \mathbb{R}^{n \times n}$ stands for a random noise matrix. This paper focuses on the family of noise matrices satisfying the following assumptions:

Assumption 1: The noise matrix $\mathbf{H} \in \mathbb{R}^{n \times n}$ satisfies the following properties.

- 1) **Independence and zero mean.** The entries $\{H_{ij}\}_{1 \leq i, j \leq n}$ are independent zero-mean random variables; this indicates that the matrices $\mathbf{H}$ and $\mathbf{M}$ are, in general, *asymmetric*.
- 2) **Heteroscedasticity and unknown variances.** Let $\sigma_{ij}^2 := \text{Var}(H_{ij})$ denote the variance of H_{ij} . To accommodate more realistic scenarios, we allow σ_{ij}^2 to vary across entries — commonly referred to as *heteroscedastic noise*. In addition, we do not a priori know $\{\sigma_{ij}^2\}_{1 \leq i, j \leq n}$. Throughout this paper, we assume

$$0 \leq \sigma_{\min}^2 \leq \sigma_{ij}^2 \leq \sigma_{\max}^2, \quad 1 \leq i, j \leq n. \quad (8)$$

- 3) **Magnitudes.** Each H_{ij} ($1 \leq i, j \leq n$) satisfies either of the following conditions:

- (a) $|H_{ij}| \leq B$;
- (b) H_{ij} has a symmetric distribution obeying $\mathbb{P}\{|H_{ij}| > B\} \leq c_b n^{-12}$ for some constant $c_b > 0$, and $\mathbb{E}[H_{ij}^2 \mathbf{1}_{\{|H_{ij}| > B\}}] = o(\sigma_{ij}^2)$.

Remark 1: Here, the quantities B , $\{\sigma_{i,j}\}$, $\sigma_{\min}$ and $\sigma_{\max}$ may all depend on n .

Remark 2: As we shall see momentarily, the lower bound $\sigma_{\min}^2$ on the noise variance is imposed only for our statistical inference theory (or more precisely, it is imposed in order to ensure the plausibility to estimate the variance of the estimation error in an accurate manner). This lower bound $\sigma_{\min}^2$ is not needed at all in our eigenvector estimation guarantees (e.g., Theorems 1 and 2).

The careful reader would naturally ask when we would have an asymmetric noise matrix $\mathbf{H}$. This might happen when, for example, one has collected two independent samples about each entry of $\mathbf{M}^*$ and chooses to arrange the observed data in an asymmetric manner. Moreover, when the noise matrix is a symmetric Gaussian random matrix and possibly contains missing data, [17, Appendix J] points out some simple asymmetrization tricks that allow one to convert a symmetric data matrix $\mathbf{M}$ to an asymmetric counterpart with independent components; see Appendix G for an example.

In addition to the above assumptions on the noise, we assume that the unknown matrix $\mathbf{M}^*$ satisfies a certain incoherence condition, as commonly seen in the low-rank matrix estimation literature.

Definition 1 (Incoherence): The matrix $\mathbf{M}^*$ with eigen-decomposition $\mathbf{M}^* = \mathbf{U}^* \mathbf{\Sigma}^* \mathbf{U}^{*\top}$ is said to be μ -incoherent if $\|\mathbf{U}^*\|_\infty \leq \sqrt{\mu/n}$.

B. Goal

The primary goal of the current paper is to perform certain “fine-grained” statistical estimation and inference on the unknown eigenvectors $\{\mathbf{u}_1^*, \dots, \mathbf{u}_r^*\}$. To be more specific, we aim at developing statistically efficient methods to estimate and construct confidence intervals for linear functionals taking the form $\mathbf{a}^\top \mathbf{u}_l^*$ ($1 \leq l \leq r$), where $\mathbf{a} \in \mathbb{R}^n$ is a pre-determined fixed vector. Along the way, we shall also demonstrate how to perform estimation and inference on the unknown eigenvalues $\{\lambda_1^*, \dots, \lambda_r^*\}$. Ideally, all these tasks can be accomplished even when the associated eigen-gaps are very small, without requiring prior knowledge about the noise distributions and noise levels.

III. ESTIMATION

This section presents our algorithms and the accompanying theory for estimating an eigenvector $\mathbf{u}_l^*$, a linear functional $\mathbf{a}^\top \mathbf{u}_l^*$ of this eigenvector, as well as the associated eigenvalue λ_l^* .

Two popular estimation schemes immediately come into mind: (1) performing eigen-decomposition after symmetrizing the data matrix (e.g. replacing $\mathbf{M}$ by $\frac{1}{2}(\mathbf{M} + \mathbf{M}^\top)$); (2) computing singular value decomposition (SVD) of the asymmetric data matrix $\mathbf{M}$. Contrastingly, this paper adopts a far less widely used, and in fact far less widely studied, strategy based on eigen-decomposition without symmetrization; namely, we attempt estimation of $\{\mathbf{u}_l^*\}$ and $\{\lambda_l^*\}$ via the eigenvectors and the eigenvalues of $\mathbf{M}$, respectively, despite the asymmetric nature of $\mathbf{M}$ in general. While conventional wisdom does not advocate eigen-decomposition of asymmetric

matrices (due to, say, potential numerical instability), our investigation uncovers remarkable advantages of this approach under the statistical context considered in the present paper.

A. Notation: Eigen-Decomposition Without Symmetrization

Before continuing, we introduce several additional notation that shall be adopted throughout. Owing to the asymmetry of M , the left and the right eigenvectors of M do not coincide.

Notation 1: Let $\lambda_1, \dots, \lambda_r$ denote the top- r leading eigenvalues of M (so that $\min_{1 \leq l \leq r} |\lambda_l|$ is larger in magnitude than any other eigenvalue of M). Assume that they are sorted by their real parts, namely, $\text{Re}(\lambda_1) \geq \text{Re}(\lambda_2) \geq \dots \geq \text{Re}(\lambda_r)$, and denote by $\mathbf{u}_l$ (resp. $\mathbf{w}_l$) the right (resp. left) eigenvector of M associated with λ_l ; this means

$$M\mathbf{u}_l = \lambda_l \mathbf{u}_l \quad \text{and} \quad M^\top \mathbf{w}_l = \lambda_l \mathbf{w}_l. \quad (9)$$

In addition, if $\mathbf{u}_l$ and $\mathbf{w}_l$ are both real-valued, then we assume without loss of generality that $\langle \mathbf{u}_l, \mathbf{w}_l \rangle \geq 0$.

Remark 3: As we shall justify shortly in Theorem 1, even though M is in general asymmetric, the eigenvalue λ_l and the eigenvectors $\mathbf{u}_l$ and $\mathbf{w}_l$ are, with high probability, real-valued under the assumptions imposed in this paper.

B. Estimation Algorithms

We are now ready to present our procedures for estimating the l th eigenvector of M^* , on the basis of eigen-decomposition of M (see Notation 1).

- **Estimator for $\mathbf{u}_l^*$ (with the aim of achieving low ℓ_2 risk):**

$$\hat{\mathbf{u}}_l := \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l); \quad (10)$$

- **Estimator for $\mathbf{a}^\top \mathbf{u}_l^*$ for a preconceived direction $\mathbf{a}$:**

$$\hat{u}_{\mathbf{a},l} := \min \left\{ \sqrt{\left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{w}_l^\top \mathbf{u}_l} \right|}, \|\mathbf{a}\|_2 \right\}. \quad (11)$$

Here, the rationale behind $\hat{u}_{\mathbf{a},l}$ is this: both $\mathbf{a}^\top \mathbf{u}_l$ and $\mathbf{a}^\top \mathbf{w}_l$ might systematically *under-estimate* the quantity of interest $\mathbf{a}^\top \mathbf{u}_l^*$. As a result, one is advised to first alleviate the bias via proper de-shrinking, which is the role played by the factor $\frac{1}{\mathbf{w}_l^\top \mathbf{u}_l}$. The choice of this factor is based on in-depth understanding of the behavior of $\mathbf{a}^\top \mathbf{u}_l$ and $\mathbf{a}^\top \mathbf{w}_l$, and will be better understood after we delve into technical details. As an important feature, the proposed estimation procedures do not rely on prior knowledge about the noise distributions.

Remark 4: The estimator (11) involves the term $\|\mathbf{a}\|_2$ due to the trivial upper bound $|\mathbf{a}^\top \mathbf{u}_l^*| \leq \|\mathbf{a}\|_2 \|\mathbf{u}_l^*\|_2 = \|\mathbf{a}\|_2$.

C. Theoretical Guarantees

We now embark on theoretical development for the proposed estimators. As alluded to previously, whether we can reliably estimate and infer an eigenvector $\mathbf{u}_l^*$ depends largely upon the spacing between the l th eigenvalue and its adjacent

eigenvalues. In light of this, we define formally the eigen-gap metric w.r.t. the l th eigenvalue of M^* as follows

$$\Delta_l^* := \begin{cases} \min_{1 \leq k \leq r, k \neq l} |\lambda_l^* - \lambda_k^*|, & \text{if } r > 1; \\ \infty, & \text{otherwise.} \end{cases} \quad (12)$$

Most of our theoretical guarantees rely on this crucial metric. Moreover, our theory in this section is built upon a set of assumptions on the noise levels:

Assumption 2: Suppose that the noise parameters defined in Assumption 1 satisfy

$$\Delta_l^* > 2c_4 \kappa^2 r^2 \sigma_{\max} \sqrt{\mu \log n} \quad (13a)$$

$$B \log n \leq \sigma_{\max} \sqrt{n \log n} \leq c_5 \lambda_{\min}^* / \kappa^3 \quad (13b)$$

for some sufficiently large (resp. small) universal constant $c_4 > 0$ (resp. $c_5 > 0$).

Remark 5: Consider, for example, the case where $r, \kappa, \mu \asymp 1$: the lower bound on Δ_l^* in (13a) is $O(\sqrt{n})$ times smaller than the lower bound on $\lambda_{\min}^*$ in (13b), meaning that the eigen-gap can be considerably smaller than the minimum eigenvalue of the truth.

The following theorem delivers statistical guarantees for the proposed eigenvector estimators, with the proof postponed to Appendix B. We recall the notation $\min \|\mathbf{z} \pm \mathbf{u}_l^*\| := \min\{\|\mathbf{z} - \mathbf{u}_l^*\|, \|\mathbf{z} + \mathbf{u}_l^*\|\}$ for any norm $\|\cdot\|$.

Theorem 1 (Eigenvector Estimation): Suppose that $\mu \kappa^2 r^4 \lesssim n$, and that Assumptions 1-2 hold. With probability at least $1 - O(n^{-6})$, the eigenvalue λ_l and the associated eigenvectors $\mathbf{u}_l$ and $\mathbf{w}_l$ (see Notation 1) are all real-valued, and one has the following:

- 1) (ℓ_2 and ℓ_∞ guarantees)

$$\begin{aligned} & \min \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_2 \\ & \lesssim \frac{\sigma_{\max} \sqrt{\kappa^6 n \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r^2 \log n}, \end{aligned} \quad (14a)$$

$$\begin{aligned} & \min \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_\infty \\ & \lesssim \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^4 r \log n} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}}; \end{aligned} \quad (14b)$$

these hold unchanged if $\mathbf{u}_l$ is replaced by either $\mathbf{w}_l$ or $\hat{\mathbf{u}}_l$ (cf. (10));

- 2) (statistical guarantees for linear forms of an eigenvector) for any fixed vector $\mathbf{a}$ with $\|\mathbf{a}\|_2 = 1$, the proposed estimator (11) obeys¹

$$\begin{aligned} & \min |\hat{u}_{\mathbf{a},l} \pm \mathbf{a}^\top \mathbf{u}_l^*| \\ & \lesssim \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} |\mathbf{a}^\top \mathbf{u}_l^*| \\ & \quad + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|}. \end{aligned} \quad (14c)$$

Interestingly, even though we work with eigen-decomposition of asymmetric matrices, the obtained eigenvalue and eigenvector are provably real-valued as long as a fairly mild eigen-gap condition is satisfied. This

¹If the rank $r = 1$, then the 2nd and the 3rd terms on the right-hand side of (14c) are set to be zero.

presents an important feature that cannot be derived from generic matrix perturbation theory.

To interpret the effectiveness of this theorem, we find it convenient to concentrate on the case with $r, \kappa, \mu \asymp 1$ under i.i.d. Gaussian noise. The implications in this case are summarized below.

- ℓ_2 and ℓ_∞ guarantees. The ℓ_2 and ℓ_∞ statistical guarantees derived in Theorem 1 read

$$\min \|u_l \pm u_l^*\|_2 \lesssim \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max} \sqrt{\log n}}{\Delta_l^*} =: \mathcal{E}_2; \quad (15a)$$

$$\min \|u_l \pm u_l^*\|_\infty \lesssim \frac{1}{\sqrt{n}} \mathcal{E}_2. \quad (15b)$$

Our ℓ_∞ error bound is about $O(\sqrt{n})$ times smaller than the ℓ_2 risk bound, implying that the energy of the estimation error is more or less dispersed across all entries.

- *Improved eigen-gap requirements.* Consistent estimation of u_l^* — in the sense of $\min \|u_l \pm u_l^*\|_2 = o(\|u_l^*\|_2)$ and $\min \|u_l \pm u_l^*\|_\infty = o(\|u_l^*\|_\infty)$ — is guaranteed as long as²

$$\|H\| \log n \lesssim \|M^*\|; \quad (16a)$$

$$\Delta_l^* \gtrsim \frac{\|H\| \log n}{\sqrt{n}}. \quad (16b)$$

While the condition (16a) is commonly seen in prior literature (up to some log factor), the eigen-gap requirement (16b) is in stark contrast to classical matrix perturbation theory (e.g. the Davis-Kahan $\sin \Theta$ theorem or the Wedin theorem [23]–[25]). In fact, prior theory typically requires the spacing between λ_l^* and the rest of the eigenvalues to at least exceed

$$\Delta_l^* \gtrsim \|H\| \quad (\text{prior theory}), \quad (17)$$

which is about $O(\sqrt{n}/\log n)$ times more stringent than our requirement (16b).

- *The influence of eigen-gaps and $\{a^\top u_i^*\}$ upon estimation accuracy.* For any fixed unit vector a , our theoretical guarantees for estimating $a^\top u_l^*$ read

$$\begin{aligned} \min |\hat{u}_{a,l} \pm a^\top u_l^*| &\lesssim \frac{\sigma_{\max} \sqrt{\log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}^2 \log n}{(\Delta_l^*)^2} |a^\top u_l^*| \\ &\quad + \underbrace{\sigma_{\max} \sqrt{\log n} \max_{k \neq l} \frac{|a^\top u_k^*|}{|\lambda_l^* - \lambda_k^*|}}_{\text{"interferers"}}, \end{aligned} \quad (18)$$

The first term on the right-hand side of (18) is a universal term that is no larger than $O(1/\sqrt{n})$ times the ℓ_2 bound (15a). The other two terms are more complicated, which depend on not only the spacing of the eigenvalues but also the sizes of $\{a^\top u_i^*\}_{1 \leq i \leq r}$. More concretely, (1) the influence of the target quantity $a^\top u_l^*$ upon the estimation error is captured by the multiplicative factor $\frac{\sigma_{\max}^2}{(\Delta_l^*)^2}$,

which scales inverse quadratically in Δ_l^* ; (2) the linear functionals of other eigenvectors (namely, $\{a^\top u_k^*\}_{k \neq l}$) essentially behave as “interferers” that might degrade estimation fidelity; in particular, the influence of $a^\top u_k^*$ ($k \neq l$) upon estimation loss can be understood through the coefficient $\sigma_{\max}/|\lambda_l^* - \lambda_k^*|$, which is inversely proportional to the associated eigen-gap $|\lambda_l^* - \lambda_k^*|$. Intuitively, if $a^\top u_k^*$ ($k \neq l$) becomes large, it results in stronger “interference” when estimating $a^\top u_l^*$; this adverse effect would be easier to mitigate if the gap $|\lambda_l^* - \lambda_k^*|$ increases (so that it becomes easier to differentiate the l th and the k th eigenvectors).

- *The rank-1 case.* When $r = 1$, the preceding theory can be significantly simplified as follows

$$\min \|u_l \pm u_l^*\|_2 \lesssim \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*},$$

$$\min \|u_l \pm u_l^*\|_\infty \lesssim \frac{\sigma_{\max} \sqrt{\log n}}{\lambda_{\min}^*},$$

$$\min |\hat{u}_{a,l} \pm a^\top u_l^*| \lesssim \frac{\sigma_{\max} \sqrt{\log n}}{\lambda_{\min}^*}.$$

In particular, the same estimation error bound — which is about $\sqrt{n}$ times smaller than the ℓ_2 loss — holds for any arbitrary direction as specified by a . In other words, the estimation error is more or less identical over any pre-determined direction.

Encouragingly, the above performance guarantees are minimax optimal up to some logarithmic factor, as we shall elucidate in Section III-D.

As it turns out, the feasibility of faithful eigenvector estimation is largely dictated by our ability to locate the l th eigenvalue λ_l^* and to disentangle it from the rest of the spectrum, which becomes particularly challenging if the eigen-gap Δ_l^* is very small. In light of this, we develop the following eigenvalue perturbation theory that significantly improves upon generic linear algebra theory.

Theorem 2 (Eigenvalue Estimation): Suppose that $\mu \kappa^2 r^4 \lesssim n$, and that Assumptions 1-2 hold. With probability at least $1 - O(n^{-6})$, one has

$$|\lambda_l - \lambda_l^*| \leq c_4 \sigma_{\max} \sqrt{\mu \kappa^2 r^4 \log n}. \quad (19)$$

In words, as long as the eigen-gap Δ_l^* exceeds (13a), the vanilla eigen-decomposition approach (without symmetrization) produces an estimate of λ_l^* with an estimation error at most $\Delta_l^*/2$, meaning that we have managed to locate λ_l^* with a high resolution. As mentioned previously, this eigen-gap requirement can be interpreted as $\Delta_l^* \gtrsim \|H\| \frac{\log n}{\sqrt{n}}$ in the Gaussian noise case with $r, \mu, \kappa \asymp 1$. For the sake of comparison, we remind the reader that classical linear algebra theory (e.g. Weyl’s inequality and the Bauer-Fike inequality) only provides perturbation bounds with fairly low resolution — namely, bounds at best on the order $|\lambda_l - \lambda_l^*| \leq \|H\|$ — which are highly insufficient and in fact suboptimal for our purpose.

It is worth emphasizing that the ability to obtain highly accurate eigenvalue estimates is the key to tackling small eigen-gaps. Given that we can estimate λ_l^* with a precision

²Note that when $\{H_{ij}\}$ are i.i.d. Gaussian, we have $\|H\| \asymp \sigma_{\max} \sqrt{n}$ with high probability.

$\sigma_{\max} \sqrt{\mu \kappa^2 r^4 \log n}$, we can distinguish λ_l^* from the rest of the spectrum even when the associated eigen-gap is on the same order. This also helps explain why a large body of prior theory fell short of accommodating small eigen-gaps. As discussed in [17, Section 4.3], the standard eigen-decomposition approach when applied to symmetric matrices suffers from a non-negligible bias issue, which fails to yield a high-precision eigenvalue estimate. While the recent work [28] developed a de-biasing scheme that allows one to cope with small eigen-gaps in the presence of an i.i.d. symmetric Gaussian noise matrix, the current paper covers a remarkably broader class of noise distributions by exploiting the special property of statistical asymmetry.

Finally, as mentioned before, one might be tempted to first symmetrize the data matrix (i.e. replacing $\mathbf{M}$ with $\frac{1}{2}(\mathbf{M} + \mathbf{M}^\top)$) before computing eigen-decomposition. While [17] has discussed the non-negligible bias of this approach in eigenvalue estimation, it remains unclear whether or not this natural approach suffers from a bias issue as well when estimating the eigenvectors, and, more crucially, whether or not we can still expect reliable eigenvector estimation after symmetrizing the data matrix. Addressing these questions is instrumental in understanding whether asymmetry plays an important role or it merely provides theoretical convenience. While a theory for this is yet to be developed, we shall compare these two approaches numerically and demonstrate the gains of our approach in Section V.

D. Minimax Lower Bounds

To assess the tightness of our statistical guarantees, we develop localized minimax lower bounds on eigenvector estimation, aimed at providing in-depth understanding about how the estimation difficulty changes as a function of certain salient parameters of the problem instance. To derive these lower bounds, we consider the non-asymptotic local minimax framework (see, e.g. [29]–[31]), an approach built upon the concept of the hardest local alternatives that finds its roots in [32].

Before stating our main results, let us first define several sets that are properly localized around the truth. Let $\mathbb{S}^n \subseteq \mathbb{R}^n$ represent the set of $n \times n$ symmetric matrices. Denoting by $\lambda_l(\mathbf{A})$ the l th largest eigenvalue of $\mathbf{A}$ and $\mathbf{u}_l(\mathbf{A})$ the associated eigenvector of a symmetric matrix $\mathbf{A}$, we define

$$\mathcal{M}_0(\mathbf{M}^*) := \left\{ \mathbf{A} \in \mathbb{S}^n \mid \text{rank}(\mathbf{A}) = r, \lambda_i(\mathbf{A}) = \lambda_i^* \right. \\ \left. (1 \leq i \leq r), \|\mathbf{A} - \mathbf{M}^*\|_F \leq \frac{\sigma_{\min}}{2} \right\} \quad (20a)$$

$$\mathcal{M}_1(\mathbf{M}^*) := \left\{ \mathbf{A} \in \mathbb{S}^n \mid \text{rank}(\mathbf{A}) = r, \lambda_i(\mathbf{A}) = \lambda_i^* \right. \\ \left. (1 \leq i \leq r), \|\mathbf{u}_l(\mathbf{A}) - \mathbf{u}_l^*\|_2 \leq \frac{\sigma_{\min}}{4|\lambda_l^*|} \right\} \quad (20b)$$

$$\mathcal{M}_2(\mathbf{M}^*) := \left\{ \mathbf{A} \in \mathbb{S}^n \mid \text{rank}(\mathbf{A}) = r, \lambda_i(\mathbf{A}) = \lambda_i^* \right. \\ \left. (1 \leq i \leq r), \|\mathbf{u}_l(\mathbf{A}) - \mathbf{u}_l^*\|_2 \leq c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|} \right\} \quad (20c)$$

for some sufficiently large constant $c_4 > 0$.

Now we are ready to state our results, the proof of which is deferred to Appendix E.

Theorem 3 (Minimax Lower Bounds): Consider any $1 \leq l \leq r \leq n/2$. Suppose that $H_{ij} \stackrel{\text{ind.}}{\sim} \mathcal{N}(0, \sigma_{ij}^2)$ with $\sigma_{ij}^2 \geq \sigma_{\min}^2$, and assume that $4\sigma_{\min}\sqrt{n} \leq |\lambda_l^*|$ and $\sigma_{\min} \leq \Delta_l^*$.

- 1) There exist some constants $c_0, c_1 > 0$ such that

$$\inf_{\hat{\mathbf{u}}_l} \sup_{\mathbf{A} \in \mathcal{M}_0(\mathbf{M}^*)} \mathbb{E} \left[\min \|\hat{\mathbf{u}}_l \pm \mathbf{u}_l(\mathbf{A})\|_2 \right] \geq c_0 \frac{\sigma_{\min}}{\Delta_l^*}; \quad (21a)$$

$$\inf_{\hat{\mathbf{u}}_{a,l}} \sup_{\mathbf{A} \in \mathcal{M}_0(\mathbf{M}^*)} \mathbb{E} \left[\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l(\mathbf{A})| \right] \\ \geq c_1 \left\{ |\mathbf{a}^\top \mathbf{u}_l^*| \frac{\sigma_{\min}^2}{\Delta_l^{*2}} + \sigma_{\min} \max_{k:k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \right\}. \quad (21b)$$

- 2) There exists some constant $c_2 > 0$ such that

$$\inf_{\hat{\mathbf{u}}_{a,l}} \sup_{\mathbf{A} \in \mathcal{M}_1(\mathbf{M}^*)} \mathbb{E} \left[\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l(\mathbf{A})| \right] \\ \geq c_2 \frac{\sigma_{\min} \|\mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{a}\|_2}{|\lambda_l^*|}. \quad (21c)$$

- 3) Then there exists some constant $c_3 > 0$ such that

$$\inf_{\hat{\mathbf{u}}_l} \sup_{\mathbf{A} \in \mathcal{M}_2(\mathbf{M}^*)} \mathbb{E} \left[\min \|\hat{\mathbf{u}}_l \pm \mathbf{u}_l(\mathbf{A})\|_2 \right] \geq c_3 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|}. \quad (21d)$$

Here, the infimum in (21a) and (21d) is over all eigenvector estimators, while the infimum in (21b) and (21c) is over all estimators for the linear form of the l th eigenvector.

Remark 6: Note that each part of Theorem 3 captures one bottleneck for the estimation task. Given that enlarging the set of matrices under consideration can only lead to increased minimax lower bound, we can take the supremum over the union $\mathcal{M}_0(\mathbf{M}^*) \cup \mathcal{M}_1(\mathbf{M}^*) \cup \mathcal{M}_2(\mathbf{M}^*)$ in order to yield a minimax lower bound that reflects all these bottlenecks at once.

Consider again the case with $r, \kappa, \mu \asymp 1$ under i.i.d. Gaussian noise, and the case when $|\mathbf{a}^\top \mathbf{u}_l^*| \leq (1-\epsilon)\|\mathbf{a}\|_2$ for an arbitrarily small constant $\epsilon > 0$ (so that $\mathbf{a}$ is not perfectly aligned with $\mathbf{u}_l^*$). The theoretical guarantees derived for our estimators (see Theorem 1, or (15) and (18)) match the minimax lower bounds in Theorem 3 up to some logarithmic factor, including both ℓ_2 loss and the risk for estimating an arbitrary linear form $\mathbf{a}^\top \mathbf{u}_l^*$. In particular, the dependencies of the estimation risk on Δ_l^* , $|\lambda_l^* - \lambda_k^*|$, and $\{\mathbf{a}^\top \mathbf{u}_k^*\}$ in Theorem 1 — which might seem complicated at first glance — are all optimal modulo some log factor. All this confirms the effectiveness and optimality of the proposed estimator in fine-grained eigenvector estimation.

IV. INFERENCE AND UNCERTAINTY QUANTIFICATION

This section moves one step further to the task of statistical inference, with the aim of constructing valid and efficient confidence intervals for the linear functionals $\mathbf{a}^\top \mathbf{u}_l^*$ and the eigenvalues λ_l^* ($1 \leq l \leq r$). More precisely, for any target

coverage level $0 < 1 - \alpha < 1$, we seek to compute intervals such that

$$\begin{aligned} \mathbb{P}\{\mathbf{a}^\top \mathbf{u}_l^* \in [c_{lb}^{\mathbf{a}}, c_{ub}^{\mathbf{a}}]\} &\approx 1 - \alpha \quad \text{or} \\ \mathbb{P}\{-\mathbf{a}^\top \mathbf{u}_l^* \in [c_{lb}^{\mathbf{a}}, c_{ub}^{\mathbf{a}}]\} &\approx 1 - \alpha. \end{aligned} \quad (22)$$

and

$$\mathbb{P}\{\lambda_l^* \in [c_{lb}^\lambda, c_{ub}^\lambda]\} \approx 1 - \alpha \quad (23)$$

Here, we have taken into account the difficulty in distinguishing $\mathbf{u}_l^*$ from $-\mathbf{u}_l^*$ using only the data $\mathbf{M}$. As we shall see, providing precise confidence intervals for the linear functional is much more challenging than the estimation task, and our theory requires additional assumptions beyond Assumption 2 in order for our procedures to have exact coverage.

A. Algorithms

1) Confidence Intervals for Linear Forms of Eigenvectors:

We start by constructing confidence intervals for the linear form $\mathbf{a}^\top \mathbf{u}_l^*$. Towards this, two ingredients are needed: (1) a nearly unbiased estimate of $\mathbf{a}^\top \mathbf{u}_l^*$, and (2) a valid length of the interval (which is typically built upon a certain variance estimate). We describe these ingredients as follows.

a) *A modified nearly unbiased estimator* $\hat{\mathbf{u}}_{a,l}^{\text{modified}}$. While the estimator $\hat{\mathbf{u}}_{a,l}$ (cf. (11)) discussed previously enjoys minimax optimal statistical accuracy, we find it more convenient to work with a modified estimator when conducting uncertainty quantification, particularly in the regime when $\mathbf{a}^\top \mathbf{u}_l^*$ is very small.³ More specifically, we introduce

$$\hat{\mathbf{u}}_{a,l}^{\text{modified}} := \begin{cases} \frac{1}{2} \mathbf{a}^\top (\mathbf{u}_l + \mathbf{w}_l), & \text{if } |\mathbf{a}^\top \mathbf{u}_l^*| \text{ is "small",} \\ \hat{\mathbf{u}}_{a,l}, & \text{otherwise.} \end{cases} \quad (24)$$

We shall make precise in Algorithm 1 what it means by “small” in a practical and data-dependent manner. As before, the proposed procedures are fully data-driven, without requiring prior knowledge about the noise distributions. A few informal yet important remarks are in order.

- When $\mathbf{a}^\top \mathbf{u}_l^*$ is very “small” in magnitude, both $\mathbf{a}^\top \mathbf{u}_l$ and $\mathbf{a}^\top \mathbf{w}_l$ serve as unbiased estimates of $\mathbf{a}^\top \mathbf{u}_l^*$, and the averaging operation further reduces the uncertainty. Here, the quantity $\sqrt{\hat{v}_{a,l}}$ in Algorithm 1 captures the level of the smallest possible uncertainty when estimating $\mathbf{a}^\top \mathbf{u}_l^*$.
- When $\mathbf{a}^\top \mathbf{u}_l^*$ is not very “small”, the procedure is identical to estimator $\hat{\mathbf{u}}_{a,l}$ proposed previously.

b) *Construction of confidence intervals:* As it turns out, the estimation error of the modified estimator $\hat{\mathbf{u}}_{a,l}^{\text{modified}}$ is well approximated by a zero-mean Gaussian random variable with tractable variance $v_{a,l}^*$ (to be formalized shortly). Motivated by this observation, we propose to first obtain an estimate of the variance of $\hat{\mathbf{u}}_{a,l}^{\text{modified}}$ — denoted by $\hat{v}_{a,l}$. The proposed confidence interval for any given coverage level $0 < 1 - \alpha < 1$ is then given by

$$\text{CI}_{1-\alpha}^{\mathbf{a}} := \left[\hat{\mathbf{u}}_{a,l}^{\text{modified}} \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\hat{v}_{a,l}} \right], \quad (25)$$

³More precisely, in the regime where $\mathbf{a}^\top \mathbf{u}_l^*$ is very small, the uncertainty of $\hat{\mathbf{u}}_{a,l}^{\text{modified}}$ is nearly Gaussian, while that of $\hat{\mathbf{u}}_{a,l}$ is non-Gaussian and more complicated to describe.

where we abbreviate $[b \pm c] := [b - c, b + c]$, and $\Phi(\cdot)$ is the CDF of a standard Gaussian.

c) *Variance estimates:* We shall take a moment to discuss how to obtain the variance estimate $\hat{v}_{a,l}$. As will be seen momentarily, the proposed estimator $\hat{\mathbf{u}}_{a,l}^{\text{modified}}$ admits — modulo some global sign — the following first-order approximation for a broad range of settings (up to global sign):

$$\hat{\mathbf{u}}_{a,l}^{\text{modified}} \approx \mathbf{a}^\top \mathbf{u}_l^* + \underbrace{\frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}_{\text{uncertainty term}}, \quad (26)$$

where $\mathbf{a}_l^{\perp} := \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*$. For a broad family of noise distributions, the uncertainty term is approximately zero-mean Gaussian with variance

$$\begin{aligned} v_{a,l}^* &:= \text{Var} \left[\frac{1}{2\lambda_l^*} (\mathbf{a}_l^{\perp})^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right] \\ &= \frac{1}{4\lambda_l^{*2}} \sum_{1 \leq i,j \leq n} (\mathbf{a}_{l,i}^{\perp} \mathbf{u}_{l,j}^* + \mathbf{a}_{l,j}^{\perp} \mathbf{u}_{l,i}^*)^2 \sigma_{ij}^2, \end{aligned} \quad (27)$$

where $\mathbf{a}_{l,j}^{\perp}$ denotes the j th entry of $\mathbf{a}_l^{\perp}$. At first glance, evaluating this variance quantity precisely requires prior knowledge about the noise level in addition to the truth $(\mathbf{u}_l^*, \lambda_l^*)$. To enable a model-agnostic and data-driven estimate of $v_{a,l}^*$, we make the observation that $\frac{1}{4\lambda_l^{*2}} \sum_{i,j} (\mathbf{a}_{l,i}^{\perp} \mathbf{u}_{l,j}^* + \mathbf{a}_{l,j}^{\perp} \mathbf{u}_{l,i}^*)^2 H_{ij}^2$ is very close to $v_{a,l}^*$ based on the concentration of measure phenomenon. This in turn motivates us to estimate $v_{a,l}^*$ via a plug-in approach: (1) replacing $\mathbf{u}_l^*$ with an estimate $\hat{\mathbf{u}}_l$, (2) replacing $\mathbf{a}_l^{\perp}$ with a plug-in estimate $\hat{\mathbf{a}}_l^{\perp}$, (3) using λ_l in place of λ_l^* , and (4) replacing H_{ij} with an estimate $\hat{H}_{ij}$, where

$$\begin{aligned} \hat{\mathbf{u}}_l &:= \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l) \quad \text{and} \\ \hat{\mathbf{H}} &= [\hat{H}_{ij}]_{1 \leq i,j \leq n} := \mathbf{M} - \mathbf{M}_{\text{svd},r}, \end{aligned} \quad (28)$$

with $\mathbf{M}_{\text{svd},r} := \arg \min_{\text{rank}(\mathbf{Z}) \leq r} \|\mathbf{M} - \mathbf{Z}\|_F$ the best rank- r approximation of $\mathbf{M}$. As a byproduct, this variance estimate in turn allows us to specify whether $\mathbf{a}^\top \mathbf{u}_l^*$ is “small” (the case where $\mathbf{a}^\top \mathbf{u}_l^*$ is comparable to or smaller than the typical size of the uncertainty component).

Remark 7: For estimating $\mathbf{H}$ (or equivalently, $\mathbf{M}^*$), it has been shown that the SVD-based approach achieves appealing entrywise accuracy (e.g. [9], [33]).

For ease of reference, the proposed procedure is summarized in Algorithm 1. The computational cost of this procedure mostly lies in computing the eigen-decomposition and the SVD of $\mathbf{M}$.

2) *Confidence Intervals for Eigenvalues:* Moving beyond linear forms of eigenvectors, one might also be interested in performing inference on the eigenvalues of interest. As it turns out, this task is simpler than inferring linear functionals of eigenvectors; one can simply estimate λ_l^* via the l th eigenvalue λ_l (see Notation 1) and compute a confidence interval based on the distributional characterization of λ_l . There is absolutely no need for careful de-biasing, since the eigen-decomposition approach automatically exploits the asymmetry structure of noise to suppress bias. Similar to Section IV-A1, our procedure for performing inference on λ_l^* is distribution-free (i.e. it does not require prior knowledge about the noise variance) and adaptive to heteroscedasticity of noise.

Algorithm 1 Constructing a Confidence Interval for the Linear Form $\mathbf{a}^\top \mathbf{u}_l^*$

- 1: Compute the l th eigenvalue λ_l of $\mathbf{M}$, and the associated right eigenvector $\mathbf{u}_l$ and left eigenvector $\mathbf{w}_l$ such that $\text{Re}(\mathbf{w}_l^\top \mathbf{u}_l) \geq 0$ (see Notation 1).
- 2: Compute the following estimator

$$\begin{aligned} \hat{u}_{a,l} &:= \min \left\{ \sqrt{\left| \frac{1}{\mathbf{w}_l^\top \mathbf{u}_l} (\mathbf{a}^\top \mathbf{u}_l) (\mathbf{a}^\top \mathbf{w}_l) \right|}, \|\mathbf{a}\|_2 \right\}, \\ \hat{u}_{a,l}^{\text{modified}} &:= \begin{cases} \frac{1}{2} \mathbf{a}^\top (\mathbf{u}_l + \mathbf{w}_l), & \text{if } |\mathbf{a}^\top \mathbf{u}_l| \leq c_b \sqrt{\hat{v}_{a,l}} \log^{1.5} n, \\ \hat{u}_{a,l}, & \text{else,} \end{cases} \end{aligned} \quad (29)$$

where $c_b > 0$ is some sufficiently large constant. Here, $\hat{v}_{a,l} = \text{ESTIMATE-VARIANCE } \hat{\mathbf{a}}_l^\perp, \hat{\mathbf{u}}_l, \lambda_l$ with $\hat{\mathbf{a}}_l^\perp := \mathbf{a} - (\mathbf{a}^\top \hat{\mathbf{u}}_l) \hat{\mathbf{u}}_l$ and $\hat{\mathbf{u}}_l := \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l)$.

- 3: For any prescribed coverage level $1 - \alpha$, compute the confidence interval as

$$\text{CI}_{1-\alpha}^a := [\hat{u}_{a,l}^{\text{modified}} \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\hat{v}_{a,l}}]. \quad (30)$$

1: **function** ESTIMATE-VARIANCE($\mathbf{a}, \mathbf{u}, \lambda$)

2: Compute $\hat{\mathbf{H}} := \mathbf{M} - \mathbf{M}_{\text{svd},r}$ with $\mathbf{M}_{\text{svd},r} := \arg \min_{\text{rank}(\mathbf{Z}) \leq r} \|\mathbf{M} - \mathbf{Z}\|_F$.

3: **return** $v := \frac{1}{4\lambda^2} \sum_{1 \leq i,j \leq n} (a_i u_j + a_j u_i)^2 \hat{H}_{ij}^2$.

To be more specific, a crucial observation, which we will make precise shortly, is as follows

$$\lambda_l \approx \lambda_l^* + \mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*, \quad (31)$$

where the uncertainty term $\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*$ is approximately Gaussian with variance

$$v_{\lambda,l}^* := \text{Var} [\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*] = \sum_{1 \leq i,j \leq n} (u_{l,i}^* u_{l,j}^*)^2 \sigma_{ij}^2. \quad (32)$$

Similar to the estimator $\hat{v}_{a,l}$, we propose to estimate $v_{\lambda,l}^*$ via the following estimator

$$\hat{v}_{\lambda,l} := \sum_{1 \leq i,j \leq n} (\hat{u}_{l,i} \hat{u}_{l,j})^2 \hat{H}_{ij}^2, \quad (33)$$

with $\hat{\mathbf{u}}_l$ and $\hat{\mathbf{H}}$ defined in (28). All this suggests the following confidence interval for λ_l^* :

$$\text{CI}_{1-\alpha}^\lambda := [\lambda_l \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\hat{v}_{\lambda,l}}]. \quad (34)$$

See Algorithm 2 for the complete procedure for performing inference on λ_l^* .

B. Theoretical Guarantees

We now set out to justify the validity of the proposed confidence intervals. Before proceeding, we need to impose another set of assumptions on the noise levels:

Algorithm 2 Constructing a Confidence Interval for λ_l^*

- 1: Compute the l th eigenvalue λ_l of $\mathbf{M}$, and the associated right eigenvector $\mathbf{u}_l$ and left eigenvector $\mathbf{w}_l$ such that $\mathbf{w}_l^\top \mathbf{u}_l \geq 0$ (see Notation 1).
- 2: For any prescribed coverage level $1 - \alpha$, compute the confidence interval as

$$\text{CI}_{1-\alpha}^\lambda := [\lambda_l \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\hat{v}_{\lambda,l}}]. \quad (35)$$

Here, $\hat{v}_{\lambda,l} = \text{ESTIMATE-VARIANCE } \hat{\mathbf{u}}_l, \hat{\mathbf{u}}_l, 1$ with $\hat{\mathbf{u}}_l := \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l)$.

Assumption 3: Suppose that the noise parameters defined in Assumption 1 satisfy

$$\sigma_{\max} \kappa^3 \sqrt{\mu^3 r^3} \log^{1.5} n = o(\Delta_l^*), \quad (36a)$$

$$\sigma_{\max} \mu \kappa^2 r^2 \sqrt{n} \log^{1.5} n = o(\lambda_{\min}^*), \quad (36b)$$

$$\frac{\sigma_{\max}}{\sigma_{\min}} \asymp 1, \quad B = o\left(\sigma_{\min} \sqrt{\frac{n}{\mu \log n}}\right). \quad (36c)$$

Armed with the above assumptions, we are ready to present theoretical guarantees.

Theorem 4 (Validity of Confidence Intervals (Rank- r)): Assume $\mathbf{M}^*$ is rank- r and μ -incoherent with $\mu^4 \kappa^8 r^2 \log^2 n = o(n)$. Given any integer $1 \leq l \leq r$, under Assumptions 1 and 3, the confidence interval returned by Algorithm 2 obeys

$$\mathbb{P}\{\lambda_l^* \in \text{CI}_{1-\alpha}^\lambda\} = 1 - \alpha + o(1) \quad (37)$$

uniformly over all $0 < \alpha < 1$. In addition, fix an arbitrarily small constant $0 < \epsilon < 1$, and consider any fixed vector $\mathbf{a}$ with $\|\mathbf{a}\|_2 = 1$ obeying

$$|\mathbf{a}^\top \mathbf{u}_l^*| \leq 1 - \epsilon, \quad |\mathbf{a}^\top \mathbf{u}_l^*| = o\left(\frac{\Delta_l^{*2}}{|\lambda_l^*| \sigma_{\max} \kappa^4 r^2 \mu \log n}\right), \quad (38a)$$

$$|\mathbf{a}^\top \mathbf{u}_k^*| = o\left(\frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \sqrt{\mu \kappa^4 r^3 \log n}}\right), \quad \forall k \neq l. \quad (38b)$$

Then the confidence interval returned by Algorithms 1 obeys

$$\begin{aligned} \mathbb{P}\{\mathbf{a}^\top \mathbf{u}_l^* \in \text{CI}_{1-\alpha}^a\} &= 1 - \alpha + o(1) \quad \text{or} \\ \mathbb{P}\{-\mathbf{a}^\top \mathbf{u}_l^* \in \text{CI}_{1-\alpha}^a\} &= 1 - \alpha + o(1) \end{aligned} \quad (39)$$

uniformly over all $0 < \alpha < 1$.

We immediately make note of a direct consequence of Theorem 4 (by taking $\mathbf{a}$ to be the k th standard basis vector $\mathbf{e}_k$), which concerns entrywise confidence intervals for $\mathbf{u}_l^*$.

Corollary 1 (Validity of Entrywise Confidence Intervals (Rank- r)): Assume that $\mathbf{M}^*$ is rank- r and μ -incoherent with $\mu^4 = o(n / \log n)$. Suppose that Assumptions 1 and 3 hold and that $\sqrt{\frac{\mu^2 \kappa^6 r^6 \log n}{n}} |\lambda_l^*| = o(\Delta_l^*)$. Then the confidence interval constructed in Algorithms 1 with $\mathbf{a} = \mathbf{e}_k$ satisfies

$$\begin{aligned} \mathbb{P}\{u_{l,k}^* \in \text{CI}_{1-\alpha}^{e_k}\} &= 1 - \alpha + o(1) \quad \text{or} \\ \mathbb{P}\{-u_{l,k}^* \in \text{CI}_{1-\alpha}^{e_k}\} &= 1 - \alpha + o(1); \end{aligned} \quad (40)$$

this holds true uniformly over all $0 < \alpha < 1$.

The above results confirm the validity of our inferential procedure in high dimension. Our theory applies to a very broad family of noise distributions, without the need of any prior knowledge about detailed noise distributions or noise levels $\{\sigma_{ij}\}$. The results are fully adaptive to heteroscedasticity of noise, making them appealing for practical scenarios. In the sequel, we discuss several important implications of our results in a more quantitative manner. For simplicity of discussion, we shall concentrate on the case where $r, \kappa, \mu \asymp 1$ and consider Gaussian noise with $\sigma_{\max} \asymp \sigma_{\min}$. We shall also assume $\|\mathbf{a}\|_2 = 1$ without loss of generality.

- We start with the rank-1 case (i.e. $r = 1$), in which we have $\Delta_l^* = \infty$ according to the definition (12). In this case, the conditions (38) admit considerable simplification:

$$|\mathbf{a}^\top \mathbf{u}_1^*| \leq 1 - \epsilon \quad (41)$$

for any small constant $\epsilon > 0$. This means that our theory covers a very wide range of linear functionals $\mathbf{a}^\top \mathbf{u}_1^*$; basically, the proposed confidence interval finds theoretical support unless the preconceived direction $\mathbf{a}$ is already highly aligned with the truth $\mathbf{u}_1^*$.

- Going beyond the rank-1 case, the eigen-gap condition imposed in Assumption 3 and Corollary 1 reads

$$\Delta_l^* \gtrsim \sigma_{\max} \text{poly log } n \asymp \frac{\|\mathbf{H}\| \text{poly log } n}{\sqrt{n}}, \quad (42)$$

which is allowed to be substantially smaller than the spectral norm $\|\mathbf{H}\|$ of the noise matrix.

- Under such eigen-gap requirements, we develop an informative distributional theory underlying the proposed estimators $\hat{\mathbf{u}}_{\mathbf{a},l}$, provided that $\{\mathbf{a}^\top \mathbf{u}_k^*\}_{1 \leq k \leq r}$ are not too large in the sense that

$$\begin{aligned} |\mathbf{a}^\top \mathbf{u}_l^*| &\lesssim \frac{\Delta_l^{*2}}{|\lambda_l^*| \sigma_{\max} \log^{1.5} n}; \\ |\mathbf{a}^\top \mathbf{u}_k^*| &\lesssim \frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \log n}, \quad k \neq l. \end{aligned} \quad (43)$$

In words, the condition (43) requires that both the target quantity $\mathbf{a}^\top \mathbf{u}_l^*$ and the “interferers” are dominated by the respective (normalized) eigen-gaps. An important instance that automatically satisfies such conditions has been singled out in Corollary 1, leading to appealing entrywise distributional characterizations and inferential procedures.

- While our theorems focus on asymmetric noise matrices, one can combine them with a simple asymmetrization trick to yield valid confidence intervals when $\mathbf{H}$ is a symmetric matrix with homoscedastic Gaussian noise. See Appendix G for detailed discussion.

We note, however, that the additional requirement (43) makes our inference results less general than our estimation guarantees in Section III-C, except for the rank-1 case. In fact, if (43) is violated, then the influence of the eigen-gaps might become the dominant factor in the uncertainty term and needs to be quantified in a precise fashion. Achieving this calls for more refined theoretical analysis, and we leave it to future investigation.

C. Key Ingredients Behind Theorem 4

Next, we single out two key ingredients that shed light on not only the validity of, but also the statistical efficiency of, our inferential procedures. To be specific, Theorem 4 is mainly built upon distributional guarantees developed for the proposed estimator $\hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}}$ and the l th eigenvalue λ_l of $\mathbf{M}$. In a nutshell, the quantity $\hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}}$ (resp. λ_l) is a nearly unbiased estimator of $\mathbf{a}^\top \mathbf{u}_l^*$ (resp. λ_l^*) and is approximately Gaussian. We formalize this distributional theory as follows, whose proof is deferred to Appendix C.

Theorem 5 (Distributional Characterization (Rank- r)): Instate the assumptions of Theorem 4. Let $\hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}}$ (cf. Algorithm 1) be the proposed estimate for $\mathbf{a}^\top \mathbf{u}_l^*$. Then one can write

$$\begin{aligned} \hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}} &= \mathbf{a}^\top \mathbf{u}_l^* + \sqrt{v_{\mathbf{a},l}^*} W_{\mathbf{a},l} + \zeta \quad \text{or} \\ -\hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}} &= \mathbf{a}^\top \mathbf{u}_l^* + \sqrt{v_{\mathbf{a},l}^*} W_{\mathbf{a},l} + \zeta \end{aligned} \quad (44a)$$

and

$$\lambda_l = \lambda_l^* + \sqrt{v_{\lambda,l}^*} W_{\lambda,l} + \xi, \quad (44b)$$

where $v_{\mathbf{a},l}^*$ and $v_{\lambda,l}^*$ are defined respectively in (27) and (32), and

$$\begin{aligned} W_{\mathbf{a},l} &:= \frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\lambda_l^* \sqrt{v_{\mathbf{a},l}^*}} \quad \text{and} \\ W_{\lambda,l} &:= \frac{\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*}{\sqrt{v_{\lambda,l}^*}} \end{aligned} \quad (45)$$

with $\mathbf{a}_l^\perp := \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*$. The residual terms obey $|\zeta| = o(\sqrt{v_{\mathbf{a},l}^*})$ and $|\xi| = o(\sqrt{v_{\lambda,l}^*})$ with probability at least $1 - O(n^{-5})$. In addition, one has

$$\begin{aligned} \sup_{z \in \mathbb{R}} |\mathbb{P}(W_{\mathbf{a},l} \leq z) - \Phi(z)| &= o(1) \quad \text{and} \\ \sup_{z \in \mathbb{R}} |\mathbb{P}(W_{\lambda,l} \leq z) - \Phi(z)| &= o(1). \end{aligned} \quad (46)$$

Informally, this theorem reveals the tightness of the following first-order approximation (up to global sign)

$$\begin{aligned} \hat{\mathbf{u}}_{\mathbf{a},l}^{\text{modified}} &\approx \mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} (\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \quad \text{and} \\ \lambda_l &\approx \lambda_l^* + \mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^* \end{aligned} \quad (47)$$

for a wide range of directions $\mathbf{a}$. In addition, the first-order approximations are close in distribution to Gaussian random variables. Such distributional characterizations would immediately lead to $(1 - \alpha)$ -confidence intervals for any α , if an oracle had informed us of the quantities $v_{\mathbf{a},l}^*$ and $v_{\lambda,l}^*$. Fortunately, the proposed $\hat{v}_{\mathbf{a},l}$ and $\hat{v}_{\lambda,l}$ (see Algorithms 1 and 2) serve as highly accurate estimates of $v_{\mathbf{a},l}^*$ and $v_{\lambda,l}^*$, respectively, and can be employed in place of $v_{\mathbf{a},l}^*$ and $v_{\lambda,l}^*$. This observation is asserted by the following theorem; the proof is postponed to Appendix D.

Theorem 6 (Accuracy of Variance Estimates (Rank- r)): Instate the assumptions of Theorem 4. With probability at least $1 - O(n^{-10})$, the variance estimators proposed in Algorithms 1-2 satisfy

$$\hat{v}_{\mathbf{a},l} = (1 + o(1))v_{\mathbf{a},l}^* \quad \text{and} \quad \hat{v}_{\lambda,l} = (1 + o(1))v_{\lambda,l}^*. \quad (48)$$

Clearly, putting Theorems 5-6 together immediately establishes Theorem 4.

V. NUMERICAL EXPERIMENTS

This section consists of numerical experiments in various settings, in order to demonstrate the performance of our estimation and inference procedures and their accompanying theory.

A. Eigen-Decomposition After Symmetrization?

As mentioned in the discussions after Theorem 2, one may consider first symmetrizing the data matrix with $\frac{1}{2}(\mathbf{M} + \mathbf{M}^\top)$ before computing the eigen-decomposition. It is unclear whether this procedure provides sensible eigenvector estimators in the presence of small eigen-gaps and heteroscedastic noise. This subsection is devoted to understanding the potential sub-optimality of this approach via a simple example. Specifically, consider the rank-2 model where

$$\mathbf{M}^* = \lambda_1^* \mathbf{u}_1^* \mathbf{u}_1^{*\top} + \lambda_2^* \mathbf{u}_2^* \mathbf{u}_2^{*\top}, \quad (49)$$

with $\mathbf{H}$ satisfying our usual assumptions. To simplify presentation, we define

1. Spectral-asym: eigen-decomposition applied to the observed asymmetric data matrix $\mathbf{M}$;
2. Spectral-sym: eigen-decomposition applied to the symmetrized data matrix $\frac{1}{2}(\mathbf{M} + \mathbf{M}^\top)$.

a) *A Numerical Example With Heteroscedastic Gaussian Noise:* We start by looking at an example with

$$\mathbf{u}_1^* = \frac{1}{\sqrt{n}} \mathbf{1}_n \quad \text{and} \quad \mathbf{u}_2^* = \frac{1}{\sqrt{n}} \begin{bmatrix} \mathbf{1}_{n/2} \\ -\mathbf{1}_{n/2} \end{bmatrix}. \quad (50)$$

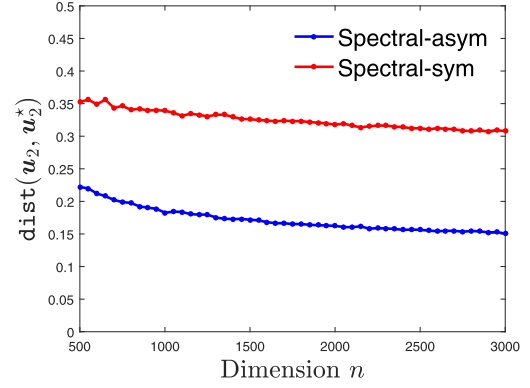
The noise matrix $\mathbf{H}$ is assumed to have independent zero-mean Gaussian entries with variance

$$\begin{aligned} \text{Var}(\mathbf{H}) &:= [\text{Var}(H_{ij})]_{1 \leq i, j \leq n} \\ &= \sigma_1^2 \begin{bmatrix} \mathbf{1}_{n/2} \mathbf{1}_{n/2}^\top - \frac{1}{2} \mathbf{I}_{n/2} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{bmatrix} + \sigma_2^2 \left(\mathbf{1}_n \mathbf{1}_n^\top - \frac{1}{2} \mathbf{I}_n \right), \end{aligned} \quad (51)$$

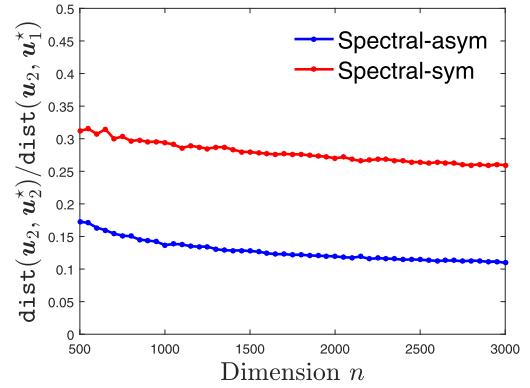
and hence the variance of the symmetrized data satisfies

$$\begin{aligned} \text{Var}\left(\frac{\mathbf{H} + \mathbf{H}^\top}{2}\right) &:= [\text{Var}(\tfrac{1}{2}(H_{ij} + H_{ji}))]_{1 \leq i, j \leq n} \\ &= \frac{\sigma_1^2}{2} \begin{bmatrix} \mathbf{1}_{n/2} \mathbf{1}_{n/2}^\top & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{bmatrix} + \frac{\sigma_2^2}{2} \mathbf{1}_n \mathbf{1}_n^\top. \end{aligned} \quad (52)$$

We plot in Fig. 1 the numerical performance of both Spectral-asym and Spectral-sym in estimating $\mathbf{u}_2^*$. For Spectral-asym, the estimator $\hat{\mathbf{u}}_2$ is constructed as in the expression (10), whereas the second eigenvector of $\frac{1}{2}(\mathbf{M} + \mathbf{M}^\top)$ is used for Spectral-sym. As can be seen, Spectral-asym strictly outperforms Spectral-sym in all cases, thus unveiling the real benefits of exploiting asymmetry in eigen-decomposition. The interested reader is deferred to Appendix H for some high-level interpretation of this phenomenon.



(a) ℓ_2 estimation error $\text{dist}(\mathbf{u}_2, \mathbf{u}_2^*)$



(b) relative ℓ_2 error $\text{dist}(\mathbf{u}_2, \mathbf{u}_2^*)/\text{dist}(\mathbf{u}_2, \mathbf{u}_1^*)$

Fig. 1. Numerical performance of Spectral-asym vs. Spectral-sym when $\lambda_1^* = 1$ and $\lambda_2^* = 0.95$. Here, we define $\text{dist}(\mathbf{u}, \mathbf{v}) := \min\{\|\mathbf{u} - \mathbf{v}\|_2, \|\mathbf{u} + \mathbf{v}\|_2\}$. (a) plots the ℓ_2 estimation error vs. the dimension n , whereas (b) plots the relative ℓ_2 error vs. n . For each n , the results are averaged over 1000 trials, with $\sigma_1 = 1/\sqrt{n \log n}$ and $\sigma_2 = 0.1/\sqrt{n \log n}$ (see the expression (51)).

B. Estimation

This section is devoted to numerically studying the efficiency of our estimators for linear functionals of the eigenvectors. Given any fixed $\mathbf{a}$ with $\|\mathbf{a}\|_2 = 1$, our estimator is constructed as in the expression (11). Theorem 1 (in particular, the upper bound (14c)) together with Theorem 3 implies that the estimation error is controlled by the eigen-gap, the true signal strength $|\mathbf{a}^\top \mathbf{u}_i^*|$, and the magnitude of the “interferers”. In the following, we examine qualitatively the effects of these quantities upon the estimation errors through some simple yet representative examples.

a) *Settings:* Consider the rank-2 case as in Eq. (49) again, where the leading eigenvalue is set to be 1 and the orthonormal pair $\mathbf{u}_1^*$ and $\mathbf{u}_2^*$ are randomly generated. We focus on estimating the linear functionals of the form $\mathbf{a}^\top \mathbf{u}_2^*$. In the following, several quantities that are selected to examine how they affect the estimation error $|\hat{u}_{a,2} - \mathbf{a}^\top \mathbf{u}_2^*|$ include the ground-truth $|\mathbf{a}^\top \mathbf{u}_2^*|$, the interferer $|\mathbf{a}^\top \mathbf{u}_1^*|$, and the eigen-gap $\Delta_2^* = \lambda_1^* - \lambda_2^*$. In the following, we consider two scenarios: the case where the observation matrix is generated from a rank-2 underlying matrix $\mathbf{M}^*$ plus heteroscedastic Gaussian noise; and the case where entries are missing at random with sampling rate 0.1 on top of the above-mentioned observation model (more details are given below).

b) *Heteroscedastic Gaussian noise*: Consider a heteroscedastic Gaussian noise matrix $\mathbf{H}$ with independent entries $H_{ij} \sim \mathcal{N}(0, \sigma_{ij}^2)$ obeying

$$\text{Var}(\mathbf{H}) := [\text{Var}(H_{ij})]_{1 \leq i, j \leq n} = \begin{bmatrix} \sigma_1^2 & & \\ & (\sigma_1 + \delta_\sigma)^2 & \\ & \vdots & \\ & & (\sigma_1 + (n-1)\delta_\sigma)^2 \end{bmatrix} \cdot \mathbf{1}_n^\top, \quad (53)$$

where δ_σ determines the spacing between adjacent standard deviation. The parameters chosen to be

$$\sigma_1 = 0.1/\sqrt{n \log n}, \quad \delta_\sigma = 0.9/((n-1)\sqrt{n \log n}).$$

c) *Missing data model*: Suppose that we only get to observe a fraction of the entries of $\mathbf{M}^*$; more precisely, each entry of $\mathbf{M}^*$ is observed independently with probability p . In this setting, we can take the observed data matrix via zero-padding and rescaling as follows

$$M_{ij} = \begin{cases} \frac{1}{p}(M_{ij}^* + \sqrt{p}\tilde{H}_{ij}), & \text{with probability } p, \\ 0, & \text{otherwise.} \end{cases} \quad (54)$$

This way we guarantee that $\mathbb{E}[\mathbf{M}] = \mathbf{M}^*$, and $\tilde{\mathbf{H}} = [\tilde{H}_{ij}]_{1 \leq i, j \leq n}$ is a heteroscedastic Gaussian noise with variance $\text{Var}(\tilde{H}_{ij}) := [\text{Var}(\tilde{H}_{ij})]_{1 \leq i, j \leq n}$ equal to (53). In this case, the matrix $\mathbf{H} := \mathbf{M} - \mathbf{M}^*$ obeys

$$|H_{ij}| \leq \left| \frac{1}{p} M_{ij}^* \right| + \left| \frac{1}{\sqrt{p}} \tilde{H}_{ij} \right| \lesssim \frac{\mu}{np} + \frac{\tilde{\sigma}_{\max} \sqrt{\log n}}{\sqrt{p}},$$

$$\mathbb{E}[H_{ij}^2] \lesssim \frac{\mu^2}{n^2 p} + \tilde{\sigma}_{\max}^2$$

with high probability, where $\tilde{\sigma}_{\max} := \sigma_1 + (n-1)\delta_\sigma$. When the sampling rate exceeds $p \geq \frac{c_p \mu \log^2 n}{n}$ and the noise size is below $\tilde{\sigma}_{\max} \leq \frac{1}{\sqrt{c_p n \log n}}$ for some constant $c_p \geq 1$, one has

$$|H_{ij}| \lesssim \frac{\mu}{np} + \frac{\tilde{\sigma}_{\max} \sqrt{\log n}}{\sqrt{p}} \lesssim \frac{1}{c_p \log n},$$

$$\sqrt{\mathbb{E}[H_{ij}^2]} \lesssim \frac{\sqrt{\mu}}{n\sqrt{p}} + \tilde{\sigma}_{\max} \lesssim \frac{1}{\sqrt{c_p n \log n}},$$

thus satisfying Assumption 2 for c_p sufficiently large. In the numerical experiments here, we shall choose $p = 0.1$ and

$$\sigma_1 = 1/\sqrt{10 n \log n}, \quad \delta_\sigma = 9/((n-1)\sqrt{10 n \log n}).$$

d) *Estimation error vs. size of the ground-truth*: To study the qualitative effect of the magnitude of the group-truth, various values of $|\mathbf{a}^\top \mathbf{u}_2^*|$ are considered. For each configuration, the dimension n is set to be 500 and we run 100 trials for the box plot. The experiments are conducted for a variety of eigen-gaps and directions $\mathbf{a}$. A clear positive correspondence is seen between the size $|\mathbf{a}^\top \mathbf{u}_2^*|$ and the estimation error; see Fig. 2.

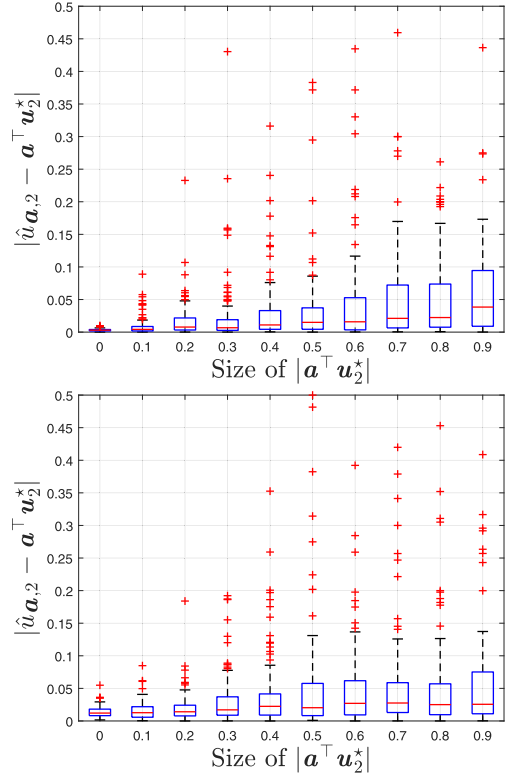


Fig. 2. Estimation error vs. size of the ground-truth $|\mathbf{a}^\top \mathbf{u}_2^*|$. The left figure corresponds to the noisy observation case where $|\Delta_2^*| = 0.01$ and the right figure corresponds to the missing observation case where $|\Delta_2^*| = 0.05$. In both cases, $\mathbf{a}$ is chosen such that $\mathbf{a}^\top \mathbf{u}_1^* = 0$.

e) *Estimation error vs size of the interferer*: To study the qualitative effect of the magnitude of the interferers, we consider a range of values for $|\mathbf{a}^\top \mathbf{u}_1^*|$ while holding $|\mathbf{a}^\top \mathbf{u}_2^*| = 0.5$ unchanged. Again, for each configuration, the dimension n is set to be 500 and we run 100 trials for the box plot. The experiments are run with various values of eigen-gaps. A negative dependency of the estimation error on the interferer is observed. In particular, Fig. 3 illustrates how the estimation errors grow as the interferer gets stronger.

f) *Estimation error vs eigen-gap*: In this part, we consider various magnitudes of the eigen-gap and study how the estimation error is influenced. Similar to what Theorem 1 and Theorem 3 predict, the estimation task becomes easier when the eigen-gap gets larger. Fig. 4 manifests this relationship in the case when $|\mathbf{a}^\top \mathbf{u}_1^*| = 0.5$ and $|\mathbf{a}^\top \mathbf{u}_2^*| = 0.2$.

C. Uncertainty Quantification

This subsection performs numerical experiments to validate our inferential procedures as well as the accompanying theory. Throughout this subsection, we consider the rank-2 case again as in the expression (11), where $\mathbf{M}^* := \lambda_1^* \mathbf{u}_1^* \mathbf{u}_1^{*\top} + \lambda_2^* \mathbf{u}_2^* \mathbf{u}_2^{*\top}$. The ground truth $\mathbf{M}^*$ is produced by setting $\lambda_1^* = 1$ with the corresponding eigenvectors $\mathbf{u}_1^*$ and $\mathbf{u}_2^*$ randomly generated. We aim to construct 95% confidence intervals (namely, $\alpha = 0.05$) for

1. the linear form $\mathbf{a}^\top \mathbf{u}_2^*$ for a given vector $\mathbf{a}$;
2. the eigenvalue λ_2^* .

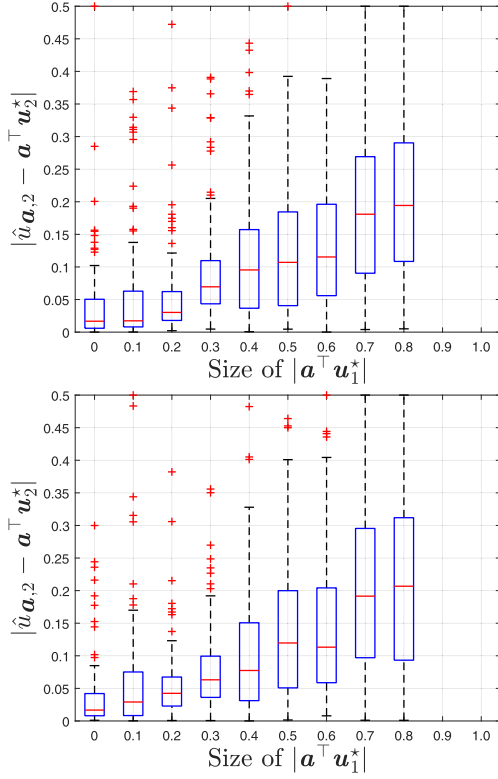


Fig. 3. Estimation error vs. size of the interferer $|\mathbf{a}^\top \mathbf{u}_1^*|$. The left figure corresponds to the full observation case where $|\Delta_2^*| = 0.01$, while the right figure corresponds to the missing observation case where $|\Delta_2^*| = 0.05$.

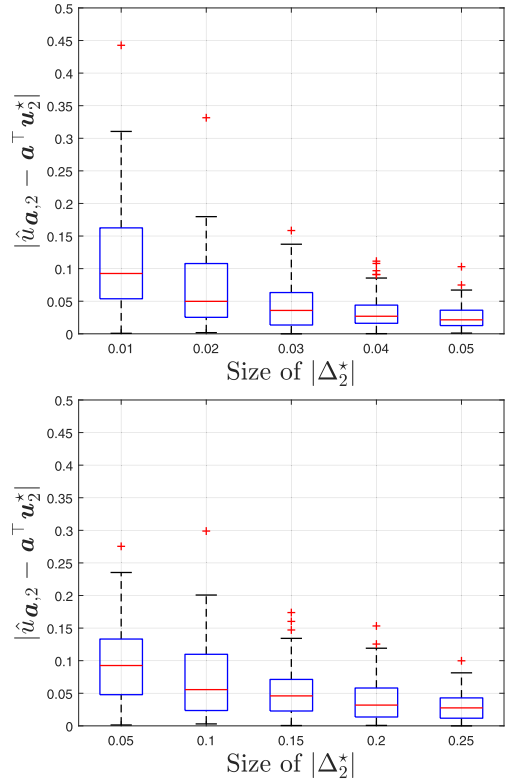


Fig. 4. Estimation error vs eigen-gap $|\Delta_2^*|$. The left figure corresponds to the full observation case, while the right figure corresponds to the missing observation case.

The unit vector $\mathbf{a}$ is chosen such that $\mathbf{a}^\top \mathbf{u}_1^* = 0$ and $\mathbf{a}^\top \mathbf{u}_2^* = 0.5$. Given that $|\mathbf{a}^\top \mathbf{u}_2^*|$ is quite large, we are expected to have $|\mathbf{a}^\top \mathbf{u}_2| \geq c_b \sqrt{\hat{v}_{a,2}} \log^{1.5} n$ with high probability, and hence we can simply take $\hat{v}_{a,2}^{\text{modified}} = \hat{u}_{a,2}$ in these numerical experiments (according to (29) and (11)).

1) *Heteroscedastic Gaussian Noise*: The noise setting is the same as in expression (53). The numerical results are displayed in Fig. 5 and Tab. I. We also examine the necessity of our requirement on the “interferers” (i.e. $|\mathbf{a}^\top \mathbf{u}_k^*| \lesssim |\lambda_l^* - \lambda_k^*| / (|\lambda_l^*| \log n)$, $k \neq l$) as stated in Theorem 4. Specifically, we make comparisons of the following two settings: (1) the “no interferer” case where $\mathbf{a}^\top \mathbf{u}_1^* = 0$ and $\mathbf{a}^\top \mathbf{u}_2^* = 0.5$, as plotted in Fig. 5(a)-5(f); and (2) the case with a strong interferer where $\mathbf{a}^\top \mathbf{u}_1^* = 0.05$ and $\mathbf{a}^\top \mathbf{u}_2^* = 0.5$, as plotted in Fig. 5(g)-5(i). Numerically, our distributional guarantees are inaccurate when there exists a strong interferer, which is consistent with what our theory predicts.

2) *Heteroscedastic Bernoulli Noise*: Consider a heteroscedastic Bernoulli noise matrix $\mathbf{H}$ with independent entries such that $H_{ij} = -\sigma_{ij}$ with probability 1/2 and $H_{ij} = \sigma_{ij}$ otherwise. The variance matrix is also chosen to satisfy (53). The numerical results are plotted in Fig. 6 and Tab. I.

3) *Missing Data Model*: In the case when we only get to observe a fraction of the entries of $\mathbf{M}^*$ as in model (54), in the same way, we aim to provide confidence intervals with precise coverages for both the linear functionals and the eigenvalues. Our numerical results are shown in Fig. 7 and Tab. I.

TABLE I

NUMERICAL COVERAGE RATES FOR OUR 95% CONFIDENCE INTERVALS OVER 10000 INDEPENDENT TRIALS

Settings	Target	Numerical coverage
heteroscedastic Gaussian	linear form $\mathbf{a}^\top \mathbf{u}_2^*$	0.9422
	eigenvalue λ_2^*	0.9496
heteroscedastic Bernoulli	linear form $\mathbf{a}^\top \mathbf{u}_2^*$	0.9470
	eigenvalue λ_2^*	0.9494
missing data	linear form $\mathbf{a}^\top \mathbf{u}_2^*$	0.9398
	eigenvalue λ_2^*	0.9519

4) *Conclusion*: In all of the above numerical experiments, the confidence intervals and the Q-Q (quantile-quantile) plots we produce match the theoretical predictions in a reasonably well manner, thus corroborating the validity and practicability of our theoretical results. We have numerically verified the need of controlling the “interferers” in Fig. 5(g)-5(i).

VI. PRIOR ART

Recent years have witnessed a flurry of activity in noisy low-rank matrix estimation [2], [13], [21], [33]–[44]. Despite the fundamental importance of estimating linear functionals of eigenvectors (or singular vectors), how to accomplish this task remains largely unknown. Only until very recently, researchers started to understand estimation errors for a special type of linear functionals, namely, the entrywise estimation error for the leading eigenvector or the $\ell_{2,\infty}$ error for the rank- r eigenspace. Examples include [9], [17], [20], [33], [45]–[52], which have been motivated by various applications

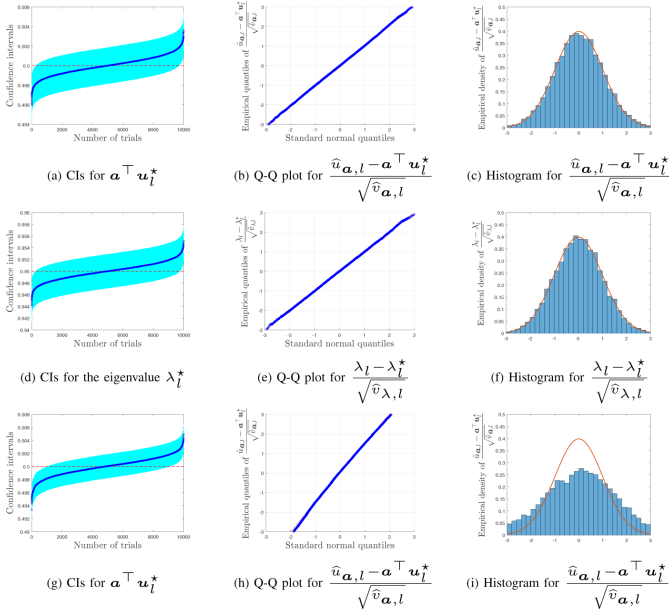


Fig. 5. Numerical results for inference for the linear form $\mathbf{a}^\top \mathbf{u}_l^*$ and the eigenvalue $\lambda_l^* = 0.95$ ($l = 2$) under heteroscedastic Gaussian noise. In (a)–(f), we take $\mathbf{a}^\top \mathbf{u}_1^* = 0$ (no interferer), while in (g), (h) and (i) $\mathbf{a}^\top \mathbf{u}_1^* = 0.05$ (with interferer). In both two settings, we take $n = 1000$, set $\sigma_1 = 0.1/\sqrt{n \log n}$ and $\delta_\sigma = 0.4/((n-1)\sqrt{n \log n})$ (cf. (53)), and run independent 10000 trials. In (a), (d) and (g), the confidence intervals are sorted respectively by the magnitudes of the estimators $\hat{u}_{a,l}$ and λ_l in these trials. Here, $\mathbf{u}_2$ is chosen such that $\mathbf{u}_2^\top \mathbf{u}_2^* \geq 0$. In (c), (f) and (i), the empirical densities are compared to the pdf of the standard normal (red curve).

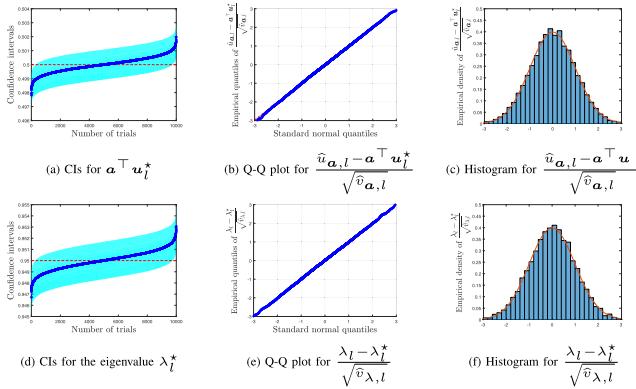


Fig. 6. Numerical results for inference for the linear form $\mathbf{a}^\top \mathbf{u}_l^*$ and the eigenvalue $\lambda_l^* = 0.95$ ($l = 2$) under heteroscedastic Bernoulli noise. We take $n = 1000$, set $\sigma_1 = 0.1/\sqrt{n \log n}$ and $\delta_\sigma = 0.4/((n-1)\sqrt{n \log n})$ (cf. (53)), and run independent 10000 trials. In (a) and (d), the confidence intervals are sorted respectively by the magnitudes of the estimators $\hat{u}_{a,l}$ and λ_l in these trials. Here, $\mathbf{u}_2$ is chosen such that $\mathbf{u}_2^\top \mathbf{u}_2^* \geq 0$. In (c) and (f), the empirical densities are compared to the pdf of the standard normal (red curve).

including matrix completion, community detection, top- K ranking, and so on. While tight ℓ_∞ eigenvector (resp. $\ell_{2,\infty}$ eigenspace) perturbation bounds in the presence of symmetric noise matrices have been derived in the prior works [9], [20], [52], all of these papers required the associated eigen-gap to exceed the spectral norm of the noise matrix, thereby falling short of addressing the scenarios with small eigen-gaps. The

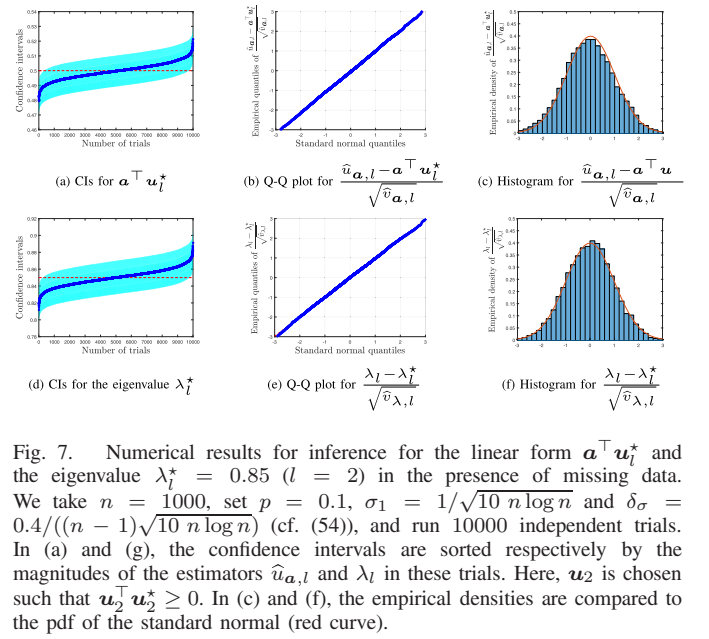


Fig. 7. Numerical results for inference for the linear form $\mathbf{a}^\top \mathbf{u}_l^*$ and the eigenvalue $\lambda_l^* = 0.85$ ($l = 2$) in the presence of missing data. We take $n = 1000$, set $p = 0.1$, $\sigma_1 = 1/\sqrt{10 n \log n}$ and $\delta_\sigma = 0.4/((n-1)\sqrt{10 n \log n})$ (cf. (54)), and run 10000 independent trials. In (a) and (g), the confidence intervals are sorted respectively by the magnitudes of the estimators $\hat{u}_{a,l}$ and λ_l in these trials. Here, $\mathbf{u}_2$ is chosen such that $\mathbf{u}_2^\top \mathbf{u}_2^* \geq 0$. In (c) and (f), the empirical densities are compared to the pdf of the standard normal (red curve).

work [33] also considered controlling the perturbation error of certain Fourier coefficients of the leading singular vector in a blind deconvolution problem, which, however, does not generalize to other linear functionals. [53] developed sharp concentration bounds for estimating linear forms of singular vectors under i.i.d. Gaussian noise, which, however, required the true singular values to be sufficiently separated (i.e. with a spacing much larger than the minimal eigen-gap studied herein). The recent work [28] is perhaps the only one that studied finite-grained eigenvector perturbation in the face of a small eigen-gap, which, however, is restricted to the case with i.i.d. Gaussian noise.

Given that estimation of linear functionals of eigenvectors is already largely under-explored, it is perhaps not surprising to see the lack of investigation about inference and uncertainty quantification for these quantities. This is in stark contrast to sparse estimation and learning problems, for which the construction of confidence regions has been extensively studied [54]–[64]. A few exceptions are worth mentioning: (1) [65]–[67] identified ℓ_2 confidence regions that are likely to cover the low-rank matrix of interest, which, however, might be loose in terms of the pre-constant; (2) focusing on low-rank matrix completion, the recent work [68] developed a de-biasing strategy that constructs both confidence regions for low-rank factors and entrywise confidence intervals for the unknown matrix, attaining statistical optimality in terms of both the pre-constant and the rate; an independent work by Xia et al. [69] analyzed a similar de-biasing strategy with the aid of double sample splitting, and shows asymptotic normality of linear forms of the matrix estimator; (3) [70], [71] developed a spectral projector to construct confidence regions for singular subspaces in the presence of i.i.d. additive noise; (4) [18] considered estimating linear forms of eigenvectors in a different covariance estimation model, whose analysis relies on the Gaussianity assumption; (5) [72] characterized the asymptotic normality of bilinear forms of eigenvectors, which

accommodates heterogeneous noise; and (6) [44] established the $\ell_{2,\infty}$ distributional guarantees for two spectral estimators (i.e., plain SVD and heteroskedastic PCA) tailored to PCA with heteroskedastic and missing data.

Additionally, the bulk distribution of the eigenvalues of i.i.d. random matrices has been studied in the physics literature (e.g. [73]–[78]), which falls short of the characterizing distributions of extreme eigenvalues. Several more recent papers started to consider a super-position of a low-rank matrix and an i.i.d. noise matrix, and studied the locations and distributions of the eigenvalues beyond the bulk [79]–[81]. These results did not cover a general class of heteroskedastic noise and did not allow the rank r to grow with n . What is more, none of these papers studied how to construct valid confidence intervals for extreme eigenvalues, let alone inference for individual eigenvectors.

VII. DISCUSSION

The present paper contributes towards “fine-grained” statistical analysis, by developing guaranteed estimation and inference algorithms for linear functionals of the unknown eigenvectors and eigenvalues. The proposed procedures are model agnostic and are able to accommodate heteroskedastic noise, without the need of prior knowledge about the noise levels. The validity of our procedures is guaranteed even when the eigen-gap is extremely small, a condition that goes significantly beyond what we have learned from generic matrix perturbation theory. The key enabler of our findings lies in an appealing bias reduction feature of eigen-decomposition when coping with asymmetric noise matrices.

Our studies leave open several interesting questions worthy of future investigation. For instance, our current theory for confidence intervals falls short of accommodating the scenarios when quantity $|\mathbf{a}^\top \mathbf{u}_l^*|$ far exceeds the associated eigen-gap — how to determine the fundamental inference limits for such scenarios and, perhaps more importantly, how to attain the limits efficiently? In addition, our theory is likely suboptimal in terms of the dependency on the rank r and the condition number κ . Can we further improve the theoretical support in these regards? Furthermore, our analysis framework shed some light on how to perform inference on functions of the eigenvectors. It would be interesting to develop a unified framework that leads to valid confidence intervals for a broader class of functions (e.g. quadratic functionals, or more general polynomials) of the eigenvectors. Moving beyond estimation and inference for individual eigenvectors, we remark that our current theory falls short of delivering useful eigenspace perturbation guarantees unless there exists a sufficient eigen-gap between any adjacent pair of eigenvalues. The challenges are at least two-fold when dealing with an asymmetric data matrix: (1) the eigenvectors are, in general, not orthogonal to each other, and (2) the eigenvectors might be complex-valued even when the observed data are real-valued, both of which are in stark contrast to what happens for a symmetric data matrix. How to extend our theory to accommodate more general eigenspace perturbation is left for future investigation.

APPENDIX A PRELIMINARIES

Denote by $\mathbf{u}_l$ (resp. λ_l) the l th leading right eigenvector (resp. eigenvalue) of $\mathbf{M} = \sum_{j=1}^r \lambda_j^* \mathbf{u}_j^* \mathbf{u}_j^\top + \mathbf{H}$. We make note of several facts about $\mathbf{u}_l$ and λ_l — previously established in [17] — that will prove useful throughout.

To begin with, a simple application of the Neumann series yields the following expansion [17, Theorem 2].

Lemma 1 (Neumann Expansion): Let $\mathbf{u}_l$ and λ_l be the l th leading right eigenvector and the l th leading eigenvalue of $\mathbf{M}$ (cf. (7)), respectively. If $\|\mathbf{H}\| < |\lambda_l^*|$, then one has

$$\mathbf{u}_l = \sum_{j=1}^r \frac{\lambda_j^*}{\lambda_l} (\mathbf{u}_j^* \mathbf{u}_l) \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \mathbf{u}_j^* \right\}. \quad (55)$$

As an immediate consequence, we have the following expansion for $\mathbf{a}^\top \mathbf{u}_l$, which forms the basis of our estimators:

$$\begin{aligned} & \frac{\lambda_l}{\lambda_l^* (\mathbf{u}_l^* \mathbf{u}_l)} \mathbf{a}^\top \mathbf{u}_l \\ &= \sum_{j=1}^r \frac{\lambda_j^*}{\lambda_l^*} \cdot \frac{\mathbf{u}_j^* \mathbf{u}_l}{\mathbf{u}_l^* \mathbf{u}_l} \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_j^* \right\} \\ &= \mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{\lambda_l} \mathbf{a}^\top \mathbf{H} \mathbf{u}_l^* + \sum_{s=2}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_l^* \\ &+ \sum_{j:j \neq l} \frac{\lambda_j^*}{\lambda_l^*} \frac{\mathbf{u}_j^* \mathbf{u}_l}{\mathbf{u}_l^* \mathbf{u}_l} \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_j^* \right\}. \end{aligned} \quad (56)$$

As we shall see shortly, our theory relies heavily on approximating (56) by the lower-order terms (w.r.t. $\mathbf{H}$). This requires controlling the influence of the high-order terms. Towards this end, we make note of several useful bounds about $\mathbf{H}$ that have been established in [17].

Lemma 2: Fix any vector $\mathbf{a} \in \mathbb{R}^n$, and suppose that $\|\mathbf{u}_l^*\|_\infty \leq \sqrt{\mu/n}$. Suppose the noise matrix $\mathbf{H}$ obeys Assumption 1, and assume the existence of some sufficiently small constant $c_1 > 0$ such that

$$\max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \} \leq c_1 \lambda_{\min}^*. \quad (57)$$

Then there exist some universal constants $c_2, c_3 > 0$ such that with probability at least $1 - O(n^{-10})$,

$$|\mathbf{a}^\top \mathbf{H}^s \mathbf{u}_l^*| \leq \sqrt{\frac{\mu}{n}} \left(c_2 \max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \} \right)^s \|\mathbf{a}\|_2, \quad (58a)$$

$$\|\mathbf{H}\| \leq c_3 \max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \}. \quad (58b)$$

Remark 8: The work [17] established the bound (58a) only for the case with $s \leq 20 \log n$. Fortunately, the case with $s > 20 \log n$ follows immediately by combining the crude bound $|\mathbf{a}^\top \mathbf{H}^s \mathbf{u}^*| \leq \|\mathbf{H}\|^s \|\mathbf{a}\|_2$ and the inequality (58b) (by choosing $c_2 = 2c_3$ and using the fact that $(1/2)^s \ll \sqrt{1/n}$).

Combining Lemmas 1-2, the paper [17] establishes the following result.

Lemma 3: Consider a rank- r symmetric matrix $\mathbf{M}^* = \sum_{i=1}^r \lambda_i^* \mathbf{u}_i^* \mathbf{u}_i^\top \in \mathbb{R}^{n \times n}$ with incoherence parameter μ (cf. Definition 1). Suppose that the noise matrix $\mathbf{H}$ obeys Assumption 1 and Condition (57). Then for any fixed vector

$\mathbf{a} \in \mathbb{R}^n$ and any $1 \leq i \leq r$, with probability at least $1 - O(n^{-10})$ one has

$$\left| \mathbf{a}^\top \left(\mathbf{u}_i - \sum_{k=1}^r \frac{\mathbf{u}_k^\top \mathbf{u}_i}{\lambda_i / \lambda_k^*} \mathbf{u}_k^* \right) \right| \leq c_2 \frac{\max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \}}{\lambda_{\min}^*} \sqrt{\frac{\mu \kappa^2 r}{n}} \|\mathbf{a}\|_2 \quad (59)$$

for some universal constant $c_2 > 0$. This result holds unchanged if the i th right eigenvector $\mathbf{u}_i$ is replaced by the i th left eigenvector $\mathbf{w}_i$.

Remark 9: Under the additional condition that $B \log n \leq \sigma_{\max} \sqrt{n \log n}$, the preceding bounds simplify to

$$|\mathbf{a}^\top \mathbf{H}^s \mathbf{u}_i^*| \leq \sqrt{\frac{\mu}{n}} \left(c_2 \sigma_{\max} \sqrt{n \log n} \right)^s \|\mathbf{a}\|_2, \quad (60a)$$

$$\|\mathbf{H}\| \leq c_3 \sigma_{\max} \sqrt{n \log n}, \quad (60b)$$

$$\left| \mathbf{a}^\top \left(\mathbf{u}_i - \sum_{k=1}^r \frac{\mathbf{u}_k^\top \mathbf{u}_i}{\lambda_i / \lambda_k^*} \mathbf{u}_k^* \right) \right| \leq c_2 \frac{\sigma_{\max} \sqrt{\mu \kappa^2 r \log n}}{\lambda_{\min}^*} \|\mathbf{a}\|_2. \quad (60c)$$

Another immediate consequence under our assumptions is that, with probability at least $1 - O(n^{-10})$,

$$\|\mathbf{H}\| \leq \lambda_{\min}^* / 10. \quad (61)$$

APPENDIX B

PROOF FOR EIGENVECTOR AND EIGENVALUE ESTIMATION

In this section, we shall start by proving the eigenvalue perturbation bound stated in Theorem 2, which plays a pivotal role in establishing the eigenvector estimation guarantees in Theorem 1.

A. Proof of Theorem 2

This section aims to establish a slightly stronger version of Theorem 2, stated as follows.

Theorem 7: Assume that $\mu \kappa^2 r^4 \leq c_3 n$ for some sufficiently small constant $c_3 > 0$. Suppose that the noise parameters defined in Assumption 1 satisfy

$$\Delta_l^* > 2c_4 \kappa^2 r^2 \max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \} \sqrt{\frac{\mu}{n}} \quad (62a)$$

$$\max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \} \leq c_5 \lambda_{\min}^* \quad (62b)$$

for some sufficiently large (resp. small) constant $c_4 > 0$ (resp. $c_5 > 0$). Then given any integer $1 \leq l \leq r$, with probability $1 - O(n^{-8})$, the eigenvalue λ_l and the associated eigenvectors $\mathbf{u}_l$ and $\mathbf{w}_l$ (see Notation 1) are all real-valued, and one has

$$|\lambda_l - \lambda_l^*| \leq c_4 \kappa r^2 \max \{ \sigma_{\max} \sqrt{n \log n}, B \log n \} \sqrt{\frac{\mu}{n}}. \quad (63)$$

A few immediate consequences of this theorem and the Bauer-Fike theorem are summarized as follows.

Corollary 2: Suppose that $\mu \kappa^2 r^4 \lesssim n$, and that Assumptions 1-2 hold. With probability at least $1 - O(n^{-6})$,

$$|\lambda_l - \lambda_l^*| \leq \min \left\{ \frac{\Delta_l^*}{2}, c_3 \sigma_{\max} \sqrt{n \log n} \right\} \quad \text{and} \quad \left| \frac{\lambda_l - \lambda_l^*}{\lambda_l^*} \right| \leq \frac{1}{100} \quad (64)$$

for some sufficiently small constant $c_3 > 0$. In addition, if either $\sigma_{\max} \sqrt{n \log n} = o(\lambda_{\min}^*)$ or $\Delta_l^* = o(\lambda_{\min}^*)$ holds, then with probability at least $1 - O(n^{-6})$,

$$\lambda_l = (1 + o(1)) \lambda_l^*. \quad (65)$$

1) Proof Outline: To establish this theorem, we borrow a powerful idea from [79] that converts eigenvalue analysis to zero counting of certain complex-valued functions. Specifically, consider the following functions

$$f(z) := \det(\mathbf{I} + \mathbf{U}^{\star\top} (\mathbf{H} - z\mathbf{I})^{-1} \mathbf{U}^* \mathbf{\Sigma}^*), \quad (66a)$$

$$g(z) := \det(\mathbf{I} + \mathbf{U}^{\star\top} (-z)^{-1} \mathbf{U}^* \mathbf{\Sigma}^*). \quad (66b)$$

The intimate connection between these two functions and the eigenvalues of $\mathbf{M}$ and $\mathbf{M}^*$ is formalized in the following observation made by [79].

Claim 1: If $\lambda_{\min} > 2\|\mathbf{H}\|$, then the zeros of $f(\cdot)$ (resp. $g(\cdot)$) on the region $\mathcal{K} := \{z \in \mathbb{C} : |z| > \|\mathbf{H}\|\} \cup \{\infty\}$ are exactly the r leading eigenvalues of $\mathbf{M}$ (resp. $\mathbf{M}^*$).

With this claim in mind, we turn attention to studying the zeros of $f(\cdot)$ and $g(\cdot)$. Given that $f(\cdot)$ can be viewed as a perturbed version of $g(\cdot)$ (since $f(\cdot)$ can be obtained by adding a perturbation $\mathbf{H}$ to $g(\cdot)$ in a certain way), we hope that the zeros of $f(\cdot)$ do not deviate by much from the zeros of $g(\cdot)$. Towards justifying this, we look at the following γ -neighborhood of $\{\lambda_1^*, \lambda_2^*, \dots, \lambda_r^*\}$:

$$\mathcal{D}(\gamma) := \bigcup_{k=1}^r \mathcal{B}(\lambda_k^*, \gamma), \quad (67)$$

where $\mathcal{B}(\lambda, \gamma)$ is a ball of radius $\gamma > 0$ centered at λ , namely, $\mathcal{B}(\lambda, \gamma) := \{z \in \mathbb{C} : |z - \lambda| \leq \gamma\}$. A crucial part of the proof is to demonstrate that all zeros of $f(\cdot)$ lie in $\mathcal{D}(\gamma)$ — or equivalently, in the γ -neighborhood of the zeros of $g(\cdot)$ — for sufficiently small γ .

Remark 10: Somewhat surprisingly, γ is allowed to be as small as $\frac{\text{poly} \log(n)}{\sqrt{n}} \|\mathbf{H}\|$ when $r, \kappa, \mu \asymp 1$.

In what follows, we shall assume that

$$\gamma < \lambda_{\min}^* / 4 \quad \text{and} \quad \gamma < \Delta_l^* / 2. \quad (68)$$

In view of (61), one has $\|\mathbf{H}\| < \lambda_{\min}^* / 4$, thus indicating that

$$|z| \geq \lambda_{\min}^* - \gamma > \lambda_{\min}^* / 2 > \|\mathbf{H}\| \quad \text{for all } z \in \mathcal{D}(\gamma). \quad (69)$$

Our proof is based on the following observations:

- (i) Given that $\gamma < \Delta_l^* / 2$, one has $\mathcal{B}(\lambda_l^*, \gamma) \cap \mathcal{B}(\lambda_k^*, \gamma) = \emptyset$ for any $k \neq l$, and hence $g(\cdot)$ has exactly 1 zero in $\mathcal{B}(\lambda_l^*, \gamma)$; in addition, it is clear that $g(\cdot)$ has $l - 1$ (resp. $r - l$) zeros in $\bigcup_{k:k < l} \mathcal{B}(\lambda_k^*, \gamma)$ (resp. $\bigcup_{k:k > l} \mathcal{B}(\lambda_k^*, \gamma)$).
- (ii) Suppose that in each connected component of $\mathcal{D}(\gamma)$, the functions $f(\cdot)$ and $g(\cdot)$ always have the same number

- of zeros. If this were true, then one would have (1) a unique zero of $f(\cdot)$ in $\mathcal{B}(\lambda_l^*, \gamma)$; (2) $l-1$ zeros of $f(\cdot)$ in $\cup_{1 \leq k < l} \mathcal{B}(\lambda_k^*, \gamma)$; (3) $r-l$ zeros of $f(\cdot)$ in $\cup_{k > l} \mathcal{B}(\lambda_k^*, \gamma)$.
- (iii) Recalling how we sort $\{\lambda_k^*\}$ and $\{\lambda_k\}$, we see that the real part of any $z \in \cup_{1 \leq k < l} \mathcal{B}(\lambda_k^*, \gamma)$ must exceed $\lambda_{l-1}^* - \gamma > \lambda_l^* + \gamma$ and, similarly, the real part of any $z \in \cup_{k > l} \mathcal{B}(\lambda_k^*, \gamma)$ is at most $\lambda_{l+1}^* + \gamma < \lambda_l^* - \gamma$. Thus, the above arguments reveal that the unique zero of $f(\cdot)$ in $\mathcal{B}(\lambda_l^*, \gamma)$ is λ_l , which obeys $|\lambda_l - \lambda_l^*| \leq \gamma$.
- (iv) In addition, note that the complex conjugate $\overline{\lambda_l}$ of λ_l is also an eigenvalue of M , given that M is a real-valued matrix. However, since there is only one zero of $f(\cdot)$ residing in $\mathcal{B}(\lambda_l^*, \gamma)$, we necessarily have $\overline{\lambda_l} = \lambda_l$, meaning that λ_l is real-valued. A similar argument justifies that the associated right eigenvector u_l and left eigenvector w_l are real-valued as well.

Theorem 7 is thus established if we are allowed to pick

$$\gamma = c_4 \kappa r^2 \max \left\{ B \log n, \sigma_{\max} \sqrt{n \log n} \right\} \sqrt{\frac{\mu}{n}}. \quad (70)$$

Clearly, this choice satisfies $\gamma < \lambda_{\min}^*/4$ and $\gamma < \Delta_l^*/2$ under the assumptions of this theorem.

The remaining proof then boils down to justifying that $f(\cdot)$ and $g(\cdot)$ have the same number of zeros in each connected component of $\mathcal{D}(\gamma)$. Towards this end, we resort to Rouché's theorem in complex analysis [82].

Theorem 8 (Rouché's Theorem): Let $f(\cdot)$ and $g(\cdot)$ be two complex-valued functions that are holomorphic inside a region $\mathcal{R}$ with closed contour $\partial\mathcal{R}$. If $|f(z) - g(z)| < |g(z)|$ for all $z \in \partial\mathcal{R}$, then $f(\cdot)$ and $g(\cdot)$ have the same number of zeros inside $\mathcal{R}$.

In order to invoke Rouché's theorem to justify the claim in (ii), it suffices to fulfill the requirement $|f(z) - g(z)| < |g(z)|$ on $\mathcal{D}(\gamma)$. Given that $|z| > \|\mathbf{H}\|$ for all $z \in \mathcal{D}(\gamma)$ (see (69)), we can apply the Neumann series $(\mathbf{H} - z\mathbf{I})^{-1} = -\sum_{s=0}^{\infty} z^{-s-1} \mathbf{H}^s$ to reach

$$\begin{aligned} f(z) &= \det \left(\mathbf{I} - \mathbf{U}^{\star\top} \left(\sum_{s=0}^{\infty} z^{-s-1} \mathbf{H}^s \right) \mathbf{U}^{\star} \Sigma^{\star} \right) \\ &= \det \left(\mathbf{I} + \mathbf{U}^{\star\top} (-z)^{-1} \mathbf{U}^{\star} \Sigma^{\star} - \sum_{s=1}^{\infty} z^{-s-1} \mathbf{U}^{\star\top} \mathbf{H}^s \mathbf{U}^{\star} \Sigma^{\star} \right) \\ &= \det \left((\mathbf{I} + \mathbf{U}^{\star\top} (-z)^{-1} \mathbf{U}^{\star} \Sigma^{\star}) \left(\mathbf{I} - (\mathbf{I} + \mathbf{U}^{\star\top} (-z)^{-1} \mathbf{U}^{\star} \Sigma^{\star})^{-1} \sum_{s=1}^{\infty} z^{-s-1} \mathbf{U}^{\star\top} \mathbf{H}^s \mathbf{U}^{\star} \Sigma^{\star} \right) \right) \\ &= g(z) \det \left(\mathbf{I} - (\mathbf{I} + \mathbf{U}^{\star\top} (-z)^{-1} \mathbf{U}^{\star} \Sigma^{\star})^{-1} \sum_{s=1}^{\infty} z^{-s-1} \mathbf{U}^{\star\top} \mathbf{H}^s \mathbf{U}^{\star} \Sigma^{\star} \right) \\ &= g(z) \det \left(\mathbf{I} - \underbrace{(\mathbf{I} - \Sigma^{\star}/z)^{-1} \sum_{s=1}^{\infty} z^{-s-1} \mathbf{U}^{\star\top} \mathbf{H}^s \mathbf{U}^{\star} \Sigma^{\star}}_{=: \Delta} \right), \end{aligned}$$

where the last line holds since $\mathbf{U}^{\star\top} \mathbf{U}^{\star} = \mathbf{I}$. In addition, observe that the zeros of $g(\cdot)$ are in the interior of $\mathcal{D}$ and that $g(z) \neq 0$ for all $z \in \partial\mathcal{D}(\gamma)$. Hence, on $\partial\mathcal{D}(\gamma)$ we have

$$\begin{aligned} &|f(z) - g(z)| < |g(z)| \\ \iff &|g(z)| \cdot |\det(\mathbf{I} - \Delta) - 1| < |g(z)| \\ \iff &|\det(\mathbf{I} - \Delta) - 1| < 1. \end{aligned} \quad (71)$$

To justify (71) on $\partial\mathcal{D}$, we make the following observations:

Claim 2: The condition (71) can be guaranteed as long as $\|\Delta\| < 1/(2r)$.

Claim 3: There exists a universal constant $c_3 > 0$ such that with probability $1 - O(n^{-8})$, one has

$$\|\Delta\| \leq c_3 \frac{\max \{B \log n, \sigma_{\max} \sqrt{n \log n}\}}{\gamma} \sqrt{\frac{\mu \kappa^2 r^2}{n}}. \quad (72)$$

In summary, in order to guarantee $|f(z) - g(z)| < |g(z)|$ on $\partial\mathcal{D}(\gamma)$, it suffices to ensure that $\|\Delta\| \leq 1/(2r)$ (by (71) and Claim 2), which would hold as long as we take $\gamma = 2c_3 \max \{B \log n, \sigma_{\max} \sqrt{n \log n}\} \sqrt{\frac{\mu \kappa^2 r^4}{n}}$ (by Claim 3). This in turn establishes Theorem 7 (and hence Theorem 2) as long as $c_4 \geq 2c_3$.

Finally, the proofs of the auxiliary claims are postponed to Appendix B-A2.

2) Proofs for Auxiliary Claims in Appendix B-A1:

a) *Proof of Claim 1:* Note that both $f(\cdot)$ and $g(\cdot)$ are holomorphic on $\mathcal{K}$. Making use of the identity $\det(\mathbf{I} + \mathbf{A}\mathbf{B}) = \det(\mathbf{I} + \mathbf{B}\mathbf{A})$, we have

$$\begin{aligned} f(z) &= \det(\mathbf{I} + \mathbf{U}^{\star} \Sigma^{\star} \mathbf{U}^{\star\top} (\mathbf{H} - z\mathbf{I})^{-1}) \\ &= \det(\mathbf{I} + \mathbf{M}^{\star} (\mathbf{H} - z\mathbf{I})^{-1}) \\ &= \det((\mathbf{H} - z\mathbf{I} + \mathbf{M}^{\star})(\mathbf{H} - z\mathbf{I})^{-1}) \\ &= \frac{\det(\mathbf{M} - z\mathbf{I})}{\det(\mathbf{H} - z\mathbf{I})} \end{aligned} \quad (73)$$

$$\text{and } g(z) = \det(\mathbf{I} + \mathbf{U}^{\star} \Sigma^{\star} \mathbf{U}^{\star\top} / (-z)) = \prod_{i=1}^r \left(1 - \frac{\lambda_i^{\star}}{z} \right). \quad (74)$$

These identities make clear that the zeros of $g(\cdot)$ (resp. $f(\cdot)$) on $\mathcal{K}$ are all eigenvalues of $\mathbf{M}^{\star}$ (resp. $\mathbf{M}$). In particular, the zeros of $g(\cdot)$ are precisely $\{\lambda_i^{\star}\}_{1 \leq i \leq r}$.

It remains to show that there are exactly r zeros of $f(\cdot)$ lying in $\mathcal{K}$, and that they are exactly $\lambda_1, \dots, \lambda_r$. This is equivalent to showing that the set of eigenvalues of $\mathbf{M}$ contained in the region $\mathcal{K}$ is $\{\lambda_i\}_{1 \leq i \leq r}$. Towards this, define $\mathbf{M}(t) = \mathbf{M}^{\star} + t\mathbf{H}$. From the Bauer-Fike theorem, all eigenvalues of $\mathbf{M}(t)$ lie in the set

$$\mathcal{B}(0, t\|\mathbf{H}\|) \cup \underbrace{\left\{ \bigcup_{k=1}^r \mathcal{B}(\lambda_k^{\star}, t\|\mathbf{H}\|) \right\}}_{=: \mathcal{D}(t\|\mathbf{H}\|)}. \quad (75)$$

Given that $\|\mathbf{H}\| < \lambda_{\min}/4$ (according to (61) under our assumptions), it is easily seen that $\mathcal{B}(0, t\|\mathbf{H}\|)$ does not intersect with $\mathcal{D}(t\|\mathbf{H}\|)$ for all $0 \leq t \leq 1$. Additionally, the set of the eigenvalues of $\mathbf{M}(t)$ depends continuously on t [83, Theorem 6], requiring $\mathbf{M}(t)$ to have the same number of

zeros in $\mathcal{B}(0, t\|\mathbf{H}\|)$ for all $0 \leq t \leq 1$. Hence, $\mathbf{M} = \mathbf{M}(1)$ has exactly r eigenvalues in $\mathcal{D}(\|\mathbf{H}\|)$ (since $\mathbf{M}^* = \mathbf{M}(0)$ has r eigenvalues in this region). These eigenvalues necessarily have magnitudes larger than any point in $\mathcal{B}(0, \|\mathbf{H}\|)$ (since $\|\mathbf{H}\| < \lambda_{\min}/4$), thus indicating that the eigenvalues of $\mathbf{M}$ in $\mathcal{K}$ are precisely $\lambda_1, \dots, \lambda_r$.

b) *Proof of Claim 2:* Denoting by $\mu_1, \dots, \mu_r \in \mathbb{C}$ the r eigenvalues of $\mathbf{I} - \Delta \in \mathbb{R}^{r \times r}$, one has

$$|\det(\mathbf{I} - \Delta) - 1| = \left| \prod_{i=1}^r \mu_i - 1 \right|.$$

By virtue of the elementary inequality⁴ $|\prod_{i=1}^r (1 + a_i) - 1| \leq \prod_{i=1}^r (1 + |a_i|) - 1$, we arrive at

$$|\det(\mathbf{I} - \Delta) - 1| \leq \prod_{i=1}^r (1 + |\mu_i - 1|) - 1 \leq (1 + \|\Delta\|)^r - 1,$$

where the last inequality follows from the Bauer-Fike theorem (which forces that $|\mu_i - 1| \leq \|\Delta\|$). This means that Condition (71) holds if $(1 + \|\Delta\|)^r < 2$; the latter condition is clearly guaranteed if $\|\Delta\| \leq 1/(2r)$.

c) *Proof of Claim 3:* Since $\gamma < \lambda_{\min}/4$, it always holds that $|z| > 3\lambda_{\min}/4$ for all $z \in \partial\mathcal{D}(\gamma)$. Hence

$$\begin{aligned} \|\Delta\| &= \left\| (z\mathbf{I} - \Sigma^*)^{-1} \sum_{s=1}^{\infty} z^{-s} \mathbf{U}^{*T} \mathbf{H}^s \mathbf{U}^* \Sigma^* \right\| \\ &\leq \left\| (z\mathbf{I} - \Sigma^*)^{-1} \right\| \cdot \sum_{s=1}^{\infty} \frac{\|\mathbf{U}^{*T} \mathbf{H}^s \mathbf{U}^* \Sigma^*\|}{|z|^s} \\ &\leq \|\Sigma^*\| \cdot \left\| (z\mathbf{I} - \Sigma^*)^{-1} \right\| \cdot \sum_{s=1}^{\infty} \frac{\|\mathbf{U}^{*T} \mathbf{H}^s \mathbf{U}^*\|}{|z|^s}. \end{aligned} \quad (76)$$

Recall that $\|\Sigma^*\| = \lambda_{\max}$. Also, on $\partial\mathcal{D}(\gamma)$ we have

$$\left\| (z\mathbf{I} - \Sigma^*)^{-1} \right\| \leq \max_{1 \leq k \leq r, z \in \partial\mathcal{D}(\gamma)} \left| \frac{1}{z - \lambda_k^*} \right| \leq \frac{1}{\gamma}. \quad (77)$$

Taken collectively, the above results yield

$$\begin{aligned} \|\Delta\| &\leq \frac{\lambda_{\max}}{\gamma} \sum_{s=1}^{\infty} \frac{\|\mathbf{U}^{*T} \mathbf{H}^s \mathbf{U}^*\|}{\left(\frac{3}{4}\lambda_{\min}\right)^s} \\ &\stackrel{(i)}{\leq} \frac{\lambda_{\max} r}{\gamma} \sum_{s=1}^{\infty} \frac{\max_{1 \leq i, j \leq r} |\mathbf{u}_i^{*T} \mathbf{H}^s \mathbf{u}_j^*|}{\left(\frac{3}{4}\lambda_{\min}\right)^s}, \end{aligned} \quad (78)$$

where (i) makes use of the elementary inequality $\|\mathbf{A}\| \leq r\|\mathbf{A}\|_{\infty}$ for any $\mathbf{A} \in \mathbb{R}^{r \times r}$. Invoke Lemma 2 and a union

⁴This holds since

$$\begin{aligned} \left| \prod_{i=1}^r (1 + a_i) - 1 \right| &= \left| \sum_{k=1}^r \sum_{1 \leq i_1, \dots, i_k \leq r} a_{i_1} \cdots a_{i_k} \right| \\ &\leq \sum_{k=1}^r \sum_{1 \leq i_1, \dots, i_k \leq r} |a_{i_1}| \cdots |a_{i_k}| = \prod_{i=1}^r (1 + |a_i|) - 1. \end{aligned}$$

bound to show that: with probability $1 - O(n^{-10}r^2)$,

$$\begin{aligned} \|\Delta\| &\leq \frac{\lambda_{\max} r}{\gamma} \sum_{s=1}^{\infty} \left(\frac{c_2 \max\{\sigma_{\max} \sqrt{n \log n}, B \log n\}}{\frac{3}{4}\lambda_{\min}} \right)^s \sqrt{\frac{\mu}{n}} \\ &\stackrel{(ii)}{\leq} \frac{\lambda_{\max} r}{\gamma} \cdot \frac{c_2 \max\{\sigma_{\max} \sqrt{n \log n}, B \log n\}}{\frac{3}{4}\lambda_{\min}} \cdot \sum_{s=1}^{\infty} \left(\frac{1}{2}\right)^s \sqrt{\frac{\mu}{n}} \\ &\leq \frac{4c_2 \max\{\sigma_{\max} \sqrt{n \log n}, B \log n\}}{3\gamma} \sqrt{\frac{\mu \kappa^2 r^2}{n}}, \end{aligned}$$

where (ii) holds as long as $\frac{c_2 \max\{B \log n, \sigma_{\max} \sqrt{n \log n}\}}{\frac{3}{4}\lambda_{\min}} \leq \frac{1}{2}$.

B. Proof of Theorem 1

As mentioned previously, Theorem 2 enables us to prove the statistical guarantees stated in Theorem 1. In the sequel, we shall first establish perturbation bounds for the vanilla eigenvector estimator (i.e. $\mathbf{u}_l$ and $\mathbf{w}_l$) as well as the estimator $\hat{\mathbf{u}}_l = \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l)$.

Theorem 9: Consider any $1 \leq l \leq r$. Suppose that $\mu \kappa^2 r^4 \lesssim n$ and that Assumptions 1-2 hold. Then with probability at least $1 - O(n^{-6})$, we have:

1) (ℓ_2 perturbation of the l th eigenvector)

$$\begin{aligned} \min \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_2 &\lesssim \frac{\kappa^2 \sigma_{\max} \sqrt{\mu r^2 \log n}}{\Delta_l^*} + \frac{\kappa^2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\max}^*} \end{aligned} \quad (79a)$$

$$\begin{aligned} |\mathbf{u}_k^{*T} \mathbf{u}_l| &\lesssim \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r \log n}, \quad k \neq l, \end{aligned} \quad (79b)$$

$$\begin{aligned} |\mathbf{u}_l^{*T} \mathbf{u}_l| &\geq 1 - O\left((\kappa^4 \sigma_{\max}^2 n \log n) \left\{ \frac{1}{(\Delta_l^*)^2} \frac{\mu r^2}{n} + \frac{1}{(\lambda_{\max}^*)^2} \right\}\right); \end{aligned} \quad (79c)$$

2) (perturbation of linear forms of the l th eigenvector) for any fixed vector $\mathbf{a}$ with $\|\mathbf{a}\|_2 = 1$,

$$\begin{aligned} \min |\mathbf{a}^T (\mathbf{u}_l \pm \mathbf{u}_l^*)| &\lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^2 r^4 \log n}}{\lambda_{\min}^*} + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \cdot \max_{k \neq l} \frac{|\mathbf{a}^T \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \\ &\quad + \frac{\sigma_{\max}^2 \mu \kappa^4 r^2 \log n}{(\Delta_l^*)^2} |\mathbf{a}^T \mathbf{u}_l^*| + \frac{\sigma_{\max}^2 \kappa^4 n \log n}{(\lambda_{\max}^*)^2} |\mathbf{a}^T \mathbf{u}_l^*|; \end{aligned} \quad (79d)$$

3) (ℓ_{∞} perturbation of the l th eigenvector)

$$\min \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_{\infty} \lesssim \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^4 r \log n} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}}. \quad (79e)$$

In addition, the above results hold unchanged if $\mathbf{u}_l$ is replaced by either $\mathbf{w}_l$ or $\hat{\mathbf{u}}_l = \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l)$.

Proof: See Appendix B-B1. $\square$

In words, Theorem 9 reveals appealing statistical performance for the estimators $\mathbf{u}_l$, $\mathbf{w}_l$ and $\hat{\mathbf{u}}_l$. However, when estimating linear functionals of eigenvectors via, say, the plug-in estimator $\mathbf{a}^T \mathbf{u}_l$, the bound (79d) suffers from an additional

term $\frac{\sigma_{\max}^2 \kappa^4 n \log n}{(\lambda_{\min}^*)^2} |\mathbf{a}^\top \mathbf{u}_l^*|$ compared to the desired bound (14c) in Theorem 1. As it turns out, this extra term arises due to a systematic bias of the plug-in estimates. To compensate for the bias, one needs to properly enlarge the plug-in estimator, thus leading us to the proposed estimators $\hat{u}_{a,l}$. We now establish the claimed performance guarantees for these properly corrected estimators for $\mathbf{a}^\top \mathbf{u}_l^*$; the proof is deferred to Appendix B-B2, which builds heavily upon the analysis of Theorem 9.

Theorem 10: Instate the assumptions of Theorem 9. Fix any vector $\mathbf{a}$ with $\|\mathbf{a}\|_2 = 1$. With probability exceeding $1 - O(n^{-6})$, the estimator $\hat{u}_{a,l} := \min \left\{ \sqrt{\left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} \right|}, 1 \right\}$ satisfies

$$\begin{aligned} & \min |\hat{u}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*| \\ & \lesssim \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^2 \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} |\mathbf{a}^\top \mathbf{u}_l^*| \\ & \quad + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|}. \end{aligned} \quad (80)$$

Theorems 9-10 taken together establish Theorem 1.

1) *Proof of Theorem 9:* We shall start by proving the results for $\mathbf{u}_l$; the proofs for the results w.r.t. $\mathbf{w}_l$ are clearly identical.

a) *Proofs for ℓ_2 perturbation bounds:* To derive the ℓ_2 bound for $\mathbf{u}_l$, we start by considering the distance between $\mathbf{u}_l$ and the subspace spanned by the true eigenvectors $\{\mathbf{u}_k^*\}_{1 \leq k \leq r}$. The Neumann series (cf. Lemma 1) tells us that: when $\|\mathbf{H}\| \leq \lambda_{\min}^*/4$,

$$\begin{aligned} & \left\| \mathbf{u}_l - \sum_{k=1}^r \frac{\lambda_k^* \mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l} \mathbf{u}_k^* \right\|_2 \\ & = \left\| \sum_{k=1}^r \frac{\lambda_k^* (\mathbf{u}_k^{\top} \mathbf{u}_l)}{\lambda_l} \left\{ \sum_{s=1}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \mathbf{u}_k^* \right\} \right\|_2 \\ & \leq \sqrt{\sum_{k=1}^r \left| \frac{\lambda_k^* (\mathbf{u}_k^{\top} \mathbf{u}_l)}{\lambda_l} \right|^2} \left\| \left(\sum_{s=1}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \right) [\mathbf{u}_1^*, \dots, \mathbf{u}_r^*] \right\| \\ & \leq \sqrt{\frac{(\lambda_{\max}^*)^2}{|\lambda_l|^2} \sum_{k=1}^r |\mathbf{u}_k^{\top} \mathbf{u}_l|^2} \left\| \left(\sum_{s=1}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \right) \right\| \|\mathbf{u}_1^*, \dots, \mathbf{u}_r^*\| \\ & \leq \frac{\lambda_{\max}^*}{|\lambda_l|} \sum_{s=1}^{\infty} \left(\frac{\|\mathbf{H}\|}{|\lambda_l|} \right)^s, \end{aligned} \quad (81)$$

where the last inequality holds since $\|\mathbf{u}_1^*, \dots, \mathbf{u}_r^*\| = 1$ and $\sum_{k=1}^r |\mathbf{u}_k^{\top} \mathbf{u}_l|^2 \leq \|\mathbf{u}_l\|_2^2 = 1$ (given that the $\mathbf{u}_k^*$'s are orthonormal). In view of the Bauer-Fike theorem and the bound $\|\mathbf{H}\| \leq \lambda_{\min}^*/4$, we have the crude lower bound

$$|\lambda_l| \geq \lambda_{\min}^* - \|\mathbf{H}\| \geq 3\lambda_{\min}^*/4. \quad (82)$$

This taken collectively with the inequality (81) yields

$$\begin{aligned} & \left\| \mathbf{u}_l - \sum_{k=1}^r \frac{\lambda_k^* \mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l} \mathbf{u}_k^* \right\|_2 \\ & \leq \frac{\lambda_{\max}^*}{|\lambda_l|} \cdot \frac{\|\mathbf{H}\|}{|\lambda_l| - \|\mathbf{H}\|} \leq \frac{\lambda_{\max}^*}{\frac{3}{4}\lambda_{\min}^*} \cdot \frac{\|\mathbf{H}\|}{\frac{1}{2}\lambda_{\min}^*} = \frac{8\kappa^2}{3} \cdot \frac{\|\mathbf{H}\|}{\lambda_{\max}^*}. \end{aligned} \quad (83)$$

In addition, note that the Euclidean projection of $\mathbf{u}_l$ onto the subspace spanned by $\{\mathbf{u}_k^*\}_{1 \leq k \leq r}$ is given by $\sum_{k=1}^r (\mathbf{u}_k^{\top} \mathbf{u}_l) \mathbf{u}_k^*$. This together with (83) implies that

$$\begin{aligned} & \frac{8\kappa^2}{3} \cdot \frac{\|\mathbf{H}\|}{\lambda_{\max}^*} \\ & \geq \left\| \mathbf{u}_l - \sum_{k=1}^r \frac{\lambda_k^* \mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l} \mathbf{u}_k^* \right\|_2 \geq \min_{\mathbf{z} \in \text{span}\{\mathbf{u}_1^*, \dots, \mathbf{u}_r^*\}} \|\mathbf{u}_l - \mathbf{z}\|_2 \\ & = \left\| \mathbf{u}_l - \sum_{k=1}^r (\mathbf{u}_k^{\top} \mathbf{u}_l) \mathbf{u}_k^* \right\|_2 = \sqrt{1 - \sum_{k=1}^r |\mathbf{u}_k^{\top} \mathbf{u}_l|^2}, \end{aligned} \quad (84)$$

where the last identity arises from the Pythagorean theorem.

The above inequality (84) indicates that most of the energy of $\mathbf{u}_l$ lies in $\text{span}\{\mathbf{u}_1^*, \dots, \mathbf{u}_r^*\}$. In order to show that $|\mathbf{u}_l^{\top} \mathbf{u}_l| \approx 1$, it suffices to justify that $|\mathbf{u}_k^{\top} \mathbf{u}_l|$ is very small for any $k \neq l$. To this end, taking $\mathbf{a} = \mathbf{u}_k^*$ in Lemma 3 and making use of Condition (13b), we arrive at

$$\left| \mathbf{u}_k^{\top} \left(\mathbf{u}_l - \sum_{i=1}^r \frac{\mathbf{u}_i^{\top} \mathbf{u}_l}{\lambda_l / \lambda_i^*} \mathbf{u}_i^* \right) \right| \leq c_2 \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n}. \quad (85)$$

In addition, the orthonormality of $\{\mathbf{u}_i^*\}_{1 \leq i \leq r}$ indicates that

$$\begin{aligned} \mathbf{u}_k^{\top} \left(\mathbf{u}_l - \sum_{i=1}^r \frac{\mathbf{u}_i^{\top} \mathbf{u}_l}{\lambda_l / \lambda_i^*} \mathbf{u}_i^* \right) & = \mathbf{u}_k^{\top} \mathbf{u}_l - \frac{\lambda_k^*}{\lambda_l} \mathbf{u}_k^{\top} \mathbf{u}_l \\ & = \left(1 - \frac{\lambda_k^*}{\lambda_l} \right) \mathbf{u}_k^{\top} \mathbf{u}_l \end{aligned} \quad (86)$$

and, as a result,

$$|\mathbf{u}_k^{\top} \mathbf{u}_l| \leq \left| 1 - \frac{\lambda_k^*}{\lambda_l} \right|^{-1} c_2 \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n}. \quad (87)$$

In order to control $\left| 1 - \frac{\lambda_k^*}{\lambda_l} \right|^{-1}$, we note that: for all $k \neq l$, it follows from the triangle inequality and Theorem 7 that

$$\begin{aligned} |\lambda_k^* - \lambda_l| & \geq |\lambda_k^* - \lambda_l^*| - |\lambda_l^* - \lambda_l| \geq |\lambda_k^* - \lambda_l^*| - \Delta_l^*/2 \\ & \geq |\lambda_k^* - \lambda_l^*|/2 \geq \Delta_l^*/2, \\ \Rightarrow \left| 1 - \frac{\lambda_k^*}{\lambda_l} \right|^{-1} & = \frac{|\lambda_l|}{|\lambda_k^* - \lambda_l|} \leq \frac{\lambda_{\max}^* + \|\mathbf{H}\|}{\Delta_l^*/2} \leq \frac{4\lambda_{\max}^*}{\Delta_l^*}, \end{aligned} \quad (88)$$

where we have also used the condition $\|\mathbf{H}\| \leq \lambda_{\max}^*$. Substitution into (87) yields

$$\begin{aligned} |\mathbf{u}_k^{\top} \mathbf{u}_l| & \leq \frac{4\lambda_{\max}^*}{\Delta_l^*} \cdot c_2 \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \\ & = \frac{4c_2 \sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r \log n} \end{aligned} \quad (89)$$

for all $k \neq l$. Putting the above results together, we arrive at

$$\begin{aligned}
& \min \left\{ \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_2^2 \right\} \\
&= 2 - 2|\mathbf{u}_l^{\top} \mathbf{u}_l| \leq 2 - 2|\mathbf{u}_l^{\top} \mathbf{u}_l|^2 \\
&\leq 2 \left\{ 1 - \sum_{k=1}^r |\mathbf{u}_k^{\top} \mathbf{u}_l|^2 \right\} + 2 \left\{ \sum_{k \neq l, k=1}^r |\mathbf{u}_k^{\top} \mathbf{u}_l|^2 \right\} \\
&\lesssim \left(\kappa^2 \frac{\|\mathbf{H}\|}{\lambda_{\max}^*} \right)^2 + \left(\frac{\sigma_{\max}}{\Delta_l^*} \right)^2 \mu \kappa^4 r^2 \log n \\
&\lesssim \left(\kappa^2 \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\max}^*} \right)^2 + \left(\frac{\sigma_{\max}}{\Delta_l^*} \right)^2 \mu \kappa^4 r^2 \log n, \quad (90)
\end{aligned}$$

where the penultimate inequality relies on (84) and (89), and the last line follows from (60b).

Finally, the advertised bound on $|\mathbf{u}_l^{\top} \mathbf{u}_l|$ can be immediately established by combining the identity $|\mathbf{u}_l^{\top} \mathbf{u}_l| = 1 - \frac{1}{2} \min \{ \|\mathbf{u}_l \pm \mathbf{u}_l^*\|_2^2 \}$ and the above bound (90).

b) *Proof for perturbation of linear forms of eigenvectors:* Invoke the triangle inequality to obtain

$$\begin{aligned}
& \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right) \right| \\
&\leq \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \sum_{k=1}^r \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \mathbf{u}_k^* \right) \right| + \sum_{k \neq l, k=1}^r \left| \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} (\mathbf{a}^{\top} \mathbf{u}_k^*) \right|. \quad (91)
\end{aligned}$$

The first term on the right-hand side of (91) can be controlled via Lemma 3 and Condition (13b):

$$\left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \sum_{k=1}^r \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \mathbf{u}_k^* \right) \right| \lesssim \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \cdot \|\mathbf{a}\|_2.$$

Regarding the second term on the right-hand side of (91), we can invoke (87) to derive

$$\begin{aligned}
\left| \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \right| &\lesssim \frac{|\lambda_k^*|}{|\lambda_l - \lambda_k^*|} \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \\
&\leq \frac{\lambda_{\max}^*}{|\lambda_l^* - \lambda_k^*|/2} \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \\
&\lesssim \frac{\sigma_{\max}}{|\lambda_l^* - \lambda_k^*|} \sqrt{\mu \kappa^4 r \log n}, \quad (92)
\end{aligned}$$

where the penultimate inequality results from (88). As a consequence, one has

$$\begin{aligned}
& \sum_{k: k \neq l, 1 \leq k \leq r} \left| \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} (\mathbf{a}^{\top} \mathbf{u}_k^*) \right| \\
&\lesssim r \cdot \sigma_{\max} \sqrt{\mu \kappa^4 r \log n} \cdot \left\{ \max_{k: k \neq l} \frac{|\mathbf{a}^{\top} \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \right\}. \quad (93)
\end{aligned}$$

Substituting the above bounds into (91) and rearranging terms, we obtain

$$\begin{aligned}
& \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right) \right| \\
&\lesssim \sigma_{\max} \sqrt{\mu \kappa^2 r \log n} \left\{ \frac{1}{\lambda_{\min}^*} \|\mathbf{a}\|_2 + \kappa r \max_{k: k \neq l} \frac{|\mathbf{a}^{\top} \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \right\}. \quad (94)
\end{aligned}$$

Next, recall from Theorem 7 that

$$\begin{aligned}
& |\lambda_l - \lambda_l^*| \leq c_4 \kappa r^2 \sigma_{\max} \sqrt{\mu \log n} < \frac{1}{2} |\lambda_l^*| \\
\Rightarrow \quad & \frac{\lambda_l}{\lambda_l^*} \in \left[1 - \frac{|\lambda_l - \lambda_l^*|}{|\lambda_l^*|}, 1 + \frac{|\lambda_l - \lambda_l^*|}{|\lambda_l^*|} \right] \subset [0.5, 1.5],
\end{aligned}$$

provided that $\mu \kappa^2 r^4 \lesssim n$ and that Condition (13b) holds. We then make use of the bounds (63) and (79a) to deduce that

$$\begin{aligned}
& \min \left| 1 \pm \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right| \\
&= \left| 1 - \frac{\lambda_l^*}{\lambda_l} + \frac{\lambda_l^*}{\lambda_l} - \frac{|\mathbf{u}_l^{\top} \mathbf{u}_l|}{\lambda_l / \lambda_l^*} \right| \\
&\leq \left| 1 - \frac{\lambda_l^*}{\lambda_l} \right| + \left| \frac{\lambda_l^*}{\lambda_l} \right| \cdot |1 - |\mathbf{u}_l^{\top} \mathbf{u}_l|| \\
&\lesssim \frac{|\lambda_l - \lambda_l^*|}{\lambda_{\min}^*} + |1 - |\mathbf{u}_l^{\top} \mathbf{u}_l|| \\
&\lesssim \frac{\kappa r^2 \sigma_{\max} \sqrt{\mu \log n}}{\lambda_{\min}^*} \\
&\quad + (\kappa^4 \sigma_{\max}^2 n \log n) \left\{ \frac{1}{(\Delta_l^*)^2} \frac{\mu r^2}{n} + \frac{1}{(\lambda_{\max}^*)^2} \right\}.
\end{aligned}$$

The above bounds taken collectively demonstrate that

$$\begin{aligned}
& \min |\mathbf{a}^{\top} (\mathbf{u}_l \pm \mathbf{u}_l^*)| \\
&\leq \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right) \right| \\
&\leq \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right) \right| \\
&\quad + \left| \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{a}^{\top} \mathbf{u}_l^* - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{a}^{\top} \mathbf{u}_l^* \right| \\
&= \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right) \right| + \min \left| 1 \pm \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right| \cdot |\mathbf{a}^{\top} \mathbf{u}_l^*| \\
&\lesssim \mathcal{E}_{\text{au}}, \quad (95)
\end{aligned}$$

where

$$\begin{aligned}
& \mathcal{E}_{\text{au}} \\
&:= \sigma_{\max} \sqrt{\mu \kappa^2 r \log n} \left\{ \frac{1}{\lambda_{\min}^*} \|\mathbf{a}\|_2 + \kappa r \max_{k: k \neq l} \frac{|\mathbf{a}^{\top} \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \right\} \\
&\quad + \left\{ \frac{\kappa r^2 \sigma_{\max} \sqrt{\mu \log n}}{\lambda_{\min}^*} + (\kappa^4 \sigma_{\max}^2 n \log n) \left\{ \frac{1}{(\Delta_l^*)^2} \frac{\mu r^2}{n} \right. \right. \\
&\quad \left. \left. + \frac{1}{(\lambda_{\max}^*)^2} \right\} \right\} |\mathbf{a}^{\top} \mathbf{u}_l^*|.
\end{aligned}$$

Finally, when it comes to the ℓ_{∞} perturbation bound, we shall simply take $\mathbf{a} = \mathbf{e}_k$ ($1 \leq k \leq r$) in the above inequality and use the incoherence condition $|\mathbf{e}_k^{\top} \mathbf{u}_l^*| \leq \sqrt{\mu/n}$ ($1 \leq k \leq r$) to obtain

$$\begin{aligned}
& \left| \mathbf{e}_k^{\top} \left(\mathbf{u}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right) \right| \\
&\lesssim \sigma_{\max} \sqrt{\mu \kappa^2 r \log n} \left\{ \frac{1}{\lambda_{\min}^*} + \frac{\kappa r}{\Delta_l^*} \sqrt{\frac{\mu}{n}} \right\} \\
&\quad + \left\{ \frac{\kappa r^2 \sigma_{\max} \sqrt{\mu \log n}}{\lambda_{\min}^*} + (\kappa^4 \sigma_{\max}^2 n \log n) \left\{ \frac{1}{(\Delta_l^*)^2} \frac{\mu r^2}{n} \right. \right. \\
&\quad \left. \left. + \frac{1}{(\lambda_{\max}^*)^2} \right\} \right\} \sqrt{\mu/n}.
\end{aligned}$$

$$\begin{aligned}
& + \frac{1}{(\lambda_{\max}^*)^2} \Big\} \Big\} \sqrt{\frac{\mu}{n}} \\
& \lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^4 r \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}},
\end{aligned}$$

where the last line results from Condition (13a) and $\mu \kappa^2 r^4 \lesssim n$. This together with a union bound yields that: with high probability, one has

$$\begin{aligned}
& \min \| \mathbf{u}_l \pm \mathbf{u}_l^* \|_{\infty} \\
& \leq \left\| \mathbf{u}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right\|_{\infty} \\
& \lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^4 r \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}}
\end{aligned}$$

as claimed.

c) *Proof of the claims w.r.t. $\hat{\mathbf{u}}_l := \frac{\mathbf{u}_l + \mathbf{w}_l}{\|\mathbf{u}_l + \mathbf{w}_l\|_2}$.* With the above results in place, we can easily establish these claims w.r.t. $\hat{\mathbf{u}}_l$ as well. In what follows, we demonstrate how to establish the bound (79d) regarding the linear form; the proofs for ℓ_2 and ℓ_{∞} bounds follow from nearly identical arguments and are omitted for brevity.

To begin with, repeating the analysis (95) for $\mathbf{w}_l$ yields

$$\left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right) \right| \lesssim \mathcal{E}_{\text{au}}, \quad (96)$$

One can therefore combine (95) and (96) to obtain

$$\begin{aligned}
& \left| \mathbf{a}^{\top} \left(\frac{\mathbf{u}_l + \mathbf{w}_l}{2} \right) - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{a}^{\top} \mathbf{u}_l^* \right| \\
& \leq \frac{1}{2} \left| \mathbf{a}^{\top} \left(\mathbf{u}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right) \right| \\
& \quad + \frac{1}{2} \left| \mathbf{a}^{\top} \left(\mathbf{w}_l - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{u}_l^* \right) \right| \\
& \lesssim \mathcal{E}_{\text{au}}, \quad (97)
\end{aligned}$$

where the first inequality arises from the triangle inequality as well as the fact $\text{sign}(\mathbf{u}_l^{\top} \mathbf{w}_l) = \text{sign}(\mathbf{u}_l^{\top} \mathbf{u}_l)$ (see (200) in Lemma 15). This in turn allows us to deduce that

$$\begin{aligned}
& \left| \mathbf{a}^{\top} \left(\frac{\mathbf{u}_l + \mathbf{w}_l}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} \right) - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{a}^{\top} \mathbf{u}_l^* \right| \\
& = \left| \frac{2}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} \mathbf{a}^{\top} \left(\frac{\mathbf{u}_l + \mathbf{w}_l}{2} \right) \right. \\
& \quad \left. - \frac{2}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{a}^{\top} \mathbf{u}_l^* \right. \\
& \quad \left. + \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \left(\frac{2}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} - 1 \right) \mathbf{a}^{\top} \mathbf{u}_l^* \right| \\
& \leq \frac{2}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} \left| \mathbf{a}^{\top} \left(\frac{\mathbf{u}_l + \mathbf{w}_l}{2} \right) - \text{sign} \left(\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \right) \mathbf{a}^{\top} \mathbf{u}_l^* \right| \\
& \quad + \left(\frac{2}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} - 1 \right) |\mathbf{a}^{\top} \mathbf{u}_l^*| \\
& \lesssim \mathcal{E}_{\text{au}} + (2 - \|\mathbf{u}_l + \mathbf{w}_l\|_2) |\mathbf{a}^{\top} \mathbf{u}_l^*| \lesssim \mathcal{E}_{\text{au}}, \quad (98)
\end{aligned}$$

where the first inequality follows from the triangle inequality, the penultimate inequality relies on (97) and the fact $\|\mathbf{u}_l + \mathbf{w}_l\|_2 \asymp 1$ (see Lemma 15), and the last

inequality follows from the bound $|2 - \|\mathbf{u}_l + \mathbf{w}_l\|_2| = O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right)$ (cf. Lemma 15) in addition to the expression of $\mathcal{E}_{\text{au}}$. This establishes the bound (79d) when $\mathbf{u}_l$ is replaced by $\hat{\mathbf{u}}_l$.

2) *Proof of Theorem 10:* We start by considering the coarse estimator $\hat{\mathbf{u}}_{\mathbf{a},l}$. Recall that $\mathbf{a}$ is a fixed unit vector obeying $\|\mathbf{a}\|_2 = 1$. The following trivial bound holds true:

$$\min |\hat{\mathbf{u}}_{\mathbf{a},l} \pm \mathbf{a}^{\top} \mathbf{u}_l^*| \leq 1. \quad (99)$$

Consequently, it suffices to focus on the scenario when the upper bound on the right-hand side of (80) is bounded above by some small constant; that is, the case where

$$\frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} \leq c_7 \quad (100a)$$

$$\frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} |\mathbf{a}^{\top} \mathbf{u}_l^*| \leq c_8 \quad (100b)$$

$$\sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^{\top} \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \leq c_9 \quad (100c)$$

for some sufficiently small constants $c_7, c_8, c_9 > 0$.

The key step is to invoke the Neumann series (cf. Lemma 1) for both $\mathbf{u}_l$ and $\mathbf{w}_l$, which yields

$$\begin{aligned}
& (\hat{\mathbf{u}}_{\mathbf{a},l})^2 \\
& = \left| \frac{(\mathbf{a}^{\top} \mathbf{u}_l)(\mathbf{a}^{\top} \mathbf{w}_l)}{\left(\sum_{k=1}^r \frac{\mathbf{u}_k^{\top} \mathbf{w}_l}{\lambda_l / \lambda_k^*} \sum_{s=0}^{\infty} \frac{(\mathbf{H}^{\top})^s \mathbf{u}_k^*}{\lambda_l^s} \right)^{\top} \left(\sum_{k=1}^r \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \sum_{s=0}^{\infty} \frac{\mathbf{H}^s \mathbf{w}_k^*}{\lambda_l^s} \right)} \right| \\
& = \left| \frac{(\mathbf{a}^{\top} \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_1)(\mathbf{a}^{\top} \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} + \delta_2)}{\frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_3} \right|. \quad (101)
\end{aligned}$$

Here, $\delta_1, \delta_2, \delta_3$ are defined as follows

$$\delta_1 := \mathbf{a}^{\top} \left(\mathbf{u}_l - \frac{\mathbf{u}_l^{\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right), \quad (102a)$$

$$\delta_2 := \mathbf{a}^{\top} \left(\mathbf{w}_l - \frac{\mathbf{u}_l^{\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \mathbf{u}_l^* \right), \quad (102b)$$

$$\delta_3 := \mathcal{U}_1 + \mathcal{U}_2, \quad (102c)$$

where the quantities $\mathcal{U}_1$ and $\mathcal{U}_2$ are given by

$$\mathcal{U}_1 := \sum_{k: k \neq l} \frac{\mathbf{u}_k^{\top} \mathbf{w}_l}{\lambda_l / \lambda_k^*} \cdot \frac{\mathbf{u}_k^{\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*}, \quad (103a)$$

$$\mathcal{U}_2 := \sum_{k_1=1}^r \sum_{k_2=1}^r \sum_{s_1, s_2: s_1+s_2 \geq 1} \frac{\mathbf{u}_{k_1}^{\top} \mathbf{w}_l}{\lambda_l / \lambda_{k_1}^*} \cdot \frac{\mathbf{u}_{k_2}^{\top} \mathbf{u}_l}{\lambda_l / \lambda_{k_2}^*} \cdot \frac{\mathbf{u}_{k_1}^{\top} \mathbf{H}^{s_1+s_2} \mathbf{u}_{k_2}^*}{\lambda_l^{s_1+s_2}}. \quad (103b)$$

This motivates us to control each term on the right-hand side of (101) separately.

- To begin with, it comes directly from the inequalities (13b) and (100c) that

$$\frac{\sigma_{\max} \sqrt{n \kappa^6 \log n}}{\lambda_{\min}^*} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r^2 \log n} \leq c_5 + \frac{c_9}{\sqrt{r}},$$

thus allowing us to invoke Theorem 9 to obtain, with probability at least $1 - O(n^{-6})$, that

$$\min \{ |\mathbf{u}_l^{\top} \mathbf{w}_l|, |\mathbf{u}_l^{\top} \mathbf{u}_l| \} \geq 3/4. \quad (104)$$

- In view of (94), with probability $1 - O(n^{-6})$ one has

$$\begin{aligned} & \max\{|\delta_1|, |\delta_2|\} \\ & \lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^2 r \log n}}{\lambda_{\min}^*} + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|}. \end{aligned} \quad (105)$$

- Next, substituting (92) into (103a) yields

$$\begin{aligned} |\mathcal{U}_1| & \leq r \max_{k \neq l} \left| \frac{\mathbf{u}_k^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_k^*} \cdot \frac{\mathbf{u}_k^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \right| \\ & \lesssim r \left(\frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r \log n} \right)^2 = \frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2}. \end{aligned} \quad (106)$$

In addition, one has

$$\begin{aligned} |\mathcal{U}_2| & \stackrel{(i)}{\lesssim} \sum_{k_1=1}^r \sum_{k_2=1}^r \sum_{s_1, s_2: s_1+s_2 \geq 1} \frac{1}{|\lambda_l^* / \lambda_{k_1}^* \cdot \lambda_l^* / \lambda_{k_2}^*|} \cdot \left| \frac{\mathbf{u}_{k_1}^{*\top} \mathbf{H}^{s_1+s_2} \mathbf{u}_{k_2}^*}{(\lambda_l^* / 2)^{s_1+s_2}} \right| \\ & \stackrel{(ii)}{\leq} \kappa^2 r^2 \sum_{s_1, s_2: s_1+s_2 \geq 1} \left(\frac{2}{\lambda_{\min}^*} \right)^{s_1+s_2} |\mathbf{u}_{k_1}^{*\top} \mathbf{H}^{s_1+s_2} \mathbf{u}_{k_2}^*| \\ & \stackrel{(iii)}{\leq} \kappa^2 r^2 \sum_{s=1}^{\infty} \sum_{s_1, s_2: s_1+s_2=s} \left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^s \sqrt{\frac{\mu}{n}} \\ & \leq \sqrt{\frac{\mu}{n}} \kappa^2 r^2 \sum_{s=1}^{\infty} 2^s \left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^s \\ & = \sqrt{\frac{\mu}{n}} \kappa^2 r^2 \frac{\frac{4c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}}{1 - \frac{4c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}} \\ & \stackrel{(iv)}{\lesssim} \frac{\sigma_{\max} \sqrt{\mu \kappa^4 r^4 \log n}}{\lambda_{\min}^*}, \end{aligned}$$

where (i) follows from Corollary 2 (so that $|\lambda_l| \geq |\lambda_l^*|/2$), (ii) makes use of the fact $|\lambda_k^* / \lambda_l^*| \leq \kappa$, (iii) arises from Lemma 2, and the last line relies on the assumption (13b) (so that $1 - \frac{4c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \geq \frac{1}{2}$). The above two bounds taken together yield

$$\begin{aligned} |\delta_3| & \leq |\mathcal{U}_1| + |\mathcal{U}_2| \\ & \lesssim \frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} + \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} \leq 1/8, \end{aligned} \quad (107)$$

where the last inequality holds as long as the constants c_7 and c_8 in (100b) and (100c) are sufficiently small.

With the above bounds in mind, we are ready to control the estimation error of $\hat{\mathbf{u}}_{a,l}$. We divide the proof into two separate cases as follows.

a) *Case 1: when $|\mathbf{a}^\top \mathbf{u}_l^*|$ is “small”:* Consider the case where

$$|\mathbf{a}^\top \mathbf{u}_l^*| \leq \max\{|\delta_1|, |\delta_2|\}.$$

Continue the bound in (101) to deduce that

$$\begin{aligned} (\hat{\mathbf{u}}_{a,l})^2 & = \left| \frac{\left(\mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_1 \right) \left(\mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} + \delta_2 \right)}{\frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_3} \right| \\ & \stackrel{(i)}{\leq} \frac{\{2|\mathbf{a}^\top \mathbf{u}_l^*| + |\delta_1|\} \{2|\mathbf{a}^\top \mathbf{u}_l^*| + |\delta_2|\}}{\left(\frac{1}{2}\right)^2 - \frac{1}{8}} \\ & \leq 72 (\max\{|\delta_1|, |\delta_2|\})^2. \end{aligned} \quad (108)$$

Here, the inequality (i) makes use of the bound (104), Corollary 2 (so that $1/2 \leq |\lambda_l / \lambda_l^*| \leq 2$), and the inequality $|\delta_3| \leq 1/8$. Therefore, we can take the triangle inequality to conclude that

$$\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*| \leq |\hat{\mathbf{u}}_{a,l}| + |\mathbf{a}^\top \mathbf{u}_l^*| \lesssim \max\{|\delta_1|, |\delta_2|\}, \quad (109)$$

which combined with (105) leads to the desired bound for this case.

b) *Case 2: when $|\mathbf{a}^\top \mathbf{u}_l^*|$ is “large”:* We now move on to the case where

$$|\mathbf{a}^\top \mathbf{u}_l^*| > \max\{|\delta_1|, |\delta_2|\}. \quad (110)$$

By the same argument above in (108), we have

$$\begin{aligned} & \left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \right| \\ & = \frac{1}{\left| \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_3 \right|} \left| \left(\mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_1 \right) \right. \\ & \quad \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} + \delta_2 \right) - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \left. \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_3 \right\} \right| \\ & = \frac{\left| \mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \cdot \delta_2 + \mathbf{a}^\top \mathbf{u}_l^* \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \delta_1 + \delta_1 \delta_2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot \delta_3 \right|}{\left| \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} + \delta_3 \right|} \\ & \lesssim \frac{|\mathbf{a}^\top \mathbf{u}_l^*| \max\{|\delta_1|, |\delta_2|\} + (\max\{|\delta_1|, |\delta_2|\})^2 + (\mathbf{a}^\top \mathbf{u}_l^*)^2 |\delta_3|}{\left(\frac{1}{2}\right)^2 - \frac{1}{8}} \\ & \lesssim |\mathbf{a}^\top \mathbf{u}_l^*| \left(\frac{\sigma_{\max} \sqrt{\mu \kappa^2 r \log n}}{\lambda_{\min}^*} + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \right) \\ & \quad + (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot \left(\frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} + \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} \right), \end{aligned} \quad (111)$$

where the last line relies on our upper bounds on $|\delta_1|, |\delta_2|$ and $|\delta_3|$ (see (105) and (107)) as well as the assumption (110).

Recognizing the trivial fact (here, define $\max\{a \pm b\} = \max\{|a+b|, |a-b|\}$)

$$\max |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*| \geq |\mathbf{a}^\top \mathbf{u}_l^*|,$$

we have

$$\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*| = \frac{|(\hat{\mathbf{u}}_{a,l})^2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2|}{\max |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*|} \leq \frac{|(\hat{\mathbf{u}}_{a,l})^2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2|}{|\mathbf{a}^\top \mathbf{u}_l^*|}. \quad (112)$$

If $\frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} \geq 0$, then

$$|(\hat{u}_{a,l})^2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2| = \left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \right|.$$

If instead $\frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} < 0$ holds, then

$$\begin{aligned} |(\hat{u}_{a,l})^2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2| &\leq |(\hat{u}_{a,l})^2 + (\mathbf{a}^\top \mathbf{u}_l^*)^2| \\ &= \left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \right|. \end{aligned}$$

Therefore, putting these bounds together and invoking our bound (111), we deduce that

$$\begin{aligned} \min |\hat{u}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l^*| &\leq \frac{\left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \right|}{|\mathbf{a}^\top \mathbf{u}_l^*|} \\ &\lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^2 r \log n}}{\lambda_{\min}^*} + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \\ &\quad + |\mathbf{a}^\top \mathbf{u}_l^*| \cdot \left(\frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2} + \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} \right) \\ &\lesssim \frac{\sigma_{\max} r^2 \sqrt{\mu \kappa^4 \log n}}{\lambda_{\min}^*} + \sigma_{\max} \sqrt{\mu \kappa^4 r^3 \log n} \max_{k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \\ &\quad + |\mathbf{a}^\top \mathbf{u}_l^*| \cdot \frac{\sigma_{\max}^2 \mu r^2 \kappa^4 \log n}{(\Delta_l^*)^2}. \end{aligned} \quad (113)$$

Taking collectively (109) and (113) establishes our estimation error bound for this case.

APPENDIX C

PROOF FOR DISTRIBUTIONAL GUARANTEES (THEOREM 5)

A. Distributional Theory for Linear Forms of Eigenvectors

By virtue of Theorem 7, we know that λ_l , $\mathbf{u}_l$ and $\mathbf{w}_l$ are all real-valued. Without loss of generality, assume that $\|\mathbf{a}\|_2 = 1$ and that $\mathbf{u}_l^\top \mathbf{u}_l > 0$, which combined with Lemma 15 and Notation 1 yield

$$\mathbf{u}_l^\top \mathbf{u}_l^* = 1 - o(1), \quad \mathbf{w}_l^\top \mathbf{u}_l^* = 1 - o(1), \quad \mathbf{w}_l^\top \mathbf{u}_l = 1 - o(1). \quad (114)$$

The main step lies in establishing the following claim

$$\hat{u}_{a,l}^{\text{modified}} = \mathbf{a}^\top \mathbf{u}_l^* + \frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \quad (115)$$

$$= \mathbf{a}^\top \mathbf{u}_l^* + \frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\lambda_l^*} + o(\sqrt{v_{a,l}^*}) \quad (116)$$

with $\mathbf{a}_l^\perp := \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*$, where the last line results from Lemma 12. Let us take Claim (115) as given for now and come back to its proof in the sequel. To establish Theorem 5, it suffices to pin down the distribution of $\frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\lambda_l^*}$ and show that it matches the distributional characterizations stated in Theorem 5. This follows immediately from the classical Berry-Esseen theorem, which we defer to Lemma 14.

The rest of the proof thus boils down to justifying the claim (115), for which we divide into two cases.

1) *The Case When $|\mathbf{a}^\top \mathbf{u}_l^*|$ Is Not “small”*: In this subsection, we focus on the scenario when $\sqrt{v_{a,l}^*} \log n = o(|\mathbf{a}^\top \mathbf{u}_l^*|)$ which, according to Lemma 6, subsumes the case $\sqrt{v_{a,l}^*} \log n = o(|\mathbf{a}^\top \mathbf{u}_l^*|)$. Without loss of generality, we assume that $\mathbf{a}^\top \mathbf{u}_l^* > 0$. By virtue of Lemma 12 and the assumption $\sigma_{\max} \asymp \sigma_{\min}$, we have

$$v_{a,l}^* \asymp \frac{\sigma_{\max}^2 (1 - |\mathbf{a}^\top \mathbf{u}_l^*|^2)}{|\lambda_l^*|^2} \asymp \frac{\sigma_{\max}^2}{|\lambda_l^*|^2},$$

where the last line follows from our assumption that $|\mathbf{a}^\top \mathbf{u}_l^*|$ is bounded away from 1. As a result, the regime considered in this subsection enjoys the following property:

$$\frac{\sigma_{\max} \log n}{|\lambda_l^*|} = o(\mathbf{a}^\top \mathbf{u}_l^*). \quad (117)$$

The key starting point is the following decomposition

$$\begin{aligned} (\hat{u}_{a,l}^{\text{modified}})^2 &= \left| \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} \right| \quad (118) \end{aligned}$$

$$\begin{aligned} &= \frac{1}{\left| \frac{\mathbf{u}_l^\top \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \left(1 + \frac{\mathbf{u}_l^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_3 \right|} \left| \left\{ \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \right. \right. \\ &\quad \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_1 \left. \left\{ \frac{\mathbf{u}_l^\top \mathbf{w}_l}{\lambda_l/\lambda_l^*} \right. \right. \\ &\quad \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_2 \left. \left. \right\} \right|, \end{aligned} \quad (119)$$

where

$$\begin{cases} \tau_1 &:= \mathbf{a}^\top \mathbf{u}_l - \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right), \\ \tau_2 &:= \mathbf{a}^\top \mathbf{w}_l - \frac{\mathbf{u}_l^\top \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right), \\ \tau_3 &:= \mathbf{u}_l^\top \mathbf{w}_l - \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \left(1 + \frac{\mathbf{u}_l^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right). \end{cases} \quad (120)$$

Here, τ_1, τ_2, τ_3 encompass second- or higher-order terms in the Neumann series. As it turns out, these terms can be well-controlled, as stated in the following lemma.

Lemma 4: Instate the assumptions of Theorem 4. With probability at least $1 - O(n^{-6})$, we have

$$|\tau_1| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \quad (121a)$$

$$|\tau_2| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \quad (121b)$$

$$|\tau_3| \lesssim \left(\frac{\sigma_{\max} \kappa^2 r \sqrt{\mu \log n}}{|\Delta_l^*|} \right)^2 + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \quad (121c)$$

Here, we recall that $\Delta_l^* = \infty$ when $r = 1$, meaning that $|\tau_3| = o(\sigma_{\max}/|\lambda_l^*|)$ when $r = 1$.

Recall from Theorem 7 and Lemma 15 that $|\mathbf{u}_l^\top \mathbf{u}_l| = 1 - o(1)$, $|\mathbf{u}_l^\top \mathbf{w}_l| = 1 - o(1)$ and $\lambda_l = (1 + o(1)) \lambda_l^*$. Therefore, Lemma 4 tells us that

$$|\tau_1| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \leq o(|\mathbf{a}^\top \mathbf{u}_l^*|) = o\left(\left| \frac{\mathbf{u}_l^\top \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \mathbf{a}^\top \mathbf{u}_l^* \right|\right);$$

$$|\tau_2| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \leq o(|\mathbf{a}^\top \mathbf{u}_l^*|) = o\left(\left| \frac{\mathbf{u}_l^\top \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \mathbf{a}^\top \mathbf{u}_l^* \right|\right).$$

In addition, Lemma 4 together with the assumption $\sigma_{\max} \kappa^2 r \sqrt{\mu \log n} = o(|\Delta_l^*|)$ implies that $|\tau_3| = o(1)$.

With these bounds in place, we see that τ_1 and $\frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*}$ are indeed very small terms, meaning that the term in (119) involving τ_1 obeys $\frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_1$ and $\frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{a}^\top \mathbf{u}_l^*$ have the same signs. Similar conclusions hold as well for the terms involving τ_2 and τ_3 . As a result,

$$\begin{aligned} & \text{sign} \left\{ \frac{(\mathbf{a}^\top \mathbf{u}_l)(\mathbf{a}^\top \mathbf{w}_l)}{\mathbf{u}_l^\top \mathbf{w}_l} \right\} \\ &= \text{sign} \left\{ \frac{\left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \mathbf{a}^\top \mathbf{u}_l^* \right\} \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \mathbf{a}^\top \mathbf{u}_l^* \right\}}{\frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*}} \right\} \\ &= \text{sign} \{ (\mathbf{a}^\top \mathbf{u}_l^*)^2 \} = 1. \end{aligned}$$

This indicates that we can safely remove the absolute value function in (119) to obtain

$$\begin{aligned} & (\hat{u}_{a,l}^{\text{modified}})^2 \\ &= \frac{1}{\frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(1 + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_3} \\ & \cdot \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_1 \right\} \\ & \cdot \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_2 \right\}. \quad (122) \end{aligned}$$

Armed with this expression, the next lemma develops the distributional characterization of $\hat{u}_{a,l}^{\text{modified}}$.

Lemma 5: Instate the assumptions of Theorem 4. With probability at least $1 - O(n^{-6})$, it follows that

$$\begin{aligned} & \left| (\hat{u}_{a,l}^{\text{modified}})^2 - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} (\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right)^2 \right| \\ & \lesssim |\mathbf{a}^\top \mathbf{u}_l^*| \cdot (|\tau_1| + |\tau_2|) + (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot |\tau_3| + O\left(\frac{\sigma_{\max}^2 \log n}{\lambda_l^{*2}}\right). \quad (123) \end{aligned}$$

The proof of this result is given in Section C-D.

Given our assumption $\mathbf{a}^\top \mathbf{u}_l^* > 0$, the expression (122) together with the bounds on τ_1 , τ_2 and τ_3 immediately implies that

$$\begin{aligned} & \left| \hat{u}_{a,l}^{\text{modified}} + \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right) \right| \\ & \geq (1 - o(1)) \cdot \mathbf{a}^\top \mathbf{u}_l^*. \end{aligned}$$

This combined with Lemma 5 leads to

$$\begin{aligned} & \left| \hat{u}_{a,l}^{\text{modified}} - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right) \right| \\ &= \frac{|(\hat{u}_{a,l})^2 - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right)^2|}{\left| \hat{u}_{a,l} + \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right) \right|} \\ & \lesssim |\tau_1| + |\tau_2| + |\mathbf{a}^\top \mathbf{u}_l^*| \cdot |\tau_3| + O\left(\frac{\sigma_{\max}^2 \log n}{\lambda_l^{*2} \cdot |\mathbf{a}^\top \mathbf{u}_l^*|}\right). \end{aligned}$$

Taking this together with Lemma 4 as well as the condition (117), we arrive at

$$\begin{aligned} & \left| \hat{u}_{a,l}^{\text{modified}} - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right) \right| \\ & \lesssim o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) + |\mathbf{a}^\top \mathbf{u}_l^*| \cdot \left(\frac{\sigma_{\max} \kappa^2 r \sqrt{\mu \log n}}{|\Delta_l^*|} \right)^2 \\ & = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \quad (124) \end{aligned}$$

where the last line results from our assumption $\mathbf{a}^\top \mathbf{u}_l^* = o\left(\frac{\Delta_l^{*2}}{|\lambda_l^*| \sigma_{\max} \kappa^4 r^2 \mu \log n}\right)$.

The proof of the claim (115) is thus complete for this case.

2) *The Case When $|\mathbf{a}^\top \mathbf{u}_l^*|$ Is “small”:* We then move on to the scenario where $|\mathbf{a}^\top \mathbf{u}_l^*| \lesssim \sqrt{v_{a,l}^*} \log^{1.5} n$, which clearly subsumes the case with $|\mathbf{a}^\top \mathbf{u}_l^*| \lesssim \sqrt{\hat{v}_{a,l}} \log^{1.5} n$ (according to Lemma 6). Once again, this restriction combined with Lemma 12 and the assumption $\sigma_{\max}/\sigma_{\min} = O(1)$ requires that

$$|\mathbf{a}^\top \mathbf{u}_l^*| \lesssim \frac{\sigma_{\max} \log^{1.5} n}{|\lambda_l^*|}. \quad (125)$$

The proof is built upon the following “first-order” approximations of $\mathbf{a}^\top \mathbf{u}_l$ and $\mathbf{a}^\top \mathbf{w}_l$.

Lemma 6: Instate the assumptions of Theorem 4. Fix any vector $\mathbf{a} \in \mathbb{R}^n$ with $\|\mathbf{a}\|_2 = 1$, and assume that

$$|\mathbf{a}^\top \mathbf{u}_k^*| = o\left(\frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \sqrt{\mu \kappa^4 r^3 \log n}}\right), \quad \forall k \neq l. \quad (126)$$

Then under Assumption 3, with probability $1 - O(n^{-6})$ one has

$$\left| \frac{\lambda_l}{\lambda_l^* (\mathbf{u}_l^\top \mathbf{u}_l)} \mathbf{a}^\top \mathbf{u}_l - \mathbf{a}^\top \mathbf{u}_l^* - \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \quad (127)$$

If we further have $|\mathbf{a}^\top \mathbf{u}_l^*| = o(1/\sqrt{\log n})$, then with probability $1 - O(n^{-6})$,

$$\begin{aligned} & \left\{ \frac{1}{\mathbf{u}_l^\top \mathbf{u}_l} \mathbf{a}^\top \mathbf{u}_l = \mathbf{a}^\top \mathbf{u}_l^* + \frac{(\mathbf{a}_l^\perp)^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right); \right. \\ & \left. \frac{1}{\mathbf{u}_l^\top \mathbf{w}_l} \mathbf{a}^\top \mathbf{w}_l = \mathbf{a}^\top \mathbf{u}_l^* + \frac{(\mathbf{a}_l^\perp)^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \right\}. \quad (128) \end{aligned}$$

Proof: See Appendix C-E. $\square$

In addition, we claim that the following relations hold

$$\begin{cases} \frac{1}{\mathbf{u}_l^\top \mathbf{u}_l} \mathbf{a}^\top \mathbf{u}_l = \mathbf{a}^\top \mathbf{u}_l^* + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \\ \frac{1}{\mathbf{w}_l^\top \mathbf{u}_l} \mathbf{a}^\top \mathbf{w}_l = \mathbf{a}^\top \mathbf{u}_l^* + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \end{cases} \quad (129)$$

If these claims were valid, then one would have

$$\begin{aligned} \hat{u}_{a,l}^{\text{modified}} &= (\mathbf{a}^\top \mathbf{u}_l + \mathbf{a}^\top \mathbf{w}_l) / 2 \\ &= \frac{1}{2\mathbf{u}_l^\top \mathbf{u}_l} \mathbf{a}^\top \mathbf{u}_l + \frac{1}{2\mathbf{w}_l^\top \mathbf{u}_l} \mathbf{a}^\top \mathbf{u}_l + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \\ &= \mathbf{a}^\top \mathbf{u}_l^* + \frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \end{aligned}$$

where the last inequality follows from Lemma 6. This validates the relation (115) for this case, as long as the relations (129) hold true.

*Proof of the relations (129): As a direct consequence of Lemma 6, we can write

$$\begin{aligned} & \frac{1}{\mathbf{u}_l^\top \mathbf{u}_l^*} \mathbf{a}^\top \mathbf{u}_l - \mathbf{a}^\top \mathbf{u}_l \\ &= \frac{1}{\mathbf{u}_l^\top \mathbf{u}_l^*} \mathbf{a}^\top \mathbf{u}_l \cdot (1 - \mathbf{u}^{\star\top} \mathbf{u}_l) \\ &= \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \right) (1 - \mathbf{u}^{\star\top} \mathbf{u}_l). \end{aligned} \quad (130)$$

In view of Lemma 13, we have with probability at least $1 - O(n^{-10})$ that

$$\left| \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| = O\left(\frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|}\right).$$

In addition, Lemma 15 together with Assumption 3 guarantees that, with high probability,

$$\begin{aligned} |1 - \mathbf{u}_l^{\star\top} \mathbf{u}_l| &= O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \\ &= o\left(\frac{1}{\log^{1.5} n}\right). \end{aligned}$$

Putting everything together and using the assumption $|\mathbf{a}^\top \mathbf{u}_l^*| \lesssim \frac{\sigma_{\max} \log^{1.5} n}{|\lambda_l^*|}$, we conclude that

$$\frac{1}{\mathbf{u}_l^\top \mathbf{u}^*} \mathbf{a}^\top \mathbf{u}_l - \mathbf{a}^\top \mathbf{u} = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right)$$

as claimed. The claim w.r.t. $\mathbf{w}_l$ follows from exactly the same argument.

B. Distributional Theory for Eigenvalues

Take $\mathbf{a} = \mathbf{u}_l^*$. By definition, we have $\mathbf{a}^\top \mathbf{u}_k^* = 0$ for any $k \neq l$. The expression (127) in Lemma 6 thus indicates that, with probability $1 - O(n^{-6})$, we have

$$\left| \frac{\lambda_l}{\lambda_l^* (\mathbf{u}_l^{\star\top} \mathbf{u}_l)} \mathbf{u}_l^{\star\top} \mathbf{u}_l - \mathbf{u}_l^{\star\top} \mathbf{u}_l^* - \frac{\mathbf{u}_l^{\star\top} \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right).$$

Rearranging terms and using the fact $\mathbf{u}_l^{\star\top} \mathbf{u}_l^* = 1$ further yield

$$|\lambda_l - \lambda_l^* - \mathbf{u}_l^{\star\top} \mathbf{H} \mathbf{u}_l^*| = o(\sigma_{\max}). \quad (131)$$

Consequently, it is sufficient to characterize the distribution of $\mathbf{u}_l^{\star\top} \mathbf{H} \mathbf{u}_l^*$. Setting $\mathbf{a} = \mathbf{u}_l^*$ in Lemma 14, we see that $W_{\lambda,l} = \frac{\mathbf{u}_l^{\star\top} \mathbf{H} \mathbf{u}_l^*}{\sqrt{v_{\lambda,l}^*}} = \frac{\mathbf{u}_l^{\star\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{2\sqrt{v_{\lambda,l}^*}}$ obeys

$$\sup_{z \in \mathbb{R}} |\mathbb{P}(W_{\lambda,l} \leq z) - \Phi(z)| \leq \frac{8}{\sqrt{\log n}} \quad (132)$$

as claimed.

C. Proof of Lemma 4

First of all, invoke the Neumann series (cf. Lemma 1 and (56)) and rearrange terms to obtain

$$\begin{aligned} \tau_1 &= \mathbf{a}^\top \mathbf{u}_l - \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) \\ &= \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left[\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{\lambda_l} \mathbf{a}^\top \mathbf{H} \mathbf{u}_l^* + \sum_{s=2}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_l^* \right. \\ &\quad \left. + \sum_{k:k \neq l} \frac{\lambda_k^*}{\lambda_l^*} \frac{\mathbf{u}_k^{\star\top} \mathbf{u}_l}{\mathbf{u}_l^{\star\top} \mathbf{u}_l} \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^* \right\} \right] - \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \\ &\quad \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) \\ &= \underbrace{\frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(\frac{1}{\lambda_l} - \frac{1}{\lambda_l^*} \right) \mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}_{=: \tau_{1,1}} + \underbrace{\sum_{k:k \neq l} \frac{\mathbf{u}_k^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \sum_{s=2}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}_{=: \tau_{1,2}} \\ &\quad + \underbrace{\sum_{k=1}^r \frac{\mathbf{u}_k^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*} \left\{ \sum_{s=2}^{\infty} \frac{1}{\lambda_l^s} \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^* \right\}}_{=: \tau_{1,3}} \end{aligned}$$

and, similarly,

$$\begin{aligned} \tau_3 &= \mathbf{u}_l^\top \mathbf{w}_l - \frac{\mathbf{u}_l^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(1 + \frac{\mathbf{u}_l^{\star\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) \\ &= \left(\sum_{j=1}^r \frac{\lambda_j^*}{\lambda_l} (\mathbf{u}_j^{\star\top} \mathbf{u}_l) \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \mathbf{u}_j^* \right\} \right) \\ &\quad \cdot \left(\sum_{j=1}^r \frac{\lambda_j^*}{\lambda_l} (\mathbf{u}_j^{\star\top} \mathbf{w}_l) \left\{ \sum_{s=0}^{\infty} \frac{1}{\lambda_l^s} \mathbf{H}^s \mathbf{u}_j^* \right\} \right) \\ &\quad - \frac{\mathbf{u}_l^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \left(1 + \frac{\mathbf{u}_l^{\star\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) \\ &= \underbrace{\frac{\mathbf{u}_l^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_l^*} \cdot \frac{\mathbf{u}_l^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_l^*} \cdot \mathbf{u}_l^{\star\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \cdot \left(\frac{1}{\lambda_l} - \frac{1}{\lambda_l^*} \right)}_{=: \tau_{3,1}} \\ &\quad + \underbrace{\sum_{k:k \neq l} \frac{\mathbf{u}_k^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_k^*} \cdot \frac{\mathbf{u}_k^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_k^*}}_{=: \tau_{3,2}} \\ &\quad + \underbrace{\sum_{(k_1, k_2) \neq (l, l)} \frac{\mathbf{u}_{k_1}^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_{k_1}^*} \cdot \frac{\mathbf{u}_{k_2}^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_{k_2}^*} \cdot \frac{2 \mathbf{u}_{k_1}^{\star\top} \mathbf{H} \mathbf{u}_{k_2}^*}{\lambda_l}}_{=: \tau_{3,3}} \\ &\quad + \underbrace{\sum_{k_1=1}^r \sum_{k_2=1}^r \sum_{s_1, s_2: s_1 + s_2 \geq 2} \frac{\mathbf{u}_{k_1}^{\star\top} \mathbf{w}_l}{\lambda_l / \lambda_{k_1}^*} \cdot \frac{\mathbf{u}_{k_2}^{\star\top} \mathbf{u}_l}{\lambda_l / \lambda_{k_2}^*} \cdot \frac{\mathbf{u}_{k_1}^{\star\top} \mathbf{H}^{s_1 + s_2} \mathbf{u}_{k_2}^*}{\lambda_l^{s_1 + s_2}}}_{=: \tau_{3,4}}. \end{aligned}$$

In the special case where $r = 1$, one has $\tau_{1,2} = \tau_{3,2} = \tau_{3,3} = 0$. In what follows, we develop bounds for these terms separately.

a) *Controlling $\tau_{1,1}$ and $\tau_{3,1}$* : We have learned from Theorem 7, Lemma 13 and Lemma 15 that: under Assumption 3,

one has $\lambda_l = (1 + o(1))\lambda_l^*$, $\mathbf{u}_l^\top \mathbf{u}_l^* = 1 - o(1)$,

$$\left| \frac{\lambda_l - \lambda_l^*}{\lambda_l} \right| \lesssim \frac{\sigma_{\max} \sqrt{\mu \kappa^2 r^4 \log n}}{|\lambda_l^*|} \quad \text{and} \quad (133)$$

$$\max \{ |\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*|, |\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*| \} \lesssim \sigma_{\max} \sqrt{\log n}. \quad (134)$$

It then follows that

$$\begin{aligned} |\tau_{1,1}| &\lesssim |\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*| \cdot \left| \frac{\lambda_l - \lambda_l^*}{(\lambda_l^*)^2} \right| \lesssim \frac{\sigma_{\max}^2 \kappa r^2 \sqrt{\mu} \log n}{|\lambda_l^*|^2} \\ &\leq \frac{\sigma_{\max}}{|\lambda_l^*|} \cdot \frac{\sigma_{\max} \kappa r^2 \sqrt{\mu} \log n}{|\lambda_{\min}^*|} = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \\ |\tau_{3,1}| &\lesssim |\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*| \cdot \left| \frac{\lambda_l - \lambda_l^*}{(\lambda_l^*)^2} \right| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \end{aligned}$$

with the proviso that $\sigma_{\max} \kappa r^2 \sqrt{\mu} \log n = o(\lambda_{\min}^*)$.

b) *Controlling $\tau_{1,2}$ and $\tau_{3,3}$.* We make the observation that for any $1 \leq k \leq r$,

$$\begin{aligned} \left| \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_k^*}{\lambda_l} \right| &\stackrel{(i)}{\lesssim} \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l|} \asymp \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|} \\ &\stackrel{(ii)}{=} o\left(\frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log^2 n}} \sqrt{\log n}\right) \\ &= o\left(\frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \sqrt{\mu \kappa^4 r^3 \log n}}\right), \end{aligned} \quad (135)$$

where (i) is a consequence of Lemma 13, and (ii) holds as long as $\sigma_{\max} \sqrt{\mu \kappa^6 r^3 \log^2 n} = o(\Delta_l^*)$. Similarly, for any $1 \leq k_1, k_2 \leq r$,

$$\left| \frac{\mathbf{u}_{k_1}^{*\top} \mathbf{H} \mathbf{u}_{k_2}^*}{\lambda_l} \right| = o\left(\frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log n}}\right). \quad (136)$$

Additionally, one can invoke the inequality (87) to deduce that for any $k \neq l$,

$$\begin{aligned} |\mathbf{u}_k^{*\top} \mathbf{u}_l| &\lesssim \left| 1 - \frac{\lambda_k^*}{\lambda_l} \right|^{-1} \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \\ &\lesssim \frac{\lambda_l^*}{|\lambda_l^* - \lambda_k^*|} \cdot \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n}. \end{aligned} \quad (137)$$

Putting the above bounds together and using the conditions $|\mathbf{a}^\top \mathbf{u}_k^*| = o\left(\frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \sqrt{\mu \kappa^4 r^3 \log n}}\right)$ and $|\mathbf{u}_l^{*\top} \mathbf{u}_l| = 1 - o(1)$ give

$$\begin{aligned} |\tau_{1,2}| &\lesssim \sum_{k \neq l, 1 \leq k \leq r} \left| \frac{\lambda_k^*}{\lambda_l^*} \right| \cdot \left| \frac{\mathbf{u}_k^{*\top} \mathbf{u}_l}{\mathbf{u}_l^{*\top} \mathbf{u}_l} \right| \left(\left| \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_k^*}{\lambda_l} \right| + |\mathbf{a}^\top \mathbf{u}_k^*| \right) \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) + o\left(r \cdot \frac{\lambda_{\max}^*}{|\lambda_l^* - \lambda_k^*|} \cdot \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n}\right. \\ &\quad \cdot \left. \frac{|\lambda_l^* - \lambda_k^*|}{|\lambda_l^*| \sqrt{\mu \kappa^4 r^3 \log n}}\right) \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \end{aligned} \quad (138)$$

In addition, for every pair $(k_1, k_2) \neq (l, l)$, Theorem 9 implies that

$$\begin{aligned} &(\mathbf{u}_{k_1}^{*\top} \mathbf{u}_l) (\mathbf{u}_{k_2}^{*\top} \mathbf{u}_l) \\ &= \begin{cases} O\left(\frac{\sigma_{\max}^2}{(\Delta_l^*)^2} \mu \kappa^4 r \log n\right), & \text{if } k_1 \neq l, k_2 \neq l, \\ O\left(\frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^4 r \log n}\right), & \text{else,} \end{cases} \end{aligned} \quad (139)$$

which together with (136) and $\sigma_{\max} \sqrt{\mu \kappa^6 r^3 \log^2 n} = o(\Delta_l^*)$ indicates that

$$\begin{aligned} I_1 &:= \left| \sum_{k_1 \neq l} \sum_{k_2 \neq l} \frac{\lambda_{k_1}^* \lambda_{k_2}^*}{\lambda_l^2} (\mathbf{u}_{k_1}^{*\top} \mathbf{w}_l) (\mathbf{u}_{k_2}^{*\top} \mathbf{u}_l) \cdot \frac{2 \mathbf{u}_{k_1}^{*\top} \mathbf{H} \mathbf{u}_{k_2}^*}{\lambda_l} \right| \\ &\lesssim r^2 \kappa^2 \max_{k \neq l, j \neq l} |\mathbf{u}_k^{*\top} \mathbf{u}_l| \cdot |\mathbf{u}_j^{*\top} \mathbf{w}_l| \cdot o\left(\frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log n}}\right) \\ &= o\left(\left(\frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^6 r^3 \log n}\right)^2 \cdot \frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log n}}\right) \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|} \cdot \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^6 r^3 \log n}\right) \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right); \\ I_2 &:= \left| \sum_{k_1 \neq l} \frac{\lambda_{k_1}^* \lambda_l^*}{\lambda_l^2} (\mathbf{u}_{k_1}^{*\top} \mathbf{w}_l) (\mathbf{u}_l^{*\top} \mathbf{u}_l) \cdot \frac{2 \mathbf{u}_{k_1}^{*\top} \mathbf{H} \mathbf{u}_l^*}{\lambda_l} \right| \\ &\lesssim r \kappa \max_{k_1 \neq l} |\mathbf{u}_{k_1}^{*\top} \mathbf{w}_l| \cdot o\left(\frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log n}}\right) \\ &= o\left(\frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\mu \kappa^6 r^3 \log n} \cdot \frac{\Delta_l^*}{|\lambda_l^*| \sqrt{\mu \kappa^6 r^3 \log n}}\right) \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right); \\ I_3 &:= \left| \sum_{k_2 \neq l} \frac{\lambda_l^* \lambda_{k_2}^*}{\lambda_l^2} (\mathbf{u}_l^{*\top} \mathbf{w}_l) (\mathbf{u}_{k_2}^{*\top} \mathbf{u}_l) \cdot \frac{2 \mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_{k_2}^*}{\lambda_l} \right| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \end{aligned}$$

Putting the preceding bounds together yields

$$|\tau_{3,3}| \leq I_1 + I_2 + I_3 = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \quad (140)$$

c) *Controlling $\tau_{1,3}$ and $\tau_{3,4}$.* It can be deduced from (133) and $|\mathbf{u}_l^{*\top} \mathbf{u}_l| = 1 - o(1)$ that

$$\begin{aligned} |\tau_{1,3}| &\lesssim \sum_{k=1}^r \left| \frac{\lambda_k^*}{\lambda_l^*} \right| \cdot \left| \frac{\mathbf{u}_k^{*\top} \mathbf{u}_l}{\mathbf{u}_l^{*\top} \mathbf{u}_l} \right| \sum_{s=2}^{\infty} \left| \frac{\mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}{\lambda_l^s} \right| \\ &\lesssim \kappa \sum_{k=1}^r \sum_{s=2}^{\infty} \left| \frac{\mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}{\lambda_l^s} \right|. \end{aligned} \quad (141)$$

In addition,

$$\kappa \sum_{k=1}^r \sum_{s=2}^{\infty} \left| \frac{\mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}{\lambda_l^s} \right| \lesssim \frac{\kappa}{|\lambda_l^*|} \sum_{k=1}^r \sum_{s=2}^{\infty} \left| \frac{2^s \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}{(\lambda_{\min}^*)^{s-1}} \right|, \quad (142)$$

where the last line follows since $|\lambda_l| \geq \lambda_{\min}^* - \|\mathbf{H}\| \geq \lambda_{\min}^*/2$ (by the Bauer-Fike theorem and (61)). Applying Lemma 2

(or (60a)) yields

$$\begin{aligned}
& \frac{\kappa}{|\lambda_l^*|} \sum_{k=1}^r \sum_{s=2}^{\infty} \left| \frac{2^s \mathbf{a}^\top \mathbf{H}^s \mathbf{u}_k^*}{(\lambda_{\min}^*)^{s-1}} \right| \\
& \leq \frac{\lambda_{\min}^* \kappa}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \sum_{k=1}^r \sum_{s=2}^{\infty} \left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^s \\
& \leq \frac{\lambda_{\min}^* \kappa r}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \frac{\left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^2}{1 - \frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}} \\
& \asymp \frac{\lambda_{\min}^* \kappa r}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \left(\frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^2 = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \quad (143)
\end{aligned}$$

provided that $\sigma_{\max} \sqrt{\mu \kappa^2 r^2 n \log n} = o(\lambda_{\min}^*)$. The above bounds thus imply that

$$|\tau_{1,3}| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right). \quad (144)$$

Similarly, we can upper bound $\tau_{3,4}$ by

$$\begin{aligned}
|\tau_{3,4}| & \lesssim \kappa^2 \sum_{k_1=1}^r \sum_{k_2=1}^r \sum_{s_1, s_2: s_1+s_2 \geq 2} \left| \frac{\mathbf{u}_{k_1}^{*\top} \mathbf{H}^{s_1+s_2} \mathbf{u}_{k_2}^*}{\lambda_l^{s_1+s_2}} \right| \\
& \lesssim \frac{\kappa^2}{|\lambda_l^*|} \sum_{k_1=1}^r \sum_{k_2=1}^r \sum_{s_1, s_2: s_1+s_2 \geq 2} \left| \frac{2^{s_1+s_2} \mathbf{u}_{k_1}^{*\top} \mathbf{H}^{s_1+s_2} \mathbf{u}_{k_2}^*}{(\lambda_{\min}^*)^{s_1+s_2-1}} \right| \\
& \leq \frac{\lambda_{\min}^* \kappa^2 r^2}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \sum_{s_1, s_2: s_1+s_2 \geq 2} \left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^{s_1+s_2} \\
& \stackrel{(i)}{=} \frac{\lambda_{\min}^* \kappa^2 r^2}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \left\{ \left(\frac{1}{1 - \frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}} \right)^2 - 1 \right. \\
& \quad \left. - 2 \cdot \frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right\} \\
& \stackrel{(ii)}{=} \frac{\lambda_{\min}^* \kappa^2 r^2}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \cdot \left(\frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^2 \\
& \quad \cdot \frac{3 - \frac{4c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}}{\left(1 - \frac{2c_2 \sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^2} \\
& \asymp \frac{\lambda_{\min}^* \kappa^2 r^2}{|\lambda_l^*|} \sqrt{\frac{\mu}{n}} \left(\frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \right)^2 = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \quad (145)
\end{aligned}$$

provided that $\sigma_{\max} \kappa^2 r^2 \sqrt{n \mu \log n} = o(\lambda_{\min}^*)$. Here, (i) and (ii) make use of the elementary identity that $\sum_{s_1, s_2: s_1+s_2 \geq 2} a^{s_1+s_2} = \frac{1}{(1-a)^2} - 1 - 2a = \frac{3a^2(1-2a)}{(1-a)^2}$ for all $0 < a < 1$.

d) *Controlling $\tau_{3,2}$* : Clearly, one has $\tau_{3,2} = 0$ if $r = 1$. When $r \geq 2$, this term $\tau_{3,2}$ is the same term as $\mathcal{U}_1$ defined in (103a). Recalling our bound for inner product $|\mathbf{u}_k^{*\top} \mathbf{w}_l|$ and

$|\mathbf{u}_k^{*\top} \mathbf{u}_l|$ when $k \neq l$ from (137), we deduce that

$$\begin{aligned}
|\tau_{3,2}| & \lesssim \sum_{k:k \neq l} \left(\frac{\lambda_k^*}{\lambda_l^*} \right)^2 \cdot \left(\frac{\lambda_l^*}{|\lambda_l^* - \lambda_k^*|} \cdot \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^2 r \log n} \right)^2 \\
& \leq r \left(\frac{\sigma_{\max} \kappa^2 \sqrt{\mu r \log n}}{|\Delta_l^*|} \right)^2 \leq \left(\frac{\sigma_{\max} \kappa^2 r \sqrt{\mu \log n}}{|\Delta_l^*|} \right)^2.
\end{aligned}$$

Putting all this together, we arrive at

$$\begin{aligned}
|\tau_1| & \leq |\tau_{1,1}| + |\tau_{1,2}| + |\tau_{1,3}| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \\
|\tau_3| & \leq |\tau_{3,1}| + |\tau_{3,2}| + |\tau_{3,3}| + |\tau_{3,4}| \\
& = O\left(\frac{\sigma_{\max} \kappa^2 r \sqrt{\mu \log n}}{|\Delta_l^*|}\right)^2 + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right)
\end{aligned}$$

as claimed. The bound on $|\tau_2|$ can be established using exactly the same way as for $|\tau_1|$.

D. Proof of Lemma 5

Direct calculation yields

$$\begin{aligned}
& (\hat{\mathbf{u}}_{a,l}^{\text{modified}})^2 - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right)^2 \\
& = \frac{\mathcal{R}}{\frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l/\lambda_l^*} \left(1 + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_3}, \quad (146)
\end{aligned}$$

where $\mathcal{R} = \mathcal{R}_1 - \mathcal{R}_2$ and

$$\begin{aligned}
\mathcal{R}_1 & = \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_1 \right\} \\
& \quad \cdot \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_2 \right\}, \\
\mathcal{R}_2 & = \left\{ \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l/\lambda_l^*} \left(1 + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) + \tau_3 \right\} \\
& \quad \cdot \left\{ \mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right\}^2.
\end{aligned}$$

Rearranging terms and using the above bounds, we can derive

$$\begin{aligned}
\mathcal{R}_1 & = \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \left\{ (\mathbf{a}^\top \mathbf{u}_l^*)^2 + (\mathbf{a}^\top \mathbf{u}_l^*) \cdot \frac{\mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right. \\
& \quad \left. + \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \cdot \frac{\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} \right\} + \mathbf{a}^\top \mathbf{u}_l^* \cdot (\tau_1 + \tau_2) \\
& \quad + o(|\mathbf{a}^\top \mathbf{u}_l^*| \cdot (|\tau_1| + |\tau_2|)), \\
\mathcal{R}_2 & = \frac{\mathbf{u}_l^{*\top} \mathbf{w}_l}{\lambda_l/\lambda_l^*} \cdot \frac{\mathbf{u}_l^{*\top} \mathbf{u}_l}{\lambda_l/\lambda_l^*} \cdot \left\{ (\mathbf{a}^\top \mathbf{u}_l^*)^2 + (\mathbf{a}^\top \mathbf{u}_l^*) \right. \\
& \quad \cdot \frac{(\mathbf{a}_l^{\perp\top} + (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \\
& \quad \left. + \left(1 + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) \cdot \frac{(\mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*)^2}{4(\lambda_l^*)^2} \right. \\
& \quad \left. + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \cdot \frac{\mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right\} \\
& \quad + (1 + o(1)) (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot \tau_3.
\end{aligned}$$

By definition, we have $\mathbf{a}_l^\perp + (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^* = \mathbf{a}$, which further gives

$$\begin{aligned} |\mathcal{R}| &= |\mathcal{R}_1 - \mathcal{R}_2| \\ &= \left| (1 + o(1)) \left\{ \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \cdot \frac{\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*}{\lambda_l^*} \right. \right. \\ &\quad \left. \left. - \left(1 + \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right) \frac{(\mathbf{a}_l^\perp)^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{4\lambda_l^{*2}} - \frac{\mathbf{u}_l^{*\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \cdot \frac{\mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*}{\lambda_l^*} \right\} \right. \\ &\quad \left. + \mathbf{a}^\top \mathbf{u}_l^* (\tau_1 + \tau_2) + o(|\mathbf{a}^\top \mathbf{u}_l^*|(|\tau_1| + |\tau_2|)) \right. \\ &\quad \left. - (1 + o(1))(\mathbf{a}^\top \mathbf{u}_l^*)^2 \tau_3 \right| \\ &\lesssim |\mathbf{a}^\top \mathbf{u}_l^*| \cdot (|\tau_1| + |\tau_2|) + (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot |\tau_3| + \frac{\sigma_{\max}^2 \log n}{\lambda_l^{*2}} \end{aligned}$$

with probability exceeding $1 - O(n^{-6})$. Here, the last line uses Lemma 13. Substitution into (146) yields

$$\begin{aligned} \left| (\hat{w}_{a,l}^{\text{modified}})^2 - \left(\mathbf{a}^\top \mathbf{u}_l^* + \frac{1}{2\lambda_l^*} \mathbf{a}_l^{\perp\top} (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^* \right)^2 \right| &\leq \frac{|\mathcal{R}|}{1 - o(1)} \\ &\lesssim |\mathbf{a}^\top \mathbf{u}_l^*| \cdot (|\tau_1| + |\tau_2|) + (\mathbf{a}^\top \mathbf{u}_l^*)^2 \cdot |\tau_3| + \frac{\sigma_{\max}^2 \log n}{\lambda_l^{*2}} \end{aligned}$$

as claimed.

E. Proof of Lemma 6

The first claim

$$\left| \frac{\lambda_l}{\lambda_l^* (\mathbf{u}_l^{*\top} \mathbf{u}_l)} \mathbf{a}^\top \mathbf{u}_l - \mathbf{a}^\top \mathbf{u}_l^* - \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| = o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \quad (147)$$

follows immediately from the bound on $|\tau_1|$ in Lemma 4 and the facts $\lambda_l = (1 + o(1))\lambda_l^*$ and $\mathbf{u}_l^\top \mathbf{u}_l^* = 1 - o(1)$. As an immediate consequence, one has

$$\begin{aligned} \left| \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} \right| &\asymp \left| \frac{\lambda_l}{\lambda_l^*} \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} \right| \\ &\leq \left| \frac{\lambda_l}{\lambda_l^*} \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} - \mathbf{a}^\top \mathbf{u}_l^* - \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| + |\mathbf{a}^\top \mathbf{u}_l^*| + \left| \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| \\ &\lesssim o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) + |\mathbf{a}^\top \mathbf{u}_l^*| + \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|} \\ &\asymp |\mathbf{a}^\top \mathbf{u}_l^*| + \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|}. \end{aligned} \quad (148)$$

Next, the bound (131) tells us that

$$\frac{\lambda_l}{\lambda_l^*} - 1 = \frac{\lambda_l - \lambda_l^*}{\lambda_l^*} = \frac{\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right).$$

This in turn allows us to bound

$$\begin{aligned} \left| \left(\frac{\lambda_l}{\lambda_l^*} - 1 \right) \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} \right| &\lesssim \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|} \left\{ |\mathbf{a}^\top \mathbf{u}_l^*| + \frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|} \right\} \\ &= o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \end{aligned}$$

as long as $|\mathbf{a}^\top \mathbf{u}_l^*| = o(1/\sqrt{\log n})$. Putting the above bounds together, we arrive at the advertised bound

$$\begin{aligned} &\left| \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} - \mathbf{a}^\top \mathbf{u}_l^* - \frac{(\mathbf{a}_l^\perp)^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| \\ &= \left| \frac{\lambda_l}{\lambda_l^*} \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} - \mathbf{a}^\top \mathbf{u}_l^* - \frac{(\mathbf{a}_l^\perp)^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} - \left(\frac{\lambda_l}{\lambda_l^*} - 1 \right) \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} \right| \\ &\leq \left| \frac{\lambda_l}{\lambda_l^*} \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} - \mathbf{a}^\top \mathbf{u}_l^* - \frac{(\mathbf{a}_l^\perp)^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} - \frac{\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \mathbf{a}^\top \mathbf{u}_l^* \right| \\ &\quad + O\left(\frac{\mathbf{u}_l^{*\top} \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \mathbf{a}^\top \mathbf{u}_l^*\right) + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \\ &\leq \left| \frac{\lambda_l}{\lambda_l^*} \frac{\mathbf{a}^\top \mathbf{u}_l}{\mathbf{u}_l^\top \mathbf{u}_l^*} - \mathbf{a}^\top \mathbf{u}_l^* - \frac{\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*}{\lambda_l^*} \right| \\ &\quad + O\left(\frac{\sigma_{\max} \sqrt{\log n}}{|\lambda_l^*|} |\mathbf{a}^\top \mathbf{u}_l^*| \right) + o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right) \\ &\leq o\left(\frac{\sigma_{\max}}{|\lambda_l^*|}\right), \end{aligned}$$

where the penultimate inequality holds with the proviso that $|\mathbf{a}^\top \mathbf{u}_l^*| = o(1/\sqrt{\log n})$.

The proof for the left eigenvector $\mathbf{w}_l$ is essentially the same and is thus omitted for brevity.

APPENDIX D

PROOF FOR VARIANCE ESTIMATION (THEOREM 6)

Without loss of generality, we assume $\|\mathbf{a}\|_2 = 1$ and $\lambda_l^* = 1$ in this section.

A. Proof Outline

e) The estimation accuracy of $\hat{v}_{a,l}$: Before continuing, we note that under the assumption that $|\mathbf{a}^\top \mathbf{u}_l^*| = o(\|\mathbf{a}\|_2) = o(1)$, we have learned from Lemma 12 (or (197)) that

$$\frac{1}{2} \sigma_{\min}^2 \|\mathbf{a}_l^\perp\|_2^2 \leq v_{a,l}^* \leq \sigma_{\max}^2 \|\mathbf{a}_l^\perp\|_2 \leq \sigma_{\max}^2 \|\mathbf{a}\|_2^2. \quad (149)$$

In addition, the assumption $|\mathbf{a}^\top \mathbf{u}_l^*| \leq (1 - \epsilon)\|\mathbf{a}\|_2$ for some constant $0 < \epsilon < 1$ implies that

$$1 = \|\mathbf{a}\|_2^2 \geq \|\mathbf{a}_l^\perp\|_2^2 = \|\mathbf{a}\|_2^2 - (\mathbf{a}^\top \mathbf{u}_l^*)^2 \geq 2\epsilon - \epsilon^2 > 0. \quad (150)$$

Consequently, under the assumptions $\sigma_{\max}/\sigma_{\min} \asymp 1$ and $\epsilon \asymp 1$ one has

$$v_{a,l}^* \asymp \sigma_{\max}^2. \quad (151)$$

Recall the construction of our variance estimator

$$\hat{v}_{a,l} := \frac{1}{4\lambda_l^2} \sum_{1 \leq i, j \leq n} (\hat{a}_{l,i}^\perp \hat{u}_{l,j} + \hat{a}_{l,j}^\perp \hat{u}_{l,i})^2 \hat{H}_{ij}^2, \quad (152)$$

where $\hat{a}_l^\perp := \mathbf{a} - (\mathbf{a}^\top \hat{\mathbf{u}}_l) \hat{\mathbf{u}}_l$ and $\mathbf{a}_l^\perp := \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*$. To control the size of $\hat{v}_{a,l}$, we find it convenient to introduce a surrogate quantity

$$\tilde{v}_{a,l} := \frac{1}{4} \sum_{1 \leq i, j \leq n} (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)^2 H_{ij}^2, \quad (153)$$

which turns out to be very close to our estimator $\hat{v}_{a,l}$.

Lemma 7: With probability exceeding $1 - O(n^{-5})$, one has $|\widehat{v}_{a,l} - \widetilde{v}_{a,l}| = o(v_{a,l}^*)$.

Proof: See Appendix D-B. $\square$

This observation allows us to turn attention to bounding $\widetilde{v}_{a,l}$ instead. The result is this:

Lemma 8: With probability exceeding $1 - O(n^{-10})$, one has $|\widetilde{v}_{a,l} - v_{a,l}^*| = o(v_{a,l}^*)$.

Proof: See Appendix D-C. $\square$

Putting Lemmas 7-8 together and invoking the triangle inequality, we establish the advertised bound

$$\widehat{v}_{a,l} = v_{a,l}^* + O(|\widehat{v}_{a,l} - \widetilde{v}_{a,l}| + |\widetilde{v}_{a,l} - v_{a,l}^*|) = (1 + o(1))v_{a,l}^*. \quad (154)$$

The rest of this section is thus mainly dedicated to establishing Lemma 7 and Lemma 8.

f) The estimation accuracy of $\widehat{v}_{\lambda,l}$: The proof for $\widehat{v}_{\lambda,l} = (1 + o(1))v_{\lambda,l}^*$ follows from a very similar argument, and is hence omitted.

B. Proof of Lemma 7

Without loss of generality, assume $\langle \mathbf{u}_l, \mathbf{u}_l^* \rangle > 0$, which combined with Lemma 15 and Notation III-A gives

$$\langle \widehat{\mathbf{u}}_l, \mathbf{u}_l^* \rangle > 0 \quad \text{and} \quad \langle \mathbf{w}_l, \mathbf{u}_l^* \rangle > 0. \quad (155)$$

To prove this lemma, let us first recall that, under the assumptions stated in Theorem 4 one has

$$\|\widehat{\mathbf{u}}_l - \mathbf{u}_l^*\|_\infty \lesssim o\left(\frac{1}{\sqrt{\mu n \log n}}\right), \quad \|\widehat{\mathbf{u}}_l\|_\infty \leq 2\sqrt{\frac{\mu}{n}}, \quad (156a)$$

$$|\mathbf{a}^\top \widehat{\mathbf{u}}_l - \mathbf{a}^\top \mathbf{u}_l^*| \lesssim o\left(\frac{1}{\mu^2 \log^2 n}\right), \quad (156b)$$

where (156a) comes from Lemma 15, and (156b) arises from (79d) and (98).

We also need the following properties that control the difference between each term of $\widehat{v}_{a,l}$ with that of $\widetilde{v}_{a,l}$. In particular, we introduce the following two lemmas.

Lemma 9: Suppose that Assumption 3 holds and that $32\kappa^2\sqrt{\mu r/n} \leq 1$. Then one has

$$\xi := \|\widehat{\mathbf{H}} - \mathbf{H}\|_\infty = \|\mathbf{M}_{\text{svd}} - \mathbf{M}^*\|_\infty \lesssim \sigma_{\max} \frac{\mu \kappa^4 r \sqrt{\log n}}{\sqrt{n}} \quad (157)$$

with probability exceeding $1 - O(n^{-9})$.

Proof: See Appendix D-D. $\square$

Lemma 10: With probability at least $1 - O(n^{-8})$, for every $1 \leq i, j \leq n$, one has

$$|a_{l,i}^\perp| \leq |a_i| + |u_{l,i}^*| \leq |a_i| + \sqrt{\mu/n}; \quad (158)$$

$$\begin{aligned} |\zeta_{ij}| &:= |a_i \widehat{u}_{l,j} + a_j \widehat{u}_{l,i} - (a_i u_{l,j}^* + a_j u_{l,i}^*)| \\ &\leq o\left(\frac{|a_i| + |a_j| + |u_{l,i}^*| + |u_{l,j}^*|}{\sqrt{\mu n \log n}} + \frac{1}{\mu n \log n}\right). \end{aligned} \quad (159)$$

Proof: Recalling that $\widehat{\mathbf{a}}_l^\perp = \mathbf{a} - (\mathbf{a}^\top \widehat{\mathbf{u}}_l) \widehat{\mathbf{u}}_l$ and $\mathbf{a}_l^\perp = \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*$, we obtain as a consequence of inequalities (156a)

and (156b) that

$$\begin{aligned} |\widehat{a}_{l,i}^\perp - a_{l,i}^\perp| &= |[\mathbf{a} - (\mathbf{a}^\top \widehat{\mathbf{u}}_l) \widehat{\mathbf{u}}_l - \mathbf{a} + (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^*]_i| \\ &= |(\mathbf{a}^\top \widehat{\mathbf{u}}_l) \widehat{u}_{l,i} - (\mathbf{a}^\top \mathbf{u}_l^*) u_{l,i}^*| \\ &\leq |\mathbf{a}^\top \widehat{\mathbf{u}}_l| \cdot |\widehat{u}_{l,i} - u_{l,i}^*| + |\mathbf{a}^\top \widehat{\mathbf{u}}_l - \mathbf{a}^\top \mathbf{u}_l^*| \cdot |u_{l,i}^*| \\ &\leq o\left(\frac{1}{\sqrt{\mu n \log n}}\right) + o\left(\frac{1}{\mu^2 \log^2 n}\right) \cdot \sqrt{\frac{\mu}{n}} \\ &= o\left(\frac{1}{\sqrt{\mu n \log n}}\right) \\ \text{and } |a_{l,i}^\perp| &\leq |a_i| + |\mathbf{a}^\top \mathbf{u}_l^*| \cdot |u_{l,i}^*| \leq |a_i| + |u_{l,i}^*| \\ &\leq |a_i| + \sqrt{\mu/n}. \end{aligned} \quad (160)$$

These in turn allow one to derive

$$\begin{aligned} &|\widehat{a}_{l,i}^\perp \widehat{u}_{l,j} + \widehat{a}_{l,j}^\perp \widehat{u}_{l,i} - (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)| \\ &\leq |\widehat{a}_{l,i}^\perp \widehat{u}_{l,j} + \widehat{a}_{l,j}^\perp \widehat{u}_{l,i} - (a_{l,i}^\perp \widehat{u}_{l,j} + a_{l,j}^\perp \widehat{u}_{l,i})| \\ &\quad + |a_{l,i}^\perp \widehat{u}_{l,j} + a_{l,j}^\perp \widehat{u}_{l,i} - (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)| \\ &\leq (|\widehat{a}_{l,i}^\perp - a_{l,i}^\perp| + |\widehat{a}_{l,j}^\perp - a_{l,j}^\perp|) (|u_{l,i}^*| + |u_{l,j}^*| + \|\widehat{\mathbf{u}}_l - \mathbf{u}_l^*\|_\infty) \\ &\quad + (|a_{l,i}^\perp| + |a_{l,j}^\perp|) \|\widehat{\mathbf{u}}_l - \mathbf{u}_l^*\|_\infty \\ &\leq o\left(\frac{1}{\sqrt{\mu n \log n}}\right) \left(|u_{l,i}^*| + |u_{l,j}^*| + o\left(\frac{1}{\sqrt{\mu n \log n}}\right)\right) \\ &\quad + (|a_i| + |a_j| + |u_{l,i}^*| + |u_{l,j}^*|) \cdot o\left(\frac{1}{\sqrt{\mu n \log n}}\right) \\ &= o\left(\frac{|a_i| + |a_j| + |u_{l,i}^*| + |u_{l,j}^*|}{\sqrt{\mu n \log n}} + \frac{1}{\mu n \log n}\right). \end{aligned}$$

$\square$

We are ready to establish the claim. Invoking the above two lemmas yields

$$\begin{aligned} &|4\lambda_l^2 \widehat{v}_{a,l} - 4\widetilde{v}_{a,l}| \\ &= \left| \sum_{i,j} (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^* + \zeta_{ij})^2 \widehat{H}_{ij}^2 \right. \\ &\quad \left. - \sum_{i,j} (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)^2 H_{ij}^2 \right| \\ &\leq \sum_{i,j} 2|\zeta_{ij}| (|a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*| + |\zeta_{ij}|) (|H_{ij}| + \xi)^2 \\ &\quad + \sum_{i,j} 2(|a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*| + |\zeta_{ij}|)^2 \xi (|H_{ij}| + \xi) \\ &\stackrel{(i)}{\lesssim} \frac{1}{n\sqrt{\log n}} \sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \\ &\quad \cdot (|H_{ij}| + \xi)^2 + \frac{\mu}{n} \sum_{i,j} \xi \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 \right. \\ &\quad \left. + \frac{1}{\mu n \log n} \right) \cdot (|H_{ij}| + \xi) \\ &\lesssim \frac{1}{n\sqrt{\log n}} \\ &\quad \cdot \underbrace{\left\{ \sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) H_{ij}^2 \right\}}_{=: \alpha_1} \\ &\quad + \frac{1}{n\sqrt{\log n}} \\ &\quad \cdot \underbrace{\left\{ \sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \xi^2 \right\}}_{=: \alpha_2} \end{aligned}$$

$$+ \frac{\mu}{n} \cdot \underbrace{\sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \xi (|H_{ij}| + \xi)}_{=: \alpha_3},$$

where (i) follows since (according to (159) and the incoherence condition)

$$\left\{ \begin{aligned} & |a_{l,i}^\perp u_j^* + a_{l,j}^\perp u_i^*| + |\zeta_{ij}| \\ & \lesssim \left(|a_i| + |a_j| + |u_{l,i}^*| + |u_{l,j}^*| + \sqrt{\frac{1}{\mu n \log n}} \right) \\ & \quad \cdot \left(|u_{l,i}^*| + |u_{l,j}^*| + \sqrt{\frac{1}{\mu n \log n}} \right) \\ & \lesssim \sqrt{a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n}} \cdot \sqrt{\frac{\mu}{n}} \\ & \lesssim \sqrt{\frac{\mu}{n} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right)}; \\ & |\zeta_{ij}| \left(|a_{l,i}^\perp u_j^* + a_{l,j}^\perp u_i^*| + |\zeta_{ij}| \right) \\ & \lesssim \left(|a_i| + |a_j| + |u_{l,i}^*| + |u_{l,j}^*| + \sqrt{\frac{1}{\mu n \log n}} \right)^2 \\ & \quad \cdot \sqrt{\frac{1}{\mu n \log n}} \cdot \sqrt{\frac{\mu}{n}} \\ & \lesssim \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \frac{1}{n \sqrt{\log n}}. \end{aligned} \right.$$

We then control α_1 , α_2 and α_3 separately.

- Regarding the term α_1 , it is first seen (using the assumption $\|\mathbf{a}\|_2^2 = 1$) that

$$\mathbb{E}[\alpha_1] \lesssim \sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \sigma_{\max}^2 \\ \asymp n \sigma_{\max}^2.$$

In addition, Bernstein's inequality tells us that with probability exceeding $1 - O(n^{-10})$,

$$\begin{aligned} & |\alpha_1 - \mathbb{E}[\alpha_1]| \\ & \lesssim \sqrt{\sum_{i,j} \text{Var} \left[\left(a_i^2 + a_j^2 + \frac{\mu}{n} \right) H_{ij}^2 \right] \log n} \\ & \quad + \max_{i,j} \left\{ \left(a_i^2 + a_j^2 + \frac{\mu}{n} \right) H_{ij}^2 \right\} \log n \\ & \lesssim \sqrt{\sum_{i,j} \left(a_i^2 + a_j^2 + \frac{\mu}{n} \right)^2 B^2 \sigma_{\max}^2 \log n} \\ & \quad + \max_{i,j} \left(a_i^2 + a_j^2 + \frac{\mu}{n} \right) B^2 \log n \\ & \lesssim \sqrt{\sum_{i,j} \left\{ \left(a_i^2 + a_j^2 \right)^2 + \left(\frac{\mu}{n} \right)^2 \right\} B^2 \sigma_{\max}^2 \log n} \\ & \quad + B^2 \log n \\ & \lesssim \sqrt{\sum_{i,j} \left\{ a_i^2 + a_j^2 + \left(\frac{\mu}{n} \right)^2 \right\} B^2 \sigma_{\max}^2 \log n} \\ & \quad + B^2 \log n \\ & \lesssim B \sigma_{\max} \sqrt{n \log n} + B^2 \log n \lesssim \sigma_{\max}^2 n, \end{aligned}$$

where the fourth line follows since $a_i^2 + a_j^2 \leq \|\mathbf{a}\|_2^2 = 1$, and the last inequality comes from Assumption 3. Consequently,

$$\alpha_1 \leq |\alpha_1 - \mathbb{E}[\alpha_1]| + \mathbb{E}[\alpha_1] \lesssim \sigma_{\max}^2 n. \quad (161)$$

- By virtue of Lemma 9 and the assumption $\|\mathbf{a}\|_2 = 1$, the second term α_2 can be upper bounded by

$$\begin{aligned} \alpha_2 &= \sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \xi^2 \\ &= 2n \left\{ \|\mathbf{a}\|_2^2 + \|\mathbf{u}_l^*\|_2^2 + \frac{1}{\mu \log n} \right\} \xi^2 \\ &\asymp n \xi^2 \lesssim \sigma_{\max}^2 \mu^2 \kappa^8 r^2 \log n. \end{aligned}$$

- When it comes to α_3 , we first make the observation that

$$\begin{aligned} & \xi (|H_{ij}| + \xi) \\ & \lesssim \begin{cases} \mu \xi^2 \log n & \text{if } |H_{ij}| \lesssim \xi \mu \log n \\ \frac{1}{\mu \log n} (|H_{ij}| + \xi)^2 & \text{if } |H_{ij}| \gtrsim \xi \mu \log n \end{cases} \\ & \lesssim \mu \xi^2 \log n + \frac{1}{\mu \log n} (|H_{ij}| + \xi)^2 \mathbb{1}\{|H_{ij}| \gtrsim \xi \mu \log n\} \\ & \asymp \mu \xi^2 \log n + \frac{1}{\mu \log n} H_{ij}^2. \end{aligned}$$

As a consequence, one can derive

$$\begin{aligned} \alpha_3 &\lesssim \mu \log n \\ &\quad \cdot \underbrace{\sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) \xi^2}_{=: \alpha_2} \\ &\quad + \frac{1}{\mu \log n} \\ &\quad \cdot \underbrace{\sum_{i,j} \left(a_i^2 + a_j^2 + |u_{l,i}^*|^2 + |u_{l,j}^*|^2 + \frac{1}{\mu n \log n} \right) H_{ij}^2}_{=: \alpha_1} \\ &= \alpha_2 \mu \log n + \frac{1}{\mu \log n} \alpha_1. \end{aligned}$$

Putting the above bounds together, we conclude that

$$\begin{aligned} & |4\lambda_l^2 \hat{v}_{\mathbf{a},l} - 4\lambda_l^{*2} \tilde{v}_{\mathbf{a},l}| \\ & \lesssim \frac{1}{n\sqrt{\log n}} \alpha_1 + \frac{1}{n\sqrt{\log n}} \alpha_2 + \frac{\mu}{n} \left(\alpha_2 \mu \log n + \frac{1}{\mu \log n} \alpha_1 \right) \\ & \asymp \frac{1}{n\sqrt{\log n}} \alpha_1 + \alpha_2 \frac{\mu^2 \log n}{n} \\ & \lesssim \left(\frac{1}{\sqrt{\log n}} + \frac{\mu^4 \kappa^8 r^2 \log^2 n}{n} \right) \sigma_{\max}^2 \\ & = o(\sigma_{\max}^2), \end{aligned}$$

provided that $\mu^4 \kappa^8 r^2 \log^2 n = o(n)$. This combined with the fact that $\lambda_l = (1 + o(1))\lambda_l^*$ (cf. Corollary 2) and the inequality (151) gives

$$\hat{v}_{\mathbf{a},l} - \tilde{v}_{\mathbf{a},l} = o(v_{\mathbf{a},l}^*).$$

C. Proof of Lemma 8

The randomness in the quantity $\tilde{v}_{\mathbf{a},l}$ purely comes from H_{ij} . Given that $\tilde{v}_{\mathbf{a},l}$ is the sum of independent random variables, it is easily seen that $\mathbb{E}[\tilde{v}_{\mathbf{a},l}] = v_{\mathbf{a},l}^*$. Invoke Bernstein's

inequality [84] to guarantee that with probability at least $1 - O(n^{-10})$,

$$\begin{aligned}
& |4(\tilde{v}_{a,l} - v_{a,l}^*)| \\
&= |4(\tilde{v}_{a,l} - \mathbb{E}[\tilde{v}_{a,l}])| \\
&= \left| \sum_{1 \leq i,j \leq n} (a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)^2 (H_{ij}^2 - \mathbb{E}[H_{ij}^2]) \right| \\
&\lesssim \sqrt{\sum_{i,j} \text{Var}[(a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)^2 H_{ij}^2] \log n} \\
&\quad + \max_{i,j} \{ |(a_{l,i}^\perp u_{l,j}^* + a_{l,j}^\perp u_{l,i}^*)^2 B^2| \} \log n \\
&\stackrel{(i)}{\lesssim} \sqrt{\sum_{i,j} (a_{l,i}^\perp)^4 (u_{l,j}^*)^4 \sigma_{\max}^2 B^2 \log n + \max_{i,j} (a_{l,i}^\perp u_{l,j}^*)^2 B^2 \log n} \\
&\lesssim \sqrt{\left\{ \max_{i,j} (a_{l,i}^\perp)^2 (u_{l,j}^*)^2 \right\} \sum_{i,j} (a_{l,i}^\perp u_{l,j}^*)^2 \sigma_{\max}^2 B^2 \log n} \\
&\quad + \max_{i,j} (a_{l,i}^\perp u_{l,j}^*)^2 B^2 \log n \\
&\stackrel{(ii)}{\lesssim} B \sigma_{\max} \sqrt{\frac{\mu \log n}{n}} + \frac{\mu B^2 \log n}{n} \stackrel{(iii)}{=} o(\sigma_{\max}^2) = o(v_{a,l}^*),
\end{aligned}$$

where the inequality (i) uses the fact that $\text{Var}[H_{ij}^2] \leq \mathbb{E}[H_{ij}^4] \leq B^2 \mathbb{E}[H_{ij}^2] \leq B^2 \sigma_{\max}^2$, (ii) follows since $\max_{i,j} (a_{l,i}^\perp u_{l,j}^*)^2 \leq \|\mathbf{u}^*\|_\infty^2 \leq \mu/n$ and $\|\mathbf{a}_l^\perp\|_2^2 \leq \|\mathbf{a}\|_2^2 = \|\mathbf{u}^*\|_2^2 = 1$, (iii) arises from Assumption 3, and the last identity follows from (151). Combining the above pieces, we obtain $|\tilde{v}_{a,l} - v_{a,l}^*| = o(v_{a,l}^*)$ as claimed.

D. Proof of Lemma 9

For notational convenience, we denote by $\mathbf{M}_{\text{svd}} = \mathbf{U}_{\text{svd}} \Sigma_{\text{svd}} \mathbf{V}_{\text{svd}}^\top$ the compact SVD of $\mathbf{M}_{\text{svd}}$ (or equivalently, the rank- r SVD of $\mathbf{M}$). We shall also define the rotation matrix

$$\mathbf{Q} := \arg \min_{\mathbf{R} \in \mathcal{O}^{r \times r}} \left\| \begin{bmatrix} \mathbf{U}_{\text{svd}} \\ \mathbf{V}_{\text{svd}} \end{bmatrix} \mathbf{R} - \begin{bmatrix} \mathbf{U}^* \\ \mathbf{V}^* \end{bmatrix} \right\|_{\text{F}}, \quad (162)$$

where $\mathcal{O}^{r \times r}$ denotes the set of $r \times r$ orthonormal matrices.

We first record a perturbation bound regarding the singular subspace of $\mathbf{M}_{\text{svd}}$.

Lemma 11: Instate the assumptions of Lemma 9. With probability exceeding $1 - O(n^{-9})$ one has

$$\max \{ \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^*\|_{2,\infty}, \|\mathbf{V}_{\text{svd}} \mathbf{Q} - \mathbf{V}^*\|_{2,\infty} \} \lesssim \kappa^2 \gamma \sqrt{\frac{\mu r}{n}}, \quad (163)$$

where $\gamma \asymp \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*}$.

As immediate consequences of Lemma 11 and Assumption 3, we have

$$\begin{cases} \|\mathbf{U}_{\text{svd}}\|_{2,\infty} &\leq \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^*\|_{2,\infty} + \|\mathbf{U}^*\|_{2,\infty} \lesssim \kappa \sqrt{\frac{\mu r}{n}}, \\ \|\mathbf{V}_{\text{svd}}\|_{2,\infty} &\leq \|\mathbf{V}_{\text{svd}} \mathbf{Q} - \mathbf{V}^*\|_{2,\infty} + \|\mathbf{V}^*\|_{2,\infty} \lesssim \kappa \sqrt{\frac{\mu r}{n}}, \end{cases} \quad (164)$$

provided that $\kappa^2 \sqrt{\mu r/n} \leq 1$.

We are now ready to control the entrywise error of $\mathbf{M}_{\text{svd}}$. Towards this, we make note of the following bound

$$\begin{aligned}
& \|\mathbf{M}_{\text{svd}} - \mathbf{M}^*\|_\infty \\
& \lesssim \|\mathbf{Q}^\top \Sigma_{\text{svd}} \mathbf{Q} - \Lambda^*\| \cdot \|\mathbf{U}_{\text{svd}}\|_{2,\infty} \|\mathbf{V}_{\text{svd}}\|_{2,\infty} \\
& \quad + \|\Lambda^*\| (\|\mathbf{U}_{\text{svd}}\|_{2,\infty} + \|\mathbf{V}^*\|_{2,\infty}) \\
& \quad \cdot (\|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^*\|_{2,\infty} + \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^*\|_{2,\infty}), \quad (165)
\end{aligned}$$

which can be obtained by combining the inequalities (C.17)-(C.18) in [9]. In addition, following the argument in [9, Appendix C.3.3], we obtain

$$\|\mathbf{Q}^\top \Sigma_{\text{svd}} \mathbf{Q} - \Lambda^*\| \lesssim \gamma \lambda_{\min}^*, \quad (166)$$

where γ is defined in (168a). Substituting (163), (164) and (166) into (165) and combining terms, we reach

$$\begin{aligned}
& \|\mathbf{M}_{\text{svd}} - \mathbf{M}^*\|_\infty \\
& \lesssim \gamma \cdot \frac{\mu \kappa^2 r}{n} \lambda_{\min}^* + \gamma \cdot \frac{\mu \kappa^4 r}{n} \lambda_{\min}^* \asymp \gamma \cdot \frac{\mu \kappa^4 r}{n} \lambda_{\min}^* \\
& \asymp \sigma_{\max} \sqrt{\frac{\mu^2 \kappa^8 r^2 \log n}{n}}. \quad (167)
\end{aligned}$$

Finally, it follows from our construction that $\widehat{\mathbf{H}} - \mathbf{H} = \mathbf{M} - \mathbf{M}_{\text{svd}} - \mathbf{H} = \mathbf{M}^* - \mathbf{M}_{\text{svd}}$, which combined with (167) establishes this lemma.

Proof of Lemma 11: In order to invoke [9, Theorem 2.1] and the symmetric dilation trick in [9, Section 3.3], we need to first verify the assumptions required in [9, Section 2.1]. To this end, we introduce the following auxiliary quantity and function

$$\gamma := c_\gamma \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \quad (168a)$$

$$\varphi(x) := \begin{cases} c_\varphi \left(\frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} x + \frac{B \log n}{\lambda_{\min}^*} \right), & \text{if } \frac{1}{\sqrt{n}} \leq x \leq 1 \\ c_\varphi \left(\frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} + \frac{B \sqrt{n \log n}}{\lambda_{\min}^*} \right) x, & \text{if } 0 \leq x < \frac{1}{\sqrt{n}} \end{cases} \quad (168b)$$

for some sufficiently large constants $c_\gamma, c_\varphi > 0$. We then make the following observations:

- To begin with, the two-to-infinity norm of the truth can be bounded by

$$\begin{aligned}
\|\mathbf{M}^*\|_{2,\infty} &= \|\mathbf{U}^* \Lambda^* \mathbf{U}^{*\top}\|_{2,\infty} \leq \|\mathbf{U}^*\|_{2,\infty} \cdot \|\Lambda^*\| \cdot \|\mathbf{U}^*\| \\
&\leq \lambda_{\max}^* \sqrt{\frac{\mu r}{n}} = \lambda_{\min}^* \sqrt{\frac{\mu \kappa^2 r}{n}}. \quad (169)
\end{aligned}$$

- In view of the matrix Bernstein inequality [84], with probability exceeding $1 - O(n^{-6})$ one has

$$\begin{aligned}
\|\mathbf{M} - \mathbf{M}^*\| &= \|\mathbf{H}\| \lesssim \sigma_{\max} \sqrt{n \log n} \\
&= \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} \lambda_{\min}^* \leq \gamma \lambda_{\min}^*. \quad (170)
\end{aligned}$$

- Denote by $\mathbf{H}_i$, the i th row of $\mathbf{H}$. For any fixed $\mathbf{W} \in \mathbb{R}^{n \times r}$, the matrix Bernstein inequality yields

$$\begin{aligned} \|\mathbf{H}_i, \mathbf{W}\|_2 &\lesssim \sqrt{\sum_{j=1}^n \sigma_{i,j}^2 \|\mathbf{W}_{j,\cdot}\|_2^2 \log n + B \|\mathbf{W}\|_{2,\infty} \log n} \\ &\lesssim \sigma_{\max} \|\mathbf{W}\|_{\text{F}} \sqrt{\log n} + B \|\mathbf{W}\|_{2,\infty} \log n \\ &= \lambda_{\min}^* \|\mathbf{W}\|_{2,\infty} \cdot \left\{ \frac{\|\mathbf{W}\|_{\text{F}}}{\sqrt{n} \|\mathbf{W}\|_{2,\infty}} \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} + \frac{B \log n}{\lambda_{\min}^*} \right\} \end{aligned} \quad (171)$$

with probability exceeding $1 - O(n^{-10})$. This combined with (168b) and $\frac{1}{\sqrt{n}} \leq \frac{\|\mathbf{W}\|_{\text{F}}}{\sqrt{n} \|\mathbf{W}\|_{2,\infty}} \leq 1$ gives

$$\begin{aligned} \mathbb{P} \left\{ \|\mathbf{H}_i, \mathbf{W}\|_2 \geq \lambda_{\min}^* \|\mathbf{W}\|_{2,\infty} \varphi \left(\frac{\|\mathbf{W}\|_{\text{F}}}{\sqrt{n} \|\mathbf{W}\|_{2,\infty}} \right) \right\} \\ \geq 1 - O(n^{-10}). \end{aligned} \quad (172)$$

With the above observations in place, [9, Theorem 2.1] together with the dilation trick in [9, Section 3.3] implies that: with probability at least $1 - O(n^{-5})$, one has

$$\begin{aligned} \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^* - \mathbf{H} \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} \\ = \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{M} \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} \\ \lesssim \kappa^2 \gamma \|\mathbf{U}^*\|_{2,\infty} + \gamma \|\mathbf{M}^*\|_{2,\infty} / \lambda_{\min}^* \\ \lesssim \gamma \kappa^2 \sqrt{\frac{\mu r}{n}} + \gamma \sqrt{\frac{\mu \kappa^2 r}{n}} \asymp \gamma \sqrt{\frac{\mu \kappa^4 r}{n}}, \end{aligned} \quad (173)$$

where the last line arises from the incoherence condition as well as the inequality (169). In addition,

$$\begin{aligned} \|\mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} &\leq \|\mathbf{U}^*\|_{2,\infty} \|(\mathbf{\Lambda}^*)^{-1}\| \leq \frac{1}{\lambda_{\min}^*} \sqrt{\frac{\mu r}{n}}, \\ \|\mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{\text{F}} &\leq \|\mathbf{U}^*\|_{\text{F}} \|(\mathbf{\Lambda}^*)^{-1}\| \leq \frac{\sqrt{r}}{\lambda_{\min}^*}, \end{aligned}$$

which taken collectively with (171) (by setting $\mathbf{W} = \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}$) demonstrate that

$$\begin{aligned} \|\mathbf{H} \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} &\lesssim \frac{\sigma_{\max} \sqrt{r \log n}}{\lambda_{\min}^*} + \frac{B}{\lambda_{\min}^*} \sqrt{\frac{\mu r \log^2 n}{n}} \\ &\lesssim \frac{\sigma_{\max} \sqrt{\mu r \log n}}{\lambda_{\min}^*} \end{aligned} \quad (174)$$

with probability $1 - O(n^{-9})$. Taking it together with (173) and using the triangle inequality immediately yield

$$\begin{aligned} \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^*\|_{2,\infty} \\ \leq \|\mathbf{U}_{\text{svd}} \mathbf{Q} - \mathbf{U}^* - \mathbf{H} \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} + \|\mathbf{H} \mathbf{U}^* (\mathbf{\Lambda}^*)^{-1}\|_{2,\infty} \\ \lesssim \gamma \sqrt{\frac{\mu \kappa^4 r}{n}} + \frac{\sigma_{\max} \sqrt{\mu r \log n}}{\lambda_{\min}^*} \asymp \gamma \sqrt{\frac{\mu \kappa^4 r}{n}} \end{aligned}$$

as claimed. The part concerning $\mathbf{V}_{\text{svd}}$ follows from the same argument. $\square$

Remark 11: The careful reader would note that the results in [9] require another relation between the incoherence condition and γ (namely, Assumption A1). This relation cannot be satisfied with the current choice of γ given in (168a), unless we increase it to

$$\gamma = c_\gamma \frac{\sigma_{\max} \sqrt{n \log n}}{\lambda_{\min}^*} + \sqrt{\frac{\mu \kappa^2 r}{n}}. \quad (175)$$

Fortunately, the second term of (175) can be removed with slightly more refined analysis, provided that Assumption 3 holds. In short, the analysis of [9] is built upon a sequence of auxiliary leave-one-out estimates $\mathbf{M}^{(m)}$ ($1 \leq m \leq n$) — obtained by zeroing out the m th row/column of $\mathbf{M}$ — which allows us to approximate $\mathbf{M}$ while being statistically independent of the data in the m th row/column. However, the approximation error $\mathbf{M}^{(m)} - \mathbf{M}$ does not decrease to zero even if $\mathbf{H} = \mathbf{0}$, leading to a bias term that is non-vanishing as $\sigma_{\max} \rightarrow 0$. To address this issue, it suffices to replace $\mathbf{M}^{(m)}$ by $\tilde{\mathbf{M}}^{(m)}$, where $\tilde{\mathbf{M}}^{(m)}$ is obtained by replacing all entries in the m th row/column of $\mathbf{M}$ by their expected values. This allows us to ensure that $\tilde{\mathbf{M}}^{(m)} - \mathbf{M} \rightarrow \mathbf{0}$ as $\sigma_{\max} \rightarrow 0$, which in turn leads to the removal of the second term of (175). The refined proof is nearly identical to the original proof in [9], and is hence omitted here for the sake of brevity.

APPENDIX E

PROOF FOR MINIMAX LOWER BOUNDS (THEOREM 3)

In the following, we prove each lower bound respectively.

A. Proof of Eqs. (21a) and (21b)

Without loss of generality, consider any $1 \leq l \leq r$ and any $k \neq l$. Consider the following hypotheses regarding the eigen-decomposition of $\mathbf{M}^* \in \mathbb{R}^{n \times n}$:

$$\begin{aligned} \mathcal{H}_0 : \mathbf{M} &= \mathbf{M}^* + \mathbf{H} = \sum_{j=1}^r \lambda_j^* \mathbf{u}_j^* \mathbf{u}_j^{*\top} + \mathbf{H}; \\ \mathcal{H}_k : \mathbf{M} &= \mathbf{M}_k + \mathbf{H} = \lambda_l^* \tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top + \lambda_k^* \tilde{\mathbf{u}}_k \tilde{\mathbf{u}}_k^\top \\ &\quad + \sum_{j:j \neq k, j \neq l} \lambda_j^* \mathbf{u}_j^* \mathbf{u}_j^{*\top} + \mathbf{H}. \end{aligned}$$

In words, $\mathcal{H}_k$ is obtained by perturbing the l th and the k th eigenvectors under $\mathcal{H}_0$, with the remaining eigenvectors unaltered. In particular, for any $k \neq l$, we shall pick $\tilde{\mathbf{u}}_l$ and $\tilde{\mathbf{u}}_k$ such that they are equivalent to $\mathbf{u}_l^*$ and $\mathbf{u}_k^*$ modulo global rotation, namely,

$$\mathbf{u}_l^* \mathbf{u}_l^{*\top} + \mathbf{u}_k^* \mathbf{u}_k^{*\top} = \tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top + \tilde{\mathbf{u}}_k \tilde{\mathbf{u}}_k^\top; \quad (176)$$

as we shall see, this rotational invariance constraint (176) plays a pivotal role in understanding the effect of the eigen-gap upon estimation accuracy. In what follows, we let $\mathbb{P}_0$ (resp. $\mathbb{P}_k$) denote the distribution of $\mathbf{M}$ under $\mathcal{H}_0$ (resp. $\mathcal{H}_k$), and let $\mathbb{P}_{0,i,j}$ (resp. $\mathbb{P}_{k,i,j}$) stand for the distribution of M_{ij} under $\mathcal{H}_0$ (resp. $\mathcal{H}_k$).

g) *Step 1: calculation of KL divergence:* We first calculate the KL divergence of $\mathbb{P}_0$ from $\mathbb{P}_k$ as follows

$$\begin{aligned} \text{KL}(\mathbb{P}_k \parallel \mathbb{P}_0) &\stackrel{(i)}{=} \sum_{1 \leq i, j \leq n} \text{KL}(\mathbb{P}_{k,i,j} \parallel \mathbb{P}_{0,i,j}) \\ &\stackrel{(ii)}{=} \sum_{1 \leq i, j \leq n} \frac{(M_{ij}^* - (M_k)_{ij})^2}{2\sigma_{ij}^2} \leq \frac{\|M^* - M_k\|_F^2}{2\sigma_{\min}^2}. \end{aligned} \quad (177)$$

Here, (i) holds since KL divergence is additive for independent distributions [85, Chapter 2.4], and (ii) follows since

$$\text{KL}(\mathcal{N}(\mu_1, \sigma^2) \parallel \mathcal{N}(\mu_2, \sigma^2)) = \frac{(\mu_1 - \mu_2)^2}{2\sigma^2}.$$

In addition, a little algebra reveals that

$$\begin{aligned} M^* - M_k &= \lambda_k^* (\mathbf{u}_l^* \mathbf{u}_l^{*\top} + \mathbf{u}_k^* \mathbf{u}_k^{*\top}) + (\lambda_l^* - \lambda_k^*) \mathbf{u}_l^* \mathbf{u}_l^{*\top} \\ &\quad - \{\lambda_k^* (\tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top + \tilde{\mathbf{u}}_k \tilde{\mathbf{u}}_k^\top) + (\lambda_l^* - \lambda_k^*) \tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top\} \\ &= (\lambda_l^* - \lambda_k^*) (\mathbf{u}_l^* \mathbf{u}_l^{*\top} - \tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top), \end{aligned}$$

where the last relation makes use of the condition (176). Therefore, we continue the bound (177) to reach

$$\begin{aligned} \text{KL}(\mathbb{P}_k \parallel \mathbb{P}_0) &\leq \frac{\|M^* - M_k\|_F^2}{2\sigma_{\min}^2} \\ &= \frac{(\lambda_l^* - \lambda_k^*)^2}{2\sigma_{\min}^2} \|\mathbf{u}_l^* \mathbf{u}_l^{*\top} - \tilde{\mathbf{u}}_l \tilde{\mathbf{u}}_l^\top\|_F^2 \\ &\leq \frac{(\lambda_l^* - \lambda_k^*)^2}{\sigma_{\min}^2} \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2^2, \end{aligned} \quad (178)$$

where the last line holds due to the following inequality that holds for any unit vectors $\mathbf{u}$ and $\mathbf{v}$:

$$\begin{aligned} \|\mathbf{u}\mathbf{u}^\top - \mathbf{v}\mathbf{v}^\top\|_F^2 &= \|\mathbf{u}\|_2^4 + \|\mathbf{v}\|_2^4 - 2\langle \mathbf{u}\mathbf{u}^\top, \mathbf{v}\mathbf{v}^\top \rangle = 2 - 2(\mathbf{u}^\top \mathbf{v})^2 \\ &= (2 - 2\mathbf{u}^\top \mathbf{v})(1 + \mathbf{u}^\top \mathbf{v}) = \frac{1}{2} \|\mathbf{u} - \mathbf{v}\|_2^2 \cdot \|\mathbf{u} + \mathbf{v}\|_2^2 \\ &\leq 2\|\mathbf{u} - \mathbf{v}\|_2^2. \end{aligned} \quad (179)$$

h) *Step 2: bounding minimax probability of error:* Define the minimax probability of error as follows

$$\begin{aligned} p_{e,k} &:= \inf_{\psi} \max \left\{ \mathbb{P}\{\psi \text{ rejects } \mathcal{H}_0 \mid \mathcal{H}_0\}, \mathbb{P}\{\psi \text{ rejects } \mathcal{H}_k \mid \mathcal{H}_k\} \right\}, \end{aligned} \quad (180)$$

where the infimum is over all tests. Standard minimax lower bounds [85, Theorem 2] tell us that: if

$$\text{KL}(\mathbb{P}_k \parallel \mathbb{P}_0) \leq 1/16,$$

then one necessarily has $p_{e,k} \geq 1/5$. This taken collectively with the upper bound (178) implies that: if $\frac{(\lambda_l^* - \lambda_k^*)^2}{\sigma_{\min}^2} \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2^2 \leq 1/16$, or equivalently, if

$$\|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \leq \frac{\sigma_{\min}}{4|\lambda_l^* - \lambda_k^*|},$$

then the minimax probability of error $p_{e,k}$ is lower bounded by $1/5$.

i) *Step 3(a): establishing minimax ℓ_2 lower bounds:* The above minimax probability of testing error has direct implications on ℓ_2 eigenvector estimation. Suppose that $\|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 = \frac{\sigma_{\min}}{4|\lambda_l^* - \lambda_k^*|}$, then the calculation in (178) indicates that

$$\|M^* - M_k\|_F \leq \sqrt{2} |\lambda_l^* - \lambda_k^*| \cdot \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 < \frac{\sigma_{\min}}{2},$$

and hence $M_k \in \mathcal{M}(M^*)$. Moreover, when $\sigma_{\min} \leq |\lambda_l^* - \lambda_k^*|$, one has

$$\begin{aligned} \|\mathbf{u}_l^* + \tilde{\mathbf{u}}_l\|_2 &= \|2\mathbf{u}_l^* - (\mathbf{u}_l^* - \tilde{\mathbf{u}}_l)\|_2 \geq \|2\mathbf{u}_l^*\|_2 - \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \\ &\geq 2 - \frac{\sigma_{\min}}{4|\lambda_l^* - \lambda_k^*|} \\ &> \frac{\sigma_{\min}}{4|\lambda_l^* - \lambda_k^*|} = \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2, \end{aligned}$$

meaning that $\|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 = \min \|\mathbf{u}_l^* \pm \tilde{\mathbf{u}}_l\|_2$. Thus, the standard reduction scheme described in [85, Chapter 2.2] leads to

$$\begin{aligned} \inf_{\tilde{\mathbf{u}}_l} \sup_{\mathbf{A} \in \mathcal{M}_0(M^*)} \mathbb{E} \left[\min \|\hat{\mathbf{u}}_l \pm \mathbf{u}_l(\mathbf{A})\|_2 \right] &\gtrsim p_{e,k} \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \\ &\asymp \frac{\sigma_{\min}}{|\lambda_l^* - \lambda_k^*|}, \end{aligned}$$

where the infimum is taken over all estimator for $\mathbf{u}_l(\mathbf{A})$. Since the preceding bound holds for all $k \neq l$, we conclude that

$$\begin{aligned} \inf_{\tilde{\mathbf{u}}_l} \sup_{\mathbf{A} \in \mathcal{M}_0(M^*)} \mathbb{E} \left[\min \|\hat{\mathbf{u}}_l \pm \mathbf{u}_l(\mathbf{A})\|_2 \right] &\gtrsim \max_{k: k \neq l} \frac{\sigma_{\min}}{|\lambda_l^* - \lambda_k^*|} \\ &= \frac{\sigma_{\min}}{\Delta_l^*}. \end{aligned}$$

j) *Step 3(b): establishing minimax lower bounds on estimating linear functionals of eigenvectors:* The preceding minimax probability of error also has direct implications on estimating linear functionals of eigenvectors. In order to satisfy the rotational invariance constraint (176), we set

$$[\tilde{\mathbf{u}}_l, \tilde{\mathbf{u}}_k] = [\mathbf{u}_l^*, \mathbf{u}_k^*] \begin{bmatrix} \cos \theta_k & \sin \theta_k \\ -\sin \theta_k & \cos \theta_k \end{bmatrix}$$

for some $\theta_k \in [-\pi/2, \pi/2]$. Before continuing, we shall also make precise the connection between $\|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2$ and θ_k . Specifically, the above relation allows one to derive

$$\mathbf{u}_l^* - \tilde{\mathbf{u}}_l = \mathbf{u}_l^* (1 - \cos \theta_k) + \mathbf{u}_k^* \sin \theta_k = 2\mathbf{u}_l^* \sin^2 \frac{\theta_k}{2} + \mathbf{u}_k^* \sin \theta_k$$

and, as a consequence,

$$\begin{cases} \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \leq (1 - \cos \theta_k) + |\sin \theta_k| = 2 \sin^2 \frac{\theta_k}{2} + |\sin \theta_k| \\ \leq 3|\theta_k|, \\ \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \geq |\sin \theta_k| \geq \frac{2}{\pi} |\theta_k|. \end{cases} \quad (181)$$

In what follows, we shall take $\text{sign}(\theta_k) = \text{sign}(\frac{\mathbf{a}^\top \mathbf{u}_k^*}{\mathbf{a}^\top \mathbf{u}_l^*})$, and generate the magnitude $|\theta_k|$ as follows

$$|\theta_k| \sim \text{Uniform} \left(\left[\frac{\sigma_{\min}}{120|\lambda_l^* - \lambda_k^*|}, \frac{\sigma_{\min}}{12|\lambda_l^* - \lambda_k^*|} \right] \right), \quad (182)$$

which combined with (181) guarantees that

$$\frac{\sigma_{\min}}{60\pi |\lambda_l^* - \lambda_k^*|} \leq \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2 \leq \frac{\sigma_{\min}}{4|\lambda_l^* - \lambda_k^*|}. \quad (183)$$

We aim to translate the difficulty in distinguishing $\mathcal{H}_0$ and $\mathcal{H}_k$ into a fundamental lower bound on estimating the linear functional. Towards this, we are in need of computing the difference of the linear functional under these two hypotheses (namely, $\mathbf{a}^\top \mathbf{u}_l^* - \mathbf{a}^\top \tilde{\mathbf{u}}_l$). The above expressions give

$$\begin{aligned} & |\mathbf{a}^\top \mathbf{u}_l^* - \mathbf{a}^\top \tilde{\mathbf{u}}_l| \\ &= |2(\mathbf{a}^\top \mathbf{u}_l^*) \sin^2 \frac{\theta_k}{2} + (\mathbf{a}^\top \mathbf{u}_k^*) \sin \theta_k| \\ &\stackrel{(i)}{=} 2|\mathbf{a}^\top \mathbf{u}_l^*| \sin^2 \frac{\theta_k}{2} + |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\sin \theta_k| \\ &\asymp \theta_k^2 |\mathbf{a}^\top \mathbf{u}_l^*| + |\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*|, \end{aligned} \quad (184)$$

where the identity (i) results from the condition $\text{sign}(\theta_k) = \text{sign}(\frac{\mathbf{a}^\top \mathbf{u}_k^*}{\mathbf{a}^\top \mathbf{u}_l^*})$.

We shall also control $|\mathbf{a}^\top \mathbf{u}_l^* + \mathbf{a}^\top \tilde{\mathbf{u}}_l|$, for which we divide into two cases:

- If $|\mathbf{a}^\top \mathbf{u}_l^*| \geq \frac{\sigma_{\min}}{|\lambda_l^* - \lambda_k^*|} |\mathbf{a}^\top \mathbf{u}_k^*| \geq 12|\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*|$, then the identity $\mathbf{u}_l^* + \tilde{\mathbf{u}}_l = 2\mathbf{u}_l^* - (\mathbf{u}_l^* - \tilde{\mathbf{u}}_l)$ together with (184) yields

$$\begin{aligned} & |\mathbf{a}^\top \mathbf{u}_l^* + \mathbf{a}^\top \tilde{\mathbf{u}}_l| \\ &= |2|\mathbf{a}^\top \mathbf{u}_l^*| - 2|\mathbf{a}^\top \mathbf{u}_l^*| \sin^2 \frac{\theta_k}{2} - |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\sin \theta_k| \\ &\geq 2|\mathbf{a}^\top \mathbf{u}_l^*| \cos^2 \frac{\theta_k}{2} - |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\theta_k| \\ &\geq |\mathbf{a}^\top \mathbf{u}_l^*| - |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\theta_k| \geq \frac{1}{2} |\mathbf{a}^\top \mathbf{u}_l^*| \\ &\gtrsim \theta_k^2 |\mathbf{a}^\top \mathbf{u}_l^*| + |\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*|, \end{aligned}$$

where we have used the fact that $\sigma_{\min} \leq |\lambda_l^* - \lambda_k^*|$ (so that $|\theta_k| \leq \frac{\sigma_{\min}}{12|\lambda_l^* - \lambda_k^*|} \leq \frac{1}{12}$ and hence $\cos^2 \frac{\theta_k}{2} \geq \frac{1}{2}$).

- Suppose now that $|\mathbf{a}^\top \mathbf{u}_l^*| < \frac{\sigma_{\min}}{|\lambda_l^* - \lambda_k^*|} |\mathbf{a}^\top \mathbf{u}_k^*|$. As one can easily verify (which we omit for brevity), the scheme (182) guarantees that with probability exceeding $1/2$, one has

$$\begin{aligned} & |\mathbf{a}^\top \mathbf{u}_l^* + \mathbf{a}^\top \tilde{\mathbf{u}}_l| \\ &= |2|\mathbf{a}^\top \mathbf{u}_l^*| - 2|\mathbf{a}^\top \mathbf{u}_l^*| \sin^2 \frac{\theta_k}{2} - |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\sin \theta_k| \\ &\gtrsim \max\{|\mathbf{a}^\top \mathbf{u}_l^*|, |\mathbf{a}^\top \mathbf{u}_k^*| \cdot |\theta_k|\} \\ &\gtrsim \theta_k^2 |\mathbf{a}^\top \mathbf{u}_l^*| + |\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*|. \end{aligned}$$

Putting the above cases together reveals that

$$|\mathbf{a}^\top \mathbf{u}_l^* + \mathbf{a}^\top \tilde{\mathbf{u}}_l| \gtrsim \theta_k^2 |\mathbf{a}^\top \mathbf{u}_l^*| + |\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*|$$

with probability exceeding $1/2$. Consequently, one can find $|\theta_k| \in \left[\frac{\sigma_{\min}}{120|\lambda_l^* - \lambda_k^*|}, \frac{\sigma_{\min}}{12|\lambda_l^* - \lambda_k^*|} \right]$ such that

$$\begin{aligned} & \min |\mathbf{a}^\top \mathbf{u}_l^* \pm \mathbf{a}^\top \tilde{\mathbf{u}}_l| \\ &\gtrsim \theta_k^2 |\mathbf{a}^\top \mathbf{u}_l^*| + |\theta_k| \cdot |\mathbf{a}^\top \mathbf{u}_k^*| \\ &\asymp |\mathbf{a}^\top \mathbf{u}_l^*| \frac{\sigma_{\min}^2}{|\lambda_l^* - \lambda_k^*|^2} + |\mathbf{a}^\top \mathbf{u}_k^*| \frac{\sigma_{\min}}{|\lambda_l^* - \lambda_k^*|}, \end{aligned}$$

where we recall that $\min |a \pm b| := \min\{|a - b|, |a + b|\}$.

Suppose for the moment that $\mathbf{M}_k \in \mathcal{M}_0(\mathbf{M}^*)$. Applying the standard reduction scheme [85, Chapter 2.2] once again yields

$$\begin{aligned} & \inf_{\hat{\mathbf{u}}_{a,l}} \sup_{\mathbf{A} \in \mathcal{M}_0(\mathbf{M}^*)} \mathbb{E} \left[\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l(\mathbf{A})| \right] \\ &\gtrsim p_{e,k} \min |\mathbf{a}^\top \mathbf{u}_l^* \pm \mathbf{a}^\top \tilde{\mathbf{u}}_l| \gtrsim \frac{\sigma_{\min}^2 |\mathbf{a}^\top \mathbf{u}_l^*|}{|\lambda_l^* - \lambda_k^*|^2} + \frac{\sigma_{\min} |\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|}, \end{aligned}$$

where the infimum is taken over all estimators for $\mathbf{a}^\top \mathbf{u}_l(\mathbf{A})$. Recognizing that the preceding inequality holds for all $k \neq l$, we immediately arrive at the advertised claim

$$\begin{aligned} & \inf_{\hat{\mathbf{u}}_{a,l}} \sup_{\mathbf{A} \in \mathcal{M}_0(\mathbf{M}^*)} \mathbb{E} \left[\min |\hat{\mathbf{u}}_{a,l} \pm \mathbf{a}^\top \mathbf{u}_l(\mathbf{A})| \right] \\ &\gtrsim \sigma_{\min}^2 \max_{k:k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_l^*|}{|\lambda_l^* - \lambda_k^*|^2} + \sigma_{\min} \max_{k:k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|} \\ &= |\mathbf{a}^\top \mathbf{u}_l^*| \frac{\sigma_{\min}^2}{\Delta_l^{*2}} + \sigma_{\min} \max_{k:k \neq l} \frac{|\mathbf{a}^\top \mathbf{u}_k^*|}{|\lambda_l^* - \lambda_k^*|}. \end{aligned}$$

Finally, it remains to justify that $\mathbf{M}_k \in \mathcal{M}_0(\mathbf{M}^*)$ for all $k \neq l$. When $|\theta_k| \leq \frac{\sigma_{\min}}{12|\lambda_l^* - \lambda_k^*|}$, invoking (178) and (181) yields

$$\begin{aligned} \|\mathbf{M}^* - \mathbf{M}_k\|_{\text{F}}^2 &\leq 2(\lambda_l^* - \lambda_k^*)^2 \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_l\|_2^2 \\ &\leq 18(\lambda_l^* - \lambda_k^*)^2 |\theta_k|^2 M < \sigma_{\min}^2/4. \end{aligned}$$

This means that $\mathbf{M}_k \in \mathcal{M}_0(\mathbf{M}^*)$, thus concluding the proof.

B. Proof of Eq. (21c)

Consider the following two hypotheses

$$\begin{aligned} \mathcal{H}_0 : \quad \mathbf{M} &= \mathbf{M}_0 + \mathbf{H} := \lambda_l^* \mathbf{u}_l^* \mathbf{u}_l^{*\top} + \sum_{j:j \neq l} \lambda_j^* \mathbf{v}_j \mathbf{v}_j^\top + \mathbf{H}, \\ \mathcal{H}_a : \quad \mathbf{M} &= \mathbf{M}_a + \mathbf{H} := \lambda_l^* \tilde{\mathbf{u}}_a \tilde{\mathbf{u}}_a^\top + \sum_{j:j \neq l} \lambda_j^* \mathbf{v}_j \mathbf{v}_j^\top + \mathbf{H}, \end{aligned}$$

where the $\mathbf{v}_j$'s are orthonormal vectors obeying $\langle \mathbf{v}_j, \mathbf{a} \rangle = \langle \mathbf{v}_j, \mathbf{u}_l^* \rangle = 0$ for any $j \neq l$, and $\tilde{\mathbf{u}}_a$ is defined as

$$\begin{aligned} \tilde{\mathbf{u}}_a &:= \frac{1}{\|\mathbf{u}_l^* + \delta \mathbf{a}_\perp\|_2} (\mathbf{u}_l^* + \delta \mathbf{a}_\perp), \quad \text{with} \\ \mathbf{a}_\perp &= \mathbf{a} - (\mathbf{a}^\top \mathbf{u}_l^*) \mathbf{u}_l^* \quad \text{and} \quad \delta = \frac{\sigma_{\min}}{12|\lambda_l^*| \cdot \|\mathbf{a}_\perp\|_2}. \end{aligned}$$

Recognizing the simple fact $\langle \mathbf{a}_\perp, \mathbf{u}_l^* \rangle = 0$, we can derive

$$\begin{aligned} \|\mathbf{u}_l^* + \delta \mathbf{a}_\perp\|_2 &= \sqrt{1 + \delta^2 \|\mathbf{a}_\perp\|_2^2} \quad \text{and} \\ \tilde{\mathbf{u}}_a &:= \frac{1}{\sqrt{1 + \delta^2 \|\mathbf{a}_\perp\|_2^2}} (\mathbf{u}_l^* + \delta \mathbf{a}_\perp). \end{aligned}$$

Our proof proceeds as follows. Without loss of generality, we shall assume that $\mathbf{a}^\top \mathbf{u}^* \geq 0$.

- Let $\mathbb{P}_0$ (resp. $\mathbb{P}_a$) denote the distribution of $\mathbf{M}$ under $\mathcal{H}_0$ (resp. $\mathcal{H}_a$). Repeating the derivation in (177) gives

$$\begin{aligned} \text{KL}(\mathbb{P}_a \| \mathbb{P}_0) &\leq \frac{\|\mathbf{M}_0 - \mathbf{M}_a\|_{\text{F}}^2}{2\sigma_{\min}^2} \\ &= \frac{(\lambda_l^*)^2}{2\sigma_{\min}^2} \|\mathbf{u}_l^* \mathbf{u}_l^{*\top} - \tilde{\mathbf{u}}_a \tilde{\mathbf{u}}_a^\top\|_{\text{F}}^2 \\ &\leq \frac{(\lambda_l^*)^2}{\sigma_{\min}^2} \|\mathbf{u}_l^* - \tilde{\mathbf{u}}_a\|_2^2, \end{aligned} \quad (185)$$

where the last inequality comes from (179). In addition,

$$\begin{aligned}
& \|u_l^* - \tilde{u}_a\|_2 \\
&= \left\| u_l^* - \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} (u_l^* + \delta a_\perp) \right\|_2 \\
&\leq \|u_l^* - (u_l^* + \delta a_\perp)\|_2 + \left(1 - \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}}\right) \\
&\quad \cdot \|u_l^* + \delta a_\perp\|_2 \\
&\leq \delta \|a_\perp\|_2 + \frac{\sqrt{1 + \delta^2 \|a_\perp\|_2^2} - 1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} \cdot (1 + \delta \|a_\perp\|_2) \\
&\leq \delta \|a_\perp\|_2 + (1 + \delta \|a_\perp\|_2) \delta \|a_\perp\|_2 \leq 3\delta \|a_\perp\|_2,
\end{aligned} \tag{186}$$

where we have made use of the fact that $\delta \|a_\perp\|_2 = \frac{\sigma_{\min}}{12|\lambda_l^*|} \leq \frac{1}{12}$. Combining this with (185) and our choice of δ , we arrive at

$$\text{KL}(\mathbb{P}_a \parallel \mathbb{P}_0) \leq \frac{9\delta^2 (\lambda_l^*)^2 \|a_\perp\|_2^2}{\sigma_{\min}^2} = \frac{1}{16}. \tag{187}$$

- Define the minimax probability of error as follows

$p_{e,a}$

$$:= \inf_{\psi} \max \left\{ \mathbb{P}\{\psi \text{ rejects } \mathcal{H}_0 \mid \mathcal{H}_0\}, \mathbb{P}\{\psi \text{ rejects } \mathcal{H}_a \mid \mathcal{H}_a\} \right\},$$

where the infimum is over all tests. It then follows from [85, Theorem 2] that: if

$$\text{KL}(\mathbb{P}_k \parallel \mathbb{P}_0) \leq 1/16,$$

one necessarily has $p_{e,a} \geq 1/5$. This taken collectively with the upper bound (187) implies that: the minimax probability of error $p_{e,a}$ in distinguishing $\mathcal{H}_0$ and $\mathcal{H}_a$ is indeed lower bounded by $1/5$.

- Combining (186) with the fact $\delta \|a_\perp\|_2 = \frac{\sigma_{\min}}{12|\lambda_l^*|}$ gives

$$\|\tilde{u}_a - u_l^*\|_2 \leq \frac{\sigma_{\min}}{4|\lambda_l^*|},$$

thus indicating that $M_a \in \mathcal{M}_1(M^*)$. Apply the standard reduction scheme [85, Chapter 2.2] to yield

$$\inf_{\hat{u}_{a,l}} \sup_{A \in \mathcal{M}_1(M^*)} \mathbb{E}[\min |\hat{u}_{a,l} \pm a^\top u_l(A)|] \gtrsim \min |a^\top (u_l^* \pm \tilde{u}_a)|. \tag{188}$$

Everything then boils down to lower bounding $\min |a^\top (u_l^* \pm \tilde{u}_a)|$. On the one hand, it is seen that

$$\begin{aligned}
& |a^\top (u_l^* - \tilde{u}_a)| \\
&= \left| a^\top u_l^* - \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} (a^\top u_l^* + \delta a^\top a_\perp) \right| \\
&\geq \frac{\delta |a^\top a_\perp|}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} - |a^\top u_l^*| \left(1 - \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}}\right) \\
&\geq \frac{1}{2} \delta \|a_\perp\|_2^2 - |a^\top u_l^*| \cdot \delta^2 \|a_\perp\|_2^2 \\
&\geq \frac{1}{4} \delta \|a_\perp\|_2^2.
\end{aligned} \tag{189}$$

On the other hand, one can employ the assumption $a^\top u^* \geq 0$ to derive

$$\begin{aligned}
& |a^\top (u_l^* + \tilde{u}_a)| \\
&= \left| a^\top u_l^* + \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} (a^\top u_l^* + \delta a^\top a_\perp) \right| \\
&\geq a^\top u_l^* + \frac{1}{\sqrt{1 + \delta^2 \|a_\perp\|_2^2}} \delta \|a_\perp\|_2^2 \\
&\geq \frac{1}{2} \delta \|a_\perp\|_2^2.
\end{aligned}$$

Putting all this together, we conclude that

$$\min |a^\top (u_l^* \pm \tilde{u}_a)| \geq \frac{1}{4} \delta \|a_\perp\|_2^2,$$

which taken collectively with (188) and our choice $\delta = \frac{\sigma_{\min}}{12|\lambda_l^*| \cdot \|a_\perp\|_2}$ yields

$$\begin{aligned}
& \inf_{\hat{u}_{a,l}} \sup_{A \in \mathcal{M}_1(M^*)} \mathbb{E}[\min |\hat{u}_{a,l} \pm a^\top u_l(A)|] \\
&\gtrsim \min |a^\top (u_l^* \pm \tilde{u}_a)| \gtrsim \frac{\sigma_{\min}}{|\lambda_l^*|} \|a_\perp\|_2.
\end{aligned}$$

C. Proof of Eq. (21d)

The proof of this part uses Fano's inequality. We start by constructing a proper packing set of the space: N unit vectors $\{v_i\}_{1 \leq i \leq N}$ within the subspace perpendicular to $\{u_i^*\}_{i \neq l}$ obeying

- $\langle v_i, u_j^* \rangle = 0$ for all $1 \leq i \leq N$, $1 \leq j \leq r$ and $j \neq l$;
- there exists some sufficiently small constant $c_3 > 0$ such that

$$\|v_i - v_j\|_2 \geq c_3 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|}, \quad 1 \leq i \neq j \leq N;$$

- there exists some sufficiently large constant $c_4 > 0$ such that

$$\|v_i - u_l^*\|_2 \leq c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|}, \quad 1 \leq i \leq N.$$

Standard packing arguments [1, Chapter 5.1] imply that N can be as large as

$$N = \exp\left(n \log \frac{c_4}{2c_3}\right). \tag{190}$$

In addition, when $4c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|} < 1$, it follows that

$$\|v_i - v_j\|_2 \leq \|v_i - u_l^*\|_2 + \|v_j - u_l^*\|_2 \leq 2c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|}, \tag{191}$$

$$\begin{aligned}
\|v_i + v_j\|_2 &\geq 2\|v_i\|_2 - \|v_i - v_j\|_2 \geq 2 - 2c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|} \\
&> 2c_4 \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|},
\end{aligned} \tag{192}$$

thus indicating that

$$\min \|v_i \pm v_j\|_2 \geq \min \{c_3, 2c_4\} \frac{\sigma_{\min} \sqrt{n}}{|\lambda_l^*|}, \quad 1 \leq i \neq j \leq N. \tag{193}$$

The next step is to associate each vector $\mathbf{v}_i$ ($1 \leq i \leq N$) with a hypothesis as follows

$$\mathcal{H}_i: \mathbf{M} = \mathbf{M}_i + \mathbf{H} := \lambda_i^* \mathbf{v}_i \mathbf{v}_i^\top + \sum_{j:j \neq i} \lambda_j^* \mathbf{u}_j^* \mathbf{u}_j^{*\top} + \mathbf{H}.$$

If we denote by $\mathbb{P}_i$ the distribution of $\mathbf{M}$ under $\mathcal{H}_i$ ($1 \leq i \leq N$), then we can invoke the argument in (177) and (179) to upper bound

$$\begin{aligned} \text{KL}(\mathbb{P}_i \| \mathbb{P}_j) &\leq \frac{\|\mathbf{M}_i - \mathbf{M}_j\|_F^2}{2\sigma_{\min}^2} = \frac{(\lambda_i^*)^2}{2\sigma_{\min}^2} \|\mathbf{v}_i \mathbf{v}_i^\top - \mathbf{v}_j \mathbf{v}_j^\top\|_F^2 \\ &\leq \frac{(\lambda_i^*)^2}{\sigma_{\min}^2} \|\mathbf{v}_i - \mathbf{v}_j\|_2^2 \\ &\leq 4c_4^2 n \end{aligned} \quad (193)$$

for any $i \neq j$, where the last inequality follows from (190). This upper bound on KL divergence plays an important role in invoking Fano's inequality. More specifically, recall that Fano's inequality [85, Corollary 2.6] asserts that if

$$\frac{1}{N} \sum_{i=2}^N \text{KL}(\mathbb{P}_i \| \mathbb{P}_1) \leq \frac{1}{4} \log N,$$

then the minimax probability of testing error must satisfy

$$p_{e,N} := \inf_{\psi} \max_{1 \leq i \leq N} \mathbb{P}\{\psi \neq i \mid \mathcal{H}_i\} \geq 1/4,$$

where the infimum is over all tests. Combining this with the upper bound (193) and the packing number (189), we see that if

$$4c_4^2 n \leq \frac{N}{4(N-1)} \log N = \frac{N}{4(N-1)} \cdot n \log \frac{c_4}{2c_3}, \quad (194)$$

then one would have $p_{e,N} \geq 1/4$. Clearly, the condition (194) would hold as long as c_4/c_3 is sufficiently large.

To finish up, it suffices to invoke the standard reduction scheme [85, Chapter 2.2] to obtain

$$\begin{aligned} \inf_{\hat{\mathbf{u}}_l} \sup_{\mathbf{A} \in \mathcal{M}_2} \mathbb{E} \left[\min \|\hat{\mathbf{u}}_l \pm \mathbf{u}_l(\mathbf{A})\|_2 \right] \\ \gtrsim \min_{1 \leq i \neq j \leq N} \{\min \|\mathbf{v}_i \pm \mathbf{v}_j\|_2\} \gtrsim \frac{\sigma_{\min} \sqrt{n}}{|\lambda_i^*|}, \end{aligned}$$

where the last inequality is a consequence of the condition (192). This concludes the proof.

APPENDIX F

A FEW MORE AUXILIARY LEMMAS

Lemma 12: Suppose that $\mathbf{H}$ satisfies Assumptions 1. Fix any $\mathbf{a} \in \mathbb{R}^n$. Then one has

$$\begin{aligned} \text{Var} \left[\frac{1}{2\lambda_i^*} \mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_i^* \right] &= \frac{1}{4(\lambda_i^*)^2} \sum_{1 \leq i, j \leq n} (a_i u_{i,j}^* + a_j u_{i,i}^*)^2 \sigma_{ij}^2 \\ &=: v_{\mathbf{a},i}^*, \end{aligned} \quad (195)$$

which satisfies

$$\frac{1}{2} \left(\|\mathbf{a}\|_2^2 + (\mathbf{a}^\top \mathbf{u}_i^*)^2 \right) \frac{\sigma_{\min}^2}{(\lambda_i^*)^2} \leq v_{\mathbf{a},i}^* \leq \frac{1}{2} \left(\|\mathbf{a}\|_2^2 + (\mathbf{a}^\top \mathbf{u}_i^*)^2 \right) \frac{\sigma_{\max}^2}{(\lambda_i^*)^2}. \quad (196)$$

Proof: See Appendix F-A. $\square$

If we further have $\mathbf{a}^\top \mathbf{u}_i^* = o(\|\mathbf{a}\|_2)$, then one has

$$\frac{1 + o(1)}{2} \frac{\sigma_{\min}^2 \|\mathbf{a}\|_2^2}{(\lambda_i^*)^2} \leq v_{\mathbf{a},i}^* \leq \frac{1 + o(1)}{2} \frac{\sigma_{\max}^2 \|\mathbf{a}\|_2^2}{(\lambda_i^*)^2}. \quad (197)$$

Lemma 13: Suppose that $\mathbf{H}$ satisfies Assumptions 1 and obeys $B \leq \sigma_{\max} \sqrt{\frac{n}{\mu \log n}}$. Fix any $\mathbf{a} \in \mathbb{R}^n$. Then with probability at least $1 - O(n^{-10})$, we have

$$\begin{aligned} \max \{ |\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*|, |\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*|, |\mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*| \} \\ \lesssim \sigma_{\max} \|\mathbf{a}\|_2 \sqrt{\log n} \end{aligned} \quad (198)$$

for all $1 \leq l \leq r$.

Proof: See Appendix F-B. $\square$

Lemma 14: Fix any unit vector $\mathbf{a} \in \mathbb{R}^n$, and consider another fixed unit vector $\mathbf{u}^*$ obeying $\|\mathbf{u}^*\|_\infty \leq \sqrt{\mu/n}$. Suppose that $\mathbf{H}$ satisfies Assumption 1 and obeys $B \leq \sigma_{\min} \sqrt{n/(\mu \log n)}$. Then the distribution of $W := \frac{\mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}^*}{2\sqrt{v^*}}$ satisfies

$$\sup_{z \in \mathbb{R}} |\mathbb{P}(W \leq z) - \Phi(z)| \leq \frac{8}{\sqrt{\log n}}, \quad (199)$$

where $\Phi(\cdot)$ is the CDF of a standard Gaussian, and $v^* := \text{Var}(\frac{1}{2} \mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}^*)$.

Proof: See Appendix F-B. $\square$

Lemma 15: Suppose that $\mathbf{H}$ satisfies Assumptions 1 and 2. Let $\mathbf{u}_l$ (resp. $\mathbf{w}_l$) be the l th right (resp. left) eigenvector of $\mathbf{M}$ obeying $\mathbf{u}_l^\top \mathbf{u}^* > 0$ and $\mathbf{u}_l^\top \mathbf{w}_l > 0$, and set $\hat{\mathbf{u}}_l := \frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} (\mathbf{u}_l + \mathbf{w}_l)$. Then with probability exceeding $1 - O(n^{-10})$, one has

$$\begin{cases} \mathbf{u}_l^{*\top} \mathbf{u}_l = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50, \\ \mathbf{u}_l^{*\top} \mathbf{w}_l = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50, \\ \hat{\mathbf{u}}_l^\top \mathbf{u}_l^* = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50, \end{cases} \quad (200)$$

and

$$\begin{aligned} \|\mathbf{u}_l + \mathbf{w}_l\|_2 &= 2 + O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \\ &\geq 9/5. \end{aligned} \quad (201)$$

In addition, if Assumption 3 holds, then with probability at least $1 - O(n^{-10})$,

$$\begin{aligned} \max \{ \|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty, \|\mathbf{w}_l - \mathbf{u}_l^*\|_\infty, \|\hat{\mathbf{u}}_l - \mathbf{u}_l^*\|_\infty \} \\ \lesssim \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^4 r \log n} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}} \\ \lesssim o\left(\frac{1}{\sqrt{\mu n \log n}}\right). \end{aligned} \quad (202)$$

Proof: See Appendix F-C. $\square$

A. Proof of Lemma 12

Without loss of generality, assume that $\lambda_l^* = 1$ and that $\|\mathbf{a}\|_2 = 1$. The first claim (195) follows from direct calculations. Regarding the second claim (196), we first establish the lower bound

$$\begin{aligned} v_{\mathbf{a},l}^* &\geq \frac{1}{4} \sum_{1 \leq i,j \leq n} (a_i u_{l,j}^* + a_j u_{l,i}^*)^2 \sigma_{\min}^2 \\ &= \frac{1}{4} \sum_{1 \leq i,j \leq n} (a_i^2 (u_{l,j}^*)^2 + a_j^2 (u_{l,i}^*)^2 + 2a_i u_{l,i}^* a_j u_{l,j}^*) \sigma_{\min}^2 \\ &= \frac{1}{2} (1 + (\mathbf{a}^\top \mathbf{u}_l^*)^2) \sigma_{\min}^2, \end{aligned}$$

where we have used the fact that $\sum_{i,j} a_i^2 (u_{l,j}^*)^2 = \sum_i a_i^2 \sum_j (u_{l,j}^*)^2 = 1$. A similar argument leads to the advertised upper bound $v_{\mathbf{a},l}^* \leq \frac{1}{2} (1 + (\mathbf{a}^\top \mathbf{u}_l^*)^2) \sigma_{\max}^2$.

B. Proofs of Lemma 13 and Lemma 14

Without loss of generality, we shall assume $\lambda_l^* = 1$ throughout the proof. We shall also assume that $|H_{ij}| \leq B$ for all $1 \leq i, j \leq n$.⁵ By direct calculations, the following quantity of interest can be expressed as

$$\frac{\mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}^*}{2} = \frac{1}{2} \sum_{i,j} (a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij} =: T_1. \quad (203)$$

The above quantity is the sum of n^2 independent zero-mean random variables, each obeying

$$\text{Var}[(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}] = (a_i u_{l,j}^* + a_j u_{l,i}^*)^2 \sigma_{ij}^2; \quad (204)$$

$$|(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}| \leq |a_i u_{l,j}^* + a_j u_{l,i}^*| B. \quad (205)$$

Proof of Lemma 13: We shall only establish the lemma for the quantity $\mathbf{a}^\top (\mathbf{H} + \mathbf{H}^\top) \mathbf{u}_l^*$; the bounds on $\mathbf{a}^\top \mathbf{H} \mathbf{u}_l^*$ and $\mathbf{a}^\top \mathbf{H}^\top \mathbf{u}_l^*$ follow similarly. Invoking Bernstein's inequality for bounded random variables [84], we guarantee that with probability at least $1 - O(n^{-11})$,

$$\begin{aligned} &\left| \sum_{i,j} (a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij} \right| \\ &\lesssim \sqrt{\sum_{i,j} \text{Var}[(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}] \log n} \\ &\quad + \max_{i,j} \{|(a_i u_{l,j}^* + a_j u_{l,i}^*) B|\} \log n \\ &\leq \sqrt{\sum_{i,j} (a_i u_{l,j}^* + a_j u_{l,i}^*)^2 \sigma_{\max}^2 \log n} \\ &\quad + \max_{i,j} \{|(a_i u_{l,j}^* + a_j u_{l,i}^*) B|\} \log n \\ &\lesssim \sqrt{(\sigma_{\max}^2 \log n) \sum_{i,j} \{(a_i u_{l,j}^*)^2 + (a_j u_{l,i}^*)^2\}} \\ &\quad + B \|\mathbf{u}_l^*\|_\infty \|\mathbf{a}\|_2 \log n \\ &\asymp \sqrt{(\sigma_{\max}^2 \log n) \|\mathbf{a}\|_2^2 \|\mathbf{u}_l^*\|_2^2} + B \|\mathbf{u}_l^*\|_\infty \|\mathbf{a}\|_2 \log n. \end{aligned}$$

⁵Otherwise, we can first look at the truncated version $\tilde{H}_{ij} := H_{ij} \mathbb{1}_{\{|H_{ij}| \leq B\}}$ (which is also zero-mean with variance obeying $\text{Var}[\tilde{H}_{ij}] = (1 + o(1)) \sigma_{ij}^2$ under our assumptions), and then argue that $H_{ij} = \tilde{H}_{ij}$ ($\forall i, j$) with high probability. This argument is fairly standard and is hence omitted for brevity.

Now, as an immediate consequence of the incoherence condition and the identity $\|\mathbf{a}\|_2 = \|\mathbf{u}_l^*\|_2 = 1$, we have

$$\left| \sum_{i,j} (a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij} \right| \lesssim \sigma_{\max} \sqrt{\log n} + B \log n \sqrt{\frac{\mu}{n}}.$$

Combining this with the expression (203) and the assumption $B \leq \sigma_{\max} \sqrt{\frac{n}{\mu \log n}}$, we complete the proof of the inequality (198) via the union bound.

Proof of Lemma 14: To establish this inequality, the key is to make use of the following Berry-Esseen type inequality [86, Theorem 3.7].

Theorem 11: Let $\xi_1, \xi_2, \dots, \xi_n$ be independent zero-mean random variables satisfying $\sum_{i=1}^n \text{Var}[\xi_i] = 1$. Then

$$\sup_{z \in \mathbb{R}} \left| \mathbb{P}\left(\sum_{i=1}^n \xi_i \leq z\right) - \Phi(z) \right| \leq 10\gamma, \quad \text{where } \gamma := \sum_{i=1}^n \mathbb{E}[|\xi_i|^3]. \quad (206)$$

In order to apply Theorem 11, let us define

$$\xi_{ij} := \frac{1}{\sqrt{\text{Var}[T_1]}} \frac{1}{2} (a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij} \quad \text{and} \quad W = \sum_{i,j} \xi_{ij},$$

where it follows from the equality (204) that

$$v^* := \text{Var}[T_1] = \frac{1}{4} \sum_{i,j} (a_i u_{l,j}^* + a_j u_{l,i}^*)^2 \sigma_{ij}^2. \quad (207)$$

By definition, it is easily seen that $\sum_{i,j} \text{Var}[\xi_{ij}] = 1$. Therefore, the property (206) follows. In order to establish inequality (199), it suffices to upper bound γ in Theorem 11.

To this end, we first make the observation that

$$\begin{aligned} \gamma &= \sum_{i,j} \mathbb{E}[|\xi_{ij}|^3] \\ &= \frac{1}{\text{Var}[T_1]^{3/2}} \frac{1}{8} \sum_{i,j} \mathbb{E}[|(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}|^3] \\ &\leq \frac{1}{8(v^*)^{3/2}} \\ &\quad \cdot \sum_{i,j} \mathbb{E}\left[\max_{i,j} |(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}| \cdot |(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}|^2\right] \\ &\leq \frac{1}{8(v^*)^{3/2}} \sum_{i,j} \mathbb{E}[\|\mathbf{a}\|_2 \|\mathbf{u}_l^*\|_\infty B \cdot |(a_i u_{l,j}^* + a_j u_{l,i}^*) H_{ij}|^2] \\ &\leq \frac{1}{8(v^*)^{3/2}} \|\mathbf{a}\|_2 \|\mathbf{u}_l^*\|_\infty B (4v^*) = \frac{1}{2(v^*)^{1/2}} \|\mathbf{a}\|_2 \|\mathbf{u}_l^*\|_\infty B. \end{aligned}$$

Re-organizing terms and using $\|\mathbf{a}\|_2 = \|\mathbf{u}_l^*\|_2 = 1$ as well as the incoherence condition, we have

$$\gamma \leq B \sqrt{\frac{\mu}{4v^* n}}. \quad (208)$$

It thus remains to lower bound v^* . Again, use Lemma 12 and the fact $\|\mathbf{a}\|_2 = 1$ to arrive at $4v^* \geq 2\sigma_{\min}^2$. Substitution into (208) allows us to further control γ as

$$\gamma \leq \frac{B}{\sigma_{\min}} \sqrt{\frac{\mu}{2n}} \leq \frac{1}{\sqrt{2 \log n}},$$

where the last inequality arises since $B \leq \sigma_{\min} \sqrt{n/(\mu \log n)}$. Applying Theorem 11 concludes the proof.

C. Proof of Lemma 15

To begin with, Theorem 7 tells that that $\mathbf{u}_l$ and $\mathbf{w}_l$ are both real-valued under Assumption 2. It has also been shown in Theorem 9 that

$$\begin{cases} \mathbf{u}_l^* \top \mathbf{u}_l = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50 \\ |\mathbf{u}_l^* \top \mathbf{w}_l| = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50 \end{cases}$$

provided that $\mathbf{u}_l^* \top \mathbf{u}_l > 0$ and that Assumption 2 holds. This immediately indicates that (i) $\|\mathbf{u}_l - \mathbf{u}_l^*\|_2 = \sqrt{2 - 2\langle \mathbf{u}_l, \mathbf{u}_l^* \rangle} \leq 1/5$ and $\|\mathbf{u}_l + \mathbf{u}_l^*\|_2 = 2\|\mathbf{u}_l^*\|_2 - \|\mathbf{u}_l - \mathbf{u}_l^*\|_2 \geq 9/5$, and (ii) one either has $\mathbf{w}_l^* \top \mathbf{u}_l^* \geq 49/50$ or $\mathbf{w}_l^* \top \mathbf{u}_l^* \leq -49/50$. If $\mathbf{w}_l^* \top \mathbf{u}_l^* \leq 49/50$, then one necessarily has

$$\begin{aligned} \mathbf{w}_l^* \top \mathbf{u}_l &= \mathbf{w}_l^* \top \mathbf{u}_l^* + \mathbf{w}_l^* \top (\mathbf{u}_l - \mathbf{u}_l^*) \leq -49/50 + \|\mathbf{w}_l\|_2 \|\mathbf{u}_l - \mathbf{u}_l^*\|_2 \\ &\leq -49/50 + 1/5 < 0, \end{aligned}$$

thus contradicting the assumption that $\mathbf{w}_l^* \top \mathbf{u}_l > 0$. As a result, one concludes that

$$\mathbf{u}_l^* \top \mathbf{w}_l = 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \geq 49/50. \quad (209)$$

Further, it follows from Theorem 9 that

$$\begin{aligned} &\min\{\|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty, \|\mathbf{u}_l + \mathbf{u}_l^*\|_\infty\} \\ &\lesssim \frac{\sigma_{\max}}{\lambda_{\min}^*} \sqrt{\mu \kappa^4 r \log n} + \frac{\sigma_{\max}}{\Delta_l^*} \sqrt{\frac{\mu^2 \kappa^4 r^3 \log n}{n}} \\ &\leq o\left(\frac{1}{\sqrt{\mu n \log n}}\right), \end{aligned}$$

where the last line holds under Assumption 3. Suppose that $\|\mathbf{u}_l + \mathbf{u}_l^*\|_\infty \leq \|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty$. Then the above inequality implies that

$$\begin{aligned} \|\mathbf{u}_l + \mathbf{u}_l^*\|_2 &\leq \sqrt{n} \|\mathbf{u}_l + \mathbf{u}_l^*\|_\infty \\ &= \sqrt{n} \min\{\|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty, \|\mathbf{u}_l + \mathbf{u}_l^*\|_\infty\} \ll 1, \end{aligned}$$

which is contradictory to the relation $\|\mathbf{u}_l + \mathbf{u}_l^*\|_2 \geq 9/5$ shown above. Therefore, we must have

$$\begin{aligned} \|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty &= \min\{\|\mathbf{u}_l - \mathbf{u}_l^*\|_\infty, \|\mathbf{u}_l + \mathbf{u}_l^*\|_\infty\} \\ &\lesssim O\left(\frac{\sigma_{\max} \sqrt{\mu \log n}}{|\lambda^*|}\right) = o\left(\frac{1}{\sqrt{\mu n \log n}}\right) \end{aligned} \quad (210)$$

as claimed. Similarly, we shall also have $\|\mathbf{u}_l - \mathbf{u}_l^*\|_2 = \min\{\|\mathbf{u}_l \pm \mathbf{u}_l^*\|_2\}$ under Assumption 2, which combined with Theorem 9 gives

$$\begin{aligned} &\max\{\|\mathbf{u}_l - \mathbf{u}_l^*\|_2, \|\mathbf{w}_l - \mathbf{w}_l^*\|_2\} \\ &\lesssim O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right). \end{aligned} \quad (211)$$

In addition, the above results further imply that

$$\begin{aligned} &\|\mathbf{u}_l + \mathbf{w}_l\|_2 - 2 \\ &= \frac{\|\mathbf{u}_l + \mathbf{w}_l\|_2^2 - 4}{\|\mathbf{u}_l + \mathbf{w}_l\|_2 + 2} \stackrel{(i)}{\asymp} \|\mathbf{u}_l + \mathbf{w}_l\|_2^2 - 4 \\ &= \|\mathbf{u}_l\|_2^2 + \|\mathbf{w}_l\|_2^2 - 2 + 2\{\langle \mathbf{u}_l, \mathbf{w}_l \rangle - 1\} \\ &\stackrel{(ii)}{=} 2\{\langle \mathbf{u}_l, \mathbf{w}_l \rangle - \langle \mathbf{u}_l^*, \mathbf{u}_l^* \rangle\} \\ &= 2\{\langle \mathbf{u}_l - \mathbf{u}_l^*, \mathbf{w}_l \rangle + \langle \mathbf{u}_l^*, \mathbf{w}_l - \mathbf{u}_l^* \rangle\} \\ &\stackrel{(iii)}{=} O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right), \end{aligned} \quad (212)$$

where (i) holds since $2 \leq \|\mathbf{u}_l + \mathbf{w}_l\|_2 + 2 \leq 4$, (ii) follows since $\|\mathbf{u}_l\|_2 = \|\mathbf{w}_l\|_2 = 1$, and (iii) comes from Cauchy-Schwarz and the ℓ_2 bound (211).

To finish up, it remains to show that the claimed bounds hold when $\mathbf{u}_l$ is replaced by $\hat{\mathbf{u}}_l$. Regarding the bound on $\hat{\mathbf{u}}_l^* \top \mathbf{u}_l^*$, we make the observation that

$$\begin{aligned} &\hat{\mathbf{u}}_l^* \top \mathbf{u}_l^* \\ &= \frac{1}{2}(\mathbf{u}_l^* \top \mathbf{u}_l^* + \mathbf{w}_l^* \top \mathbf{u}_l^*) + \hat{\mathbf{u}}_l^* \top \mathbf{u}_l^* - \frac{1}{2}(\mathbf{u}_l + \mathbf{w}_l)^* \top \mathbf{u}_l^* \\ &= \frac{1}{2}(\mathbf{u}_l^* \top \mathbf{u}_l^* + \mathbf{w}_l^* \top \mathbf{u}_l^*) + \left(\frac{1}{\|\mathbf{u}_l + \mathbf{w}_l\|_2} - \frac{1}{2}\right)(\mathbf{u}_l + \mathbf{w}_l)^* \top \mathbf{u}_l^* \\ &= 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right) \\ &\quad + \frac{2 - \|\mathbf{u}_l + \mathbf{w}_l\|_2}{2\|\mathbf{u}_l + \mathbf{w}_l\|_2} \cdot (\mathbf{u}_l + \mathbf{w}_l)^* \top \mathbf{u}_l^* \\ &= 1 - O\left(\frac{\kappa^4 \sigma_{\max}^2 n \log n}{(\lambda_{\max}^*)^2} + \frac{\mu \kappa^4 r^2 \sigma_{\max}^2 \log n}{(\Delta_l^*)^2}\right), \end{aligned}$$

where the last line comes from (212). The advertised ℓ_∞ norm bounds for $\hat{\mathbf{u}}_l$ follow from the same arguments; we omit it for brevity.

APPENDIX G

EXAMPLE: A SYMMETRIC CASE WITH HOMOSCEDASTIC GAUSSIAN NOISE

In this section, we isolate a simple example to illustrate the potential applicability of our main results in the presence of certain symmetric noise matrices. Specifically, consider the symmetric and homoscedastic case with Gaussian noise, namely,

$$H_{ij} = H_{ji} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \quad 1 \leq i \leq j \leq n; \quad (213a)$$

$$\mathbf{M}^* = \sum_{l=1}^r \lambda_l^* \mathbf{u}_l^* \mathbf{u}_l^{*\top}, \quad \mathbf{M} = \mathbf{M}^* + \mathbf{H}. \quad (213b)$$

Clearly, this setting differs from the model in Assumption 1 in that both $\mathbf{H}$ and $\mathbf{M}$ are now symmetric matrices. To invoke our theorems, we need to first asymmetrize the data matrix. Towards this, we propose a procedure as follows, motivated by a simple asymmetrization trick pointed out by [17].

- 1) Generate $\tilde{\mathbf{M}} = \mathbf{M} + \tilde{\mathbf{H}}$, where $\tilde{\mathbf{H}}$ is a skew-symmetric matrix obeying

$$\tilde{H}_{ij} = -\tilde{H}_{ji} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \quad 1 \leq i < j \leq n. \quad (214)$$

2) Run Algorithm 1 with input data matrix $\widetilde{M}$.

In short, the above scheme injects additional noise to the data matrix before implementing the inferential procedures. The theoretical support is this:

Corollary 3: Consider the model (213). Then Theorems 4, 5 and 6 continue to hold, if Algorithms 1 and 2 take $\widetilde{M}$ as the input matrix.

Proof: The key observation is that: under our construction, the quantities $H_{ij} + \widetilde{H}_{ij}$ and $H_{ij} - \widetilde{H}_{ij}$ are uncorrelated and hence independent. Therefore, the effective noise matrix $\mathbf{H}_{\text{eff}} := \mathbf{H} + \widetilde{\mathbf{H}}$ is in general asymmetric and contains i.i.d. entries drawn from $\mathcal{N}(0, 2\sigma^2)$. Given that

$$\mathbf{H}_{\text{eff}} + \mathbf{H}_{\text{eff}}^\top = \mathbf{H} + \mathbf{H}^\top + \widetilde{\mathbf{H}} + \widetilde{\mathbf{H}}^\top = \mathbf{H} + \mathbf{H}^\top, \quad (215)$$

applying Theorems 4, 5 and 6 immediately concludes the proof. $\square$

Interestingly, the above findings indicate that: a simple asymmetrization trick via proper noise injection allows one to construct confidence intervals under symmetric and homoscedastic Gaussian noise. We caution that the above procedure requires prior knowledge on the noise level σ , which can often be reliably estimated in the homoscedastic case.

While noise injection does not affect the first-order uncertainty term $\mathbf{a}^\top(\mathbf{H} + \mathbf{H}^\top)\mathbf{u}_l^*$, it does inflate the higher-order residual terms. For practical purposes, we recommend running the above procedure independently for multiple times and returning the “average” of them. This is summarized as follows, which might help improve practical performance.

- 1) Generate K matrices $\mathbf{M}^{(k)} = \mathbf{M} + \mathbf{H}^{(k)}$ ($1 \leq k \leq K$), where $\{\mathbf{H}^{(k)}\}$ are independent skew-symmetric matrices obeying

$$H_{ij}^{(k)} = -H_{ji}^{(k)} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, \sigma^2), \quad 1 \leq i < j \leq n. \quad (216)$$

- 2) For each $1 \leq k \leq K$, run Algorithm 1 with input data matrix $\mathbf{M}^{(k)}$. We denote by $\lambda_l^{(k)}$ the l th eigenvalue of $\mathbf{M}^{(k)}$, $\mathbf{u}_l^{(k)}$ the associated right eigenvector, $\widehat{\mathbf{u}}_{a,l}^{(k)}$ the resulting estimator for $\mathbf{a}^\top \mathbf{u}_l^*$, respectively. Here, we calibrate the global signs of $\{\mathbf{u}_l^{(k)}\}$ by ensuring $\langle \mathbf{u}_l^{(k)}, \mathbf{u}_l^{(k+1)} \rangle \geq 0$ for all $1 \leq k < K$. We also let $\widehat{v}_{a,l}^{(k)}$ and $\widehat{v}_{\lambda,l}^{(k)}$ represent the resulting variance estimators.
- 3) Average the estimators

$$\widehat{\mathbf{u}}_{a,l} = \frac{1}{K} \sum_{k=1}^K \widehat{\mathbf{u}}_{a,l}^{(k)} \quad \text{and} \quad \lambda_l = \frac{1}{K} \sum_{k=1}^K \lambda_l^{(k)}. \quad (217)$$

- 4) For any prescribed coverage level $1 - \alpha$, return the confidence intervals for $\mathbf{a}^\top \mathbf{u}_l^*$ and λ_l^* as follows

$$\text{CI}_{1-\alpha}^a := \left[\widehat{\mathbf{u}}_{a,l} \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\widehat{v}_{a,l}^{(1)}} \right]; \quad (218a)$$

$$\text{CI}_{1-\alpha}^\lambda := \left[\lambda_l \pm \Phi^{-1}(1 - \alpha/2) \sqrt{\widehat{v}_{\lambda,l}^{(1)}} \right]. \quad (218b)$$

APPENDIX H

SYMMETRIZE OR NOT? SOME HIGH-LEVEL INTERPRETATION

The numerical experiments reported in Section V reveal the potential benefit of Spectral-asym compared to Spectral-sym. Here, we provide some informal interpretation.

If we hope the eigenvector $\mathbf{u}_2$ to behave as a reliable estimate of $\mathbf{u}_2^*$ (meaning that $|\mathbf{u}_2^\top \mathbf{u}_2^*| = 1 - o(1)$), then we would necessarily require $|\mathbf{u}_2^\top \mathbf{u}_1^*| = o(1)$. To develop an understanding about the size of $\mathbf{u}_2^\top \mathbf{u}_1^*$, we recall from Neumann’s series that (see [17] or Appendix A)

$$\begin{aligned} \mathbf{u}_1^{*\top} \mathbf{u}_2 &= \frac{\lambda_1^* (\mathbf{u}_1^{*\top} \mathbf{u}_2)}{\lambda_2} + \frac{\lambda_2^* (\mathbf{u}_2^{*\top} \mathbf{u}_2)}{\lambda_2} \sum_{s=1}^{\infty} \frac{\mathbf{u}_1^{*\top} \mathbf{H}^s \mathbf{u}_2^*}{\lambda_2^s} \\ &\quad + \frac{\lambda_1^* (\mathbf{u}_1^{*\top} \mathbf{u}_2)}{\lambda_2} \sum_{s=1}^{\infty} \frac{\mathbf{u}_1^{*\top} \mathbf{H}^s \mathbf{u}_1^*}{\lambda_2^s}, \end{aligned} \quad (219)$$

or equivalently,

$$\mathbf{u}_1^{*\top} \mathbf{u}_2 = \frac{\lambda_2^* (\mathbf{u}_2^{*\top} \mathbf{u}_2)}{\lambda_2 - \lambda_1^* - \lambda_1^* \sum_{s=1}^{\infty} \frac{\mathbf{u}_1^{*\top} \mathbf{H}^s \mathbf{u}_1^*}{\lambda_2^s}} \sum_{s=1}^{\infty} \frac{\mathbf{u}_1^{*\top} \mathbf{H}^s \mathbf{u}_2^*}{\lambda_2^s}. \quad (220)$$

Consequently, an important condition that allows one to control quantity $\mathbf{u}_1^{*\top} \mathbf{u}_2$ is to ensure that each summand in expression (220) is as small as possible.

Towards this end, let us single out the second-order term $\mathbf{u}_1^{*\top} \mathbf{H}^2 \mathbf{u}_2^*$, which suffices for us to develop some intuition (here, we assume that $\lambda_2 \approx \lambda_2^*$ so that the denominator λ_2^2 does not affect the order of this term).

- If $\mathbf{H}$ is asymmetric, then it has been shown in [17] (see also Appendix A) that

$$|\mathbf{u}_1^{*\top} \mathbf{H}^2 \mathbf{u}_2^*| \lesssim \sigma_{\max}^2 \sqrt{\mu n} \log n \quad (221)$$

with high probability, which is exceedingly small given the assumption that $\sigma_{\max} \sqrt{n \log n} \ll 1$.

- If $\mathbf{H}$ is replaced by $\widetilde{\mathbf{H}} := \frac{1}{2}(\mathbf{H} + \mathbf{H}^\top)$, then in general there is no guarantee that $\mathbf{u}_1^{*\top} \widetilde{\mathbf{H}}^2 \mathbf{u}_2^*$ can be equally well-controlled. Take the case (50) and (51) for example: straightforward calculations reveal

$$\mathbb{E}[\mathbf{u}_1^{*\top} \widetilde{\mathbf{H}}^2 \mathbf{u}_2^*] = \mathbf{u}_1^{*\top} \text{Var}\left(\frac{1}{2}(\mathbf{H} + \mathbf{H}^\top)\right) \mathbf{u}_2^* = \frac{1}{4} \sigma_1^2 n \asymp \sigma_{\max}^2 n, \quad (222)$$

which far exceeds the order (221) with asymmetric data matrices. If this is the case, then one would have to require a better control of the multiplicative term $\frac{\lambda_2^*}{\lambda_2 - \lambda_1^* - \lambda_1^* \sum_{s=1}^{\infty} \frac{\mathbf{u}_1^{*\top} \mathbf{H}^s \mathbf{u}_1^*}{\lambda_2^s}}$ in (220), which is often equivalent to imposing a more stringent separation condition on the eigenvalue pair $(\lambda_1^*, \lambda_2^*)$.

The take-home message is this: Spectral-sym might sometimes be suboptimal when dealing with heteroscedastic noise, particularly when the variance structure of the noise matrix is somewhat “aligned” with the true eigen-structure. As a result, Spectral-sym might not be as robust as Spectral-asym when performing eigenvector estimation.

ACKNOWLEDGMENT

Chen Cheng would like to thank Lihua Lei for helpful discussions.

REFERENCES

- [1] M. J. Wainwright, *High-Dimensional Statistics: A Non-Asymptotic Viewpoint*, vol. 48. Cambridge, U.K.: Cambridge Univ. Press, 2019.
- [2] E. J. Candès and Y. Plan, "Matrix completion with noise," *Proc. IEEE*, vol. 98, no. 6, pp. 925–936, Jun. 2010.
- [3] Y. Hua and T. K. Sarkar, "Matrix pencil method for estimating parameters of exponentially damped/undamped sinusoids in noise," *IEEE Trans. Acoust., Speech, Signal Process.*, vol. 38, no. 5, pp. 814–824, May 1990.
- [4] Y. Chen and Y. Chi, "Robust spectral compressed sensing via structured matrix completion," *IEEE Trans. Inf. Theory*, vol. 60, no. 10, pp. 6576–6601, Oct. 2014.
- [5] A. Javanmard and A. Montanari, "Localization from incomplete noisy distance measurements," *Found. Comput. Math.*, vol. 13, no. 3, pp. 297–345, 2013.
- [6] Y. Chen, L. Guibas, and Q. Huang, "Near-optimal joint object matching via convex relaxation," in *Proc. Int. Conf. Mach. Learn.*, 2014, pp. 100–108.
- [7] S. Péché, "The largest eigenvalue of small rank perturbations of Hermitian random matrices," *Probab. Theory Rel. Fields*, vol. 134, no. 1, pp. 127–173, 2006.
- [8] T. T. Cai and A. Zhang, "Rate-optimal perturbation bounds for singular subspaces with applications to high-dimensional statistics," *Ann. Statist.*, vol. 46, no. 1, pp. 60–89, 2018.
- [9] E. Abbe, J. Fan, K. Wang, and Y. Zhong, "Entrywise eigenvector analysis of random matrices with low expected rank," *Ann. Statist.*, vol. 48, no. 3, p. 1452, Jun. 2020.
- [10] S. O'Rourke, V. Vu, and K. Wang, "Random perturbation of low rank matrices: Improving classical bounds," *Linear Algebra Appl.*, vol. 540, pp. 26–59, Mar. 2018.
- [11] V. Vu, "Singular vectors under random perturbation," *Random Struct. Algorithms*, vol. 39, no. 4, pp. 526–538, Dec. 2011.
- [12] Y. Yu, T. Wang, and R. J. Samworth, "A useful variant of the Davis–Kahan theorem for statisticians," *Biometrika*, vol. 102, no. 2, pp. 315–323, 2015.
- [13] R. H. Keshavan, A. Montanari, and S. Oh, "Matrix completion from noisy entries," *J. Mach. Learn. Res.*, vol. 11, pp. 2057–2078, Mar. 2010.
- [14] Y. Chi, Y. M. Lu, and Y. Chen, "Nonconvex optimization meets low-rank matrix factorization: An overview," *IEEE Trans. Signal Process.*, vol. 67, no. 20, pp. 5239–5269, Oct. 2019.
- [15] R. Wang, "Singular vector perturbation under Gaussian noise," *SIAM J. Matrix Anal. Appl.*, vol. 36, no. 1, pp. 158–177, Jan. 2015.
- [16] Y. Zhong, "Eigenvector under random perturbation: A nonasymptotic Rayleigh–Schrödinger theory," 2017, *arXiv:1702.00139*. [Online]. Available: <http://arxiv.org/abs/1702.00139>
- [17] Y. Chen, C. Cheng, and J. Fan, "Asymmetry helps: Eigenvalue and eigenvector analyses of asymmetrically perturbed low-rank matrices," *Ann. Statist.*, vol. 49, no. 1, pp. 435–458, Feb. 2021.
- [18] V. Koltchinskii, M. Löffler, and R. Nickl, "Efficient estimation of linear functionals of principal components," *Ann. Statist.*, vol. 48, no. 1, pp. 464–490, 2020.
- [19] J. Cho, D. Kim, and K. Rohe, "Asymptotic theory for estimating the singular vectors and values of a partially-observed low rank matrix with noise," *Statistica Sinica*, vol. 27, pp. 1921–1948, Oct. 2017.
- [20] C. Cai, G. Li, Y. Chi, H. V. Poor, and Y. Chen, "Subspace estimation from unbalanced and incomplete data matrices: ℓ_2, ∞ statistical guarantees," *Ann. Statist.*, vol. 49, no. 2, pp. 944–967, Apr. 2021.
- [21] A. Zhang, T. T. Cai, and Y. Wu, "Heteroskedastic PCA: Algorithm, optimality, and applications," *Ann. Statist.*, 2021.
- [22] V. Koltchinskii and K. Lounici, "Asymptotics and concentration bounds for bilinear forms of spectral projectors of sample covariance," *Annales De l'Institut Henri Poincaré, Probabilités Et Statistiques*, vol. 52, no. 4, pp. 1976–2013, 2016.
- [23] Y. Chen, Y. Chi, J. Fan, and C. Ma, "Spectral methods for data science: A statistical perspective," 2020, *arXiv:2012.08496*. [Online]. Available: <http://arxiv.org/abs/2012.08496>
- [24] C. Davis and W. M. Kahan, "The rotation of eigenvectors by a perturbation. III," *SIAM J. Numer. Anal.*, vol. 7, no. 1, pp. 1–46, 1970.
- [25] P.-Å. Wedin, "Perturbation bounds in connection with singular value decomposition," *BIT Numer. Math.*, vol. 12, no. 1, pp. 99–111, 1972.
- [26] E. Abbe, "Community detection and stochastic block models: Recent developments," *J. Mach. Learn. Res.*, vol. 18, no. 1, pp. 6446–6531, Jan. 2017.
- [27] Y. Chi and Y. Chen, "Compressive two-dimensional harmonic retrieval via atomic norm minimization," *IEEE Trans. Signal Process.*, vol. 63, no. 4, pp. 1030–1042, Feb. 2014.
- [28] G. Li, C. Cai, Y. Gu, H. V. Poor, and Y. Chen, "Minimax estimation of linear functions of eigenvectors in the face of small eigen-gaps," 2021, *arXiv:2104.03298*. [Online]. Available: <http://arxiv.org/abs/2104.03298>
- [29] T. T. Cai and M. G. Low, "An adaptation theory for nonparametric confidence intervals," *Ann. Statist.*, vol. 32, no. 5, pp. 1805–1840, 2004.
- [30] T. T. Cai and M. Low, "A framework for estimation of convex functions," *Statistica Sinica*, vol. 25, pp. 423–456, Apr. 2015.
- [31] T. T. Cai, A. Guntuboyina, and Y. Wei, "Adaptive estimation of planar convex sets," *Ann. Statist.*, vol. 46, no. 3, pp. 1018–1049, Jun. 2018.
- [32] C. Stein, "Efficient nonparametric testing and estimation," in *Proc. 3rd Berkeley Symp. Math. Statist. Probab., Contrib. Statist.*, vol. 1. Oakland, CA, USA: Regents of the Univ. of California, 1956, pp. 187–196.
- [33] C. Ma, K. Wang, Y. Chi, and Y. Chen, "Implicit regularization in nonconvex statistical estimation: Gradient descent converges linearly for phase retrieval, matrix completion, and blind deconvolution," *Found. Comput. Math.*, vol. 20, no. 3, pp. 451–632, Jun. 2020.
- [34] S. Negahban and M. J. Wainwright, "Restricted strong convexity and weighted matrix completion: Optimal bounds with noise," *J. Mach. Learn. Res.*, vol. 13, no. 1, pp. 1665–1697, 2012.
- [35] V. Koltchinskii, K. Lounici, and A. Tsybakov, "Nuclear-norm penalization and optimal rates for noisy low-rank matrix completion," *Ann. Statist.*, vol. 39, no. 5, pp. 2302–2329, 2011, doi: [10.1214/11-AOS894](https://doi.org/10.1214/11-AOS894).
- [36] A. Montanari, F. Ruan, and J. Yan, "Adapting to unknown noise distribution in matrix denoising," 2018, *arXiv:1810.02954*. [Online]. Available: <http://arxiv.org/abs/1810.02954>
- [37] D. Donoho and M. Gavish, "Minimax risk of matrix denoising by singular value thresholding," *Ann. Statist.*, vol. 42, no. 6, pp. 2413–2440, 2014.
- [38] Y. Chen and M. J. Wainwright, "Fast low-rank estimation by projected gradient descent: General statistical and algorithmic guarantees," 2015, *arXiv:1509.03025*. [Online]. Available: <http://arxiv.org/abs/1509.03025>
- [39] Y. Chen and E. J. Candès, "The projected power method: An efficient algorithm for joint alignment from pairwise differences," *Commun. Pure Appl. Math.*, vol. 71, no. 8, pp. 1648–1714, Aug. 2018.
- [40] D. Donoho, M. Gavish, and I. Johnstone, "Optimal shrinkage of eigenvalues in the spiked covariance model," *Ann. Statist.*, vol. 46, no. 4, p. 1742, Aug. 2018.
- [41] L. T. Liu, E. Dobriban, and A. Singer, "ePCA: High dimensional exponential family PCA," *Ann. Appl. Statist.*, vol. 12, no. 4, pp. 2121–2150, 2018.
- [42] I. M. Johnstone and D. Paul, "PCA in high dimensions: An orientation," *Proc. IEEE*, vol. 106, no. 8, pp. 1277–1292, Aug. 2018.
- [43] Y. Chen, J. Fan, C. Ma, and Y. Yan, "Bridging convex and nonconvex optimization in robust PCA: Noise, outliers, and missing data," *Ann. Statist.*, 2020.
- [44] Y. Yan, Y. Chen, and J. Fan, "Inference for heteroskedastic PCA with missing data," 2021, *arXiv:2107.12365*. [Online]. Available: <http://arxiv.org/abs/2107.12365>
- [45] J. Fan, W. Wang, and Y. Zhong, "An ℓ_∞ eigenvector perturbation bound and its application to robust covariance estimation," *J. Mach. Learn. Res.*, vol. 18, no. 207, pp. 1–42, 2018.
- [46] Y. Chen, J. Fan, C. Ma, and K. Wang, "Spectral method and regularized MLE are both optimal for top- K ranking," *Ann. Statist.*, vol. 47, no. 4, pp. 2204–2235, Aug. 2019.
- [47] J. Cape, M. Tang, and C. E. Priebe, "The two-to-infinity norm and singular subspace geometry with applications to high-dimensional statistics," *Ann. Statist.*, vol. 47, no. 5, pp. 2405–2439, Oct. 2019.
- [48] J. Eldridge, M. Belkin, and Y. Wang, "Unperturbed: Spectral analysis beyond Davis–Kahan," in *Proc. Algorithmic Learn. Theory*, 2018, pp. 321–358.
- [49] Y. Chen, Y. Chi, J. Fan, C. Ma, and Y. Yan, "Noisy matrix completion: Understanding statistical guarantees for convex relaxation via nonconvex optimization," *SIAM J. Optim.*, vol. 30, no. 4, pp. 3098–3121, Jan. 2020.
- [50] A. Pananjady and M. J. Wainwright, "Value function estimation in Markov reward processes: Instance-dependent ℓ_∞ -bounds for policy evaluation," 2019, *arXiv:1909.08749*. [Online]. Available: <https://arxiv.org/abs/1909.08749>
- [51] J. Chen, D. Liu, and X. Li, "Nonconvex rectangular matrix completion via gradient descent without $\ell_{2,\infty}$ regularization," *IEEE Trans. Inf. Theory*, vol. 66, no. 9, pp. 5806–5841, Sep. 2020.

- [52] L. Lei, “Unified $\ell_2 \rightarrow \infty$ eigenspace perturbation theory for symmetric random matrices,” 2019, *arXiv:1909.04798*. [Online]. Available: <https://arxiv.org/abs/1909.04798>
- [53] V. Koltchinskii and D. Xia, “Perturbation of linear forms of singular vectors under Gaussian noise,” in *High Dimensional Probability VII*. Birkhäuser, Cham: Springer, 2016, pp. 397–423.
- [54] C.-H. Zhang and S. S. Zhang, “Confidence intervals for low dimensional parameters in high dimensional linear models,” *J. Roy. Stat. Soc. B, Stat. Methodol.*, vol. 76, no. 1, pp. 217–242, 2014.
- [55] S. van de Geer, P. Bühlmann, Y. Ritov, and R. Dezeure, “On asymptotically optimal confidence regions and tests for high-dimensional models,” *Ann. Statist.*, vol. 42, no. 3, pp. 1166–1202, 2014.
- [56] A. Javanmard and A. Montanari, “Confidence intervals and hypothesis testing for high-dimensional regression,” *J. Mach. Learn. Res.*, vol. 15, no. 1, pp. 2869–2909, Jan. 2014.
- [57] T. T. Cai and Z. Guo, “Confidence intervals for high-dimensional linear regression: Minimax rates and adaptivity,” *Ann. Statist.*, vol. 45, no. 2, pp. 615–646, 2017.
- [58] A. Belloni, V. Chernozhukov, and C. Hansen, “Inference for high-dimensional sparse econometric models,” 2011, *arXiv:1201.0220*. [Online]. Available: <http://arxiv.org/abs/1201.0220>
- [59] Z. Ren, T. Sun, C.-H. Zhang, and H. H. Zhou, “Asymptotic normality and optimalities in estimation of large Gaussian graphical models,” *Ann. Statist.*, vol. 43, no. 3, pp. 991–1026, Jun. 2015.
- [60] Y. Ning and H. Liu, “A general theory of hypothesis tests and confidence regions for sparse high dimensional models,” *Ann. Statist.*, vol. 45, no. 1, pp. 158–195, 2017.
- [61] C. Ma, J. Lu, and H. Liu, “Inter-subject analysis: A partial Gaussian graphical model approach,” *J. Amer. Stat. Assoc.*, vol. 116, pp. 1–57, Apr. 2020.
- [62] J. Janková and S. van de Geer, “De-biased sparse PCA: Inference and testing for eigenstructure of large covariance matrices,” 2018, *arXiv:1801.10567*. [Online]. Available: <http://arxiv.org/abs/1801.10567>
- [63] L. Miolane and A. Montanari, “The distribution of the Lasso: Uniform control over sparse balls and adaptive parameter tuning,” *Ann. Statist.*, 2020.
- [64] M. Celentano, A. Montanari, and Y. Wei, “The lasso with general Gaussian designs with applications to hypothesis testing,” 2020, *arXiv:2007.13716*. [Online]. Available: <http://arxiv.org/abs/2007.13716>
- [65] A. Carpentier, J. Eisert, D. Gross, and R. Nickl, “Uncertainty quantification for matrix compressed sensing and quantum tomography problems,” in *High Dimensional Probability VIII*. Birkhäuser, Cham: Springer, 2019, pp. 385–430.
- [66] A. Carpentier and R. Nickl, “On signal detection and confidence sets for low rank inference problems,” *Electron. J. Statist.*, vol. 9, no. 2, pp. 2675–2688, Jan. 2015.
- [67] A. Carpentier, O. Klopp, and M. Löffler, “Constructing confidence sets for the matrix completion problem,” in *Proc. Conf. Int. Soc. Non-Parametric Statist.* Cham, Switzerland: Springer, 2016, pp. 103–118.
- [68] Y. Chen, J. Fan, C. Ma, and Y. Yan, “Inference and uncertainty quantification for noisy matrix completion,” *Proc. Nat. Acad. Sci. USA*, vol. 116, no. 46, pp. 22931–22937, Nov. 2019.
- [69] D. Xia and M. Yuan, “Statistical inferences of linear forms for noisy matrix completion,” *J. Roy. Stat. Soc. B, Stat. Methodol.*, vol. 83, no. 1, pp. 58–77, Feb. 2021.
- [70] D. Xia, “Confidence region of singular subspaces for low-rank matrix regression,” *IEEE Trans. Inf. Theory*, vol. 65, no. 11, pp. 7437–7459, Nov. 2019.
- [71] D. Xia, “Normal approximation and confidence region of singular subspaces,” 2019, *arXiv:1901.00304*. [Online]. Available: <http://arxiv.org/abs/1901.00304>
- [72] J. Fan, Y. Fan, X. Han, and J. Lv, “Asymptotic theory of eigenvectors for random matrices with diverging spikes,” 2019, *arXiv:1902.06846*. [Online]. Available: <http://arxiv.org/abs/1902.06846>
- [73] H. Sommers, A. Crisanti, H. Sompolinsky, and Y. Stein, “Spectrum of large random asymmetric matrices,” *Phys. Rev. Lett.*, vol. 60, no. 19, p. 1895, 1988.
- [74] B. Khoruzhenko, “Large- N eigenvalue distribution of randomly perturbed asymmetric matrices,” *J. Phys. A, Math. Gen.*, vol. 29, no. 7, p. L165, 1996.
- [75] E. Brézin and A. Zee, “Non-Hermitian delocalization: Multiple scattering and bounds,” *Nucl. Phys. B*, vol. 509, no. 3, pp. 599–614, Jan. 1998.
- [76] J. T. Chalker and B. Mehlh, “Eigenvector statistics in non-Hermitian random matrix ensembles,” *Phys. Rev. Lett.*, vol. 81, no. 16, p. 3367, 1998.
- [77] J. Feinberg and A. Zee, “Non-Hermitian random matrix theory: Method of Hermitian reduction,” *Nucl. Phys. B*, vol. 504, no. 3, pp. 579–608, Nov. 1997.
- [78] A. Lytova and K. Tikhomirov, “On delocalization of eigenvectors of random non-Hermitian matrices,” 2018, *arXiv:1810.01590*. [Online]. Available: <http://arxiv.org/abs/1810.01590>
- [79] T. Tao, “Outliers in the spectrum of iid matrices with bounded rank perturbations,” *Probab. Theory Rel. Fields*, vol. 155, nos. 1–2, pp. 231–263, 2013.
- [80] A. B. Rajagopalan, “Outlier eigenvalue fluctuations of perturbed iid matrices,” 2015, *arXiv:1507.01441*. [Online]. Available: <http://arxiv.org/abs/1507.01441>
- [81] F. Benaych-Georges and J. Rochet, “Outliers in the single ring theorem,” *Probab. Theory Rel. Fields*, vol. 165, pp. 313–363, Jun. 2016.
- [82] J. B. Conway, *Functions of One Complex Variable II*, vol. 159. New York, NY, USA: Springer, 2012.
- [83] M. Embree and L. N. Trefethen, “Generalizing eigenvalue theorems to pseudospectra theorems,” *SIAM J. Sci. Comput.*, vol. 23, no. 2, pp. 583–590, Jan. 2001.
- [84] J. A. Tropp, “An introduction to matrix concentration inequalities,” *Found. Trends Mach. Learn.*, vol. 8, nos. 1–2, pp. 1–230, May 2015.
- [85] A. B. Tsybakov, *Introduction to Nonparametric Estimation*. New York, NY, USA: Springer, 2008.
- [86] L. H. Chen, L. Goldstein, and Q.-M. Shao, *Normal Approximation by Stein’s Method*. Heidelberg, Germany: Springer, 2010.

Chen Cheng received the bachelor’s degree from the School of Mathematical Sciences, Peking University. He is currently pursuing the Ph.D. degree with the Department of Statistics, Stanford University, jointly advised by Prof. John Duchi and Prof. Andrea Montanari. His research interests include statistical theory and algorithms for high-dimensional data, random matrix theory, machine learning theory, and reinforcement learning.

Yuting Wei received the Ph.D. degree in statistics from the University of California, Berkeley (UC Berkeley), advised by Martin Wainwright and Aditya Guntuboyina. She is currently an Assistant Professor with the Department of Statistics and Data Science, The Wharton School, University of Pennsylvania. Prior to that, she spent two years at Carnegie Mellon University as an Assistant Professor of statistics, and one year at Stanford University as a Stein Fellow. Her research interests include high-dimensional and non-parametric statistics, statistical machine learning, and reinforcement learning. She was a recipient of the 2018 Erich L. Lehmann Citation from the Department of Statistics, UC Berkeley, for her Ph.D. dissertation in theoretical statistics.

Yuxin Chen (Member, IEEE) received the B.E. degree (Hons.) from Tsinghua University in 2008, the M.S. degree in electrical and computer engineering from The University of Texas at Austin in 2010, and the M.S. degree in statistics and the Ph.D. degree in electrical engineering from Stanford University in January 2015. He is currently an Assistant Professor with the Department of Electrical Engineering, Princeton University, Princeton, NJ, USA. Prior to joining Princeton, he was a Post-Doctoral Scholar with the Department of Statistics, Stanford University. His research interests include high-dimensional data analysis, convex and non-convex optimization, reinforcement learning, statistical learning, statistical signal processing, and information theory. He was a recipient of the 2020 Princeton Graduate Mentoring Award, the 2019 AFOSR Young Investigator Award, the 2020 ARO Young Investigator Award, the ICCM Best Paper Award (Gold Medal), and the 2021 Princeton SEAS Junior Faculty Award. He was selected as a Finalist for the Best Paper Prize for Young Researchers in Continuous Optimization 2019.