

Optimally Weighted PCA for High-Dimensional Heteroscedastic Data*

David Hong[†] Fan Yang[‡] Jeffrey A. Fessler[§] Laura Balzano[§]

September 14, 2022

Abstract

Modern data are increasingly both high-dimensional and heteroscedastic. This paper considers the challenge of estimating underlying principal components from high-dimensional data with noise that is heteroscedastic across samples, i.e., some samples are noisier than others. Such heteroscedasticity naturally arises, e.g., when combining data from diverse sources or sensors. A natural way to account for this heteroscedasticity is to give noisier blocks of samples less weight in PCA by using the leading eigenvectors of a weighted sample covariance matrix. We consider the problem of choosing weights to optimally recover the underlying components. In general, one cannot know these optimal weights since they depend on the underlying components we seek to estimate. However, we show that under some natural statistical assumptions the optimal weights converge to a simple function of the signal and noise variances for high-dimensional data. Surprisingly, the optimal weights are not the inverse noise variance weights commonly used in practice. We demonstrate the theoretical results through numerical simulations and comparisons with existing weighting schemes. Finally, we briefly discuss how estimated signal and noise variances can be used when the true variances are unknown, and we illustrate the optimal weights on real data from astronomy.

Key words. principal component analysis, large-dimensional data, heterogeneous quality, optimal weighting
AMS subject classifications. 62H25

1 Introduction

Principal Component Analysis (PCA) is a fundamental technique for discovering underlying components in data and is a workhorse method for analyzing modern *high-dimensional* data. However, conventional PCA does not recover underlying principal components well when the data has *heteroscedastic* noise, as is common in practice. In particular, its performance can degrade substantially when the noise is heteroscedastic across samples, i.e., some samples are noisier than others. PCA suffers from treating all the samples uniformly with performance held back by the noisiest samples, as was rigorously characterized in [19]. Weighted PCA addresses this shortcoming by giving less weight to lower quality samples. This naturally raises a crucial question: how should the weights be chosen? Namely, *what are the optimal weights?*

*D. Hong was supported in part by NSF Graduate Research Fellowship DGE 1256260, NSF Grant ECCS-1508943, NSF BIGDATA grant IIS 1837992, the Dean's Fund for Postdoctoral Research of the Wharton School, and NSF Mathematical Sciences Postdoctoral Research Fellowship DMS 2103353. F. Yang was supported in part by the Wharton Dean's Fund for Postdoctoral Research. J. A. Fessler was supported in part by the UM-SJTU data science seed fund, NSF grant IIS 1838179, and NIH Grant U01 EB 018753. L. Balzano was supported in part by DARPA-16-43-D3M-FP-037, by NSF CAREER award CCF-1845076, by NSF Grant ECCS-1508943 and by ARO YIP Award W911NF1910027.

[†]Department of Statistics and Data Science, Wharton School, University of Pennsylvania, Philadelphia, PA, 19104 USA (dahong67@wharton.upenn.edu).

[‡]Yau Mathematical Sciences Center, Tsinghua University, Beijing, 100084 China (fyangmath@mail.tsinghua.edu.cn).

[§]Department of Electrical Engineering and Computer Science, University of Michigan, Ann Arbor, MI, 48109 USA (fessler@umich.edu, girasole@umich.edu).

This paper addresses this question by rigorously deriving optimal weights that are simple functions of the signal and noise variances. Surprisingly, they are not the inverse noise variance weights that are commonly used in practice. We now elaborate in more detail.

1.1 High-dimensional and heteroscedastic data

Modern applications of PCA span numerous and diverse areas across all of engineering and the sciences, ranging from medical imaging [3, 42] to cancer data classification [45], genetics [34], and environmental sensing [39, 50], to name just a few. Increasingly, the number of features measured is comparable to or even larger than the number of samples, i.e., the data are *high-dimensional*. Traditional asymptotic analysis of the performance of methods as the number of samples grows (with a fixed number of features) do not apply well to such settings. Modern data analysis needs new theory and methods for the high-dimensional regime where both the number of features and number of samples are large [27].

Modern datasets are also frequently composed of samples with heteroscedastic (i.e., heterogeneous) noise. In particular, we consider noise that is *heteroscedastic across samples*, namely, some samples are noisier than others. Such data arises naturally when samples are obtained at varying times or by varying means or equipment. For example, in the field of analytical chemistry, [10] considers spectrophotometric data obtained from averages taken over varying windows of time; samples from shorter windows are noisier. As another example, in the field of air quality monitoring, samples come from various sources: governments agencies provide low-noise data obtained from carefully operated instruments, while individuals provide noisier data obtained from cheaper and easy to setup sensors [17, 43]. As a final example, in the field of astronomy, measurements of astronomical objects such as stars and quasars can have various levels of noise due to atmospheric and detector effects that vary from object to object [4, 47, 48]. More generally, modern big data analysis is often performed using datasets built up by combining myriad sources, so one can expect that data with heteroscedastic noise will be the norm. Modern data analysis needs PCA methods that effectively account for this type of heteroscedasticity. Indeed, such methods may also unlock new opportunities to effectively leverage new sources of data with heteroscedastic noise.

1.2 Weighted PCA

Weighted PCA accounts for heteroscedastic noise by giving smaller weight to noisier samples. Analogous to unweighted PCA, the principal components are the leading eigenvectors of the *weighted sample covariance matrix*

$$\hat{\Sigma}_{\mathbf{w}} := \sum_{\ell=1}^L w_{\ell} \mathbf{Y}_{\ell} \mathbf{Y}_{\ell}^{\mathbf{H}},$$

where $\mathbf{Y}_1, \dots, \mathbf{Y}_L$ are L blocks of samples with associated noise variances $v_1, \dots, v_L > 0$, the superscript $\mathbf{H}$ denotes the Hermitian transpose, and $w_1, \dots, w_L \geq 0$ are the weights. Existing choices for the weights include:

- **uniform weights** ($w_{\ell} = 1$): these weights correspond to unweighted PCA and can be a natural choice when the noise is close to homoscedastic. However, its performance degrades with increasing noise heteroscedasticity, as was shown in [19, Theorem 2].
- **binary weights** ($w_{\ell} = 1$ for less noisy blocks and $w_{\ell} = 0$ for the rest): these weights correspond to performing unweighted PCA using only less noisy blocks of samples and are a natural choice when some samples are much noisier than the rest. The idea is to exclude noisier samples that do more harm than good. However, doing so also omits any useful information that was in the excluded samples. How to decide if a block of samples is better to include or exclude can be unclear.
- **inverse noise variance weights** ($w_{\ell} = 1/v_{\ell}$): these weights whiten the noise, making it homoscedastic, and can be interpreted as a maximum likelihood weighting [53]. They are a natural way to account for noise heteroscedasticity while using all the samples and are commonly used in practice.

It has been unclear which of these existing options to choose and whether any are optimal.

1.3 Contribution of this paper

The main contribution of this paper is *optimal weights* for high-dimensional heteroscedastic data, that are rigorously derived under some natural statistical assumptions. Roughly put, we show that when both dimensions of the data are large, the optimal weights converge to the following simple *asymptotic optimal weights*:

$$w_\ell = \frac{1}{v_\ell} \frac{1}{1 + v_\ell/\lambda},$$

where v_ℓ is the noise variance and λ is the signal variance. Notably, these weights are inverse noise variance weights scaled by a simple term that depends on the noise-to-signal ratio v_ℓ/λ . See [section 2](#) for the precise statement of the result and [section 6](#) for its proof.

Naturally, one wonders how well these results apply for data with finitely many samples and features. Numerical simulations in [section 3](#) illustrate that the optimal weights in finite dimensions (which are a function of the random signal coefficients and noise) concentrate around the asymptotic optimal weights as the data grows in size. As a result, the asymptotic optimal weights are often close to the optimal weights in finite dimensions when the dimensions are large enough.

We also compare the asymptotic optimal weights with the existing weights above: uniform, binary, and inverse noise variance weights. In particular, we consider how close they are to the optimal weights in finite dimensions ([subsection 4.1](#)), how well they perform in finite dimensions ([subsection 4.2](#)), and in what regimes they achieve positive asymptotic recovery ([subsection 4.3](#)). Overall, the asymptotic optimal weights outperform the existing weights.

One also wonders how to calculate the asymptotic optimal weights when the signal and noise variances are unknown. Naturally, one might consider simply using estimators of these variances. We explain that the resulting estimated weights are also asymptotically optimal as long as the estimators are consistent, and we give an example of such estimators ([section 7](#)).

Finally, we illustrate the asymptotic optimal weights on real data ([section 8](#)). The data are quasar spectra measured by the Sloan Digital Sky Survey and have heteroscedastic noise. The example exhibits some of the main themes of the paper and illustrates the potential for optimally weighted PCA to improve performance in real data.

1.4 Related works

Previous work on PCA for noise that is heteroscedastic (whether across samples or otherwise) have addressed various important questions, as elaborated below. However, to the best of our knowledge, the important question of optimal weighting was not previously considered. This paper rigorously answers the question of optimal weighting for noise that is heteroscedastic across samples. It will be interesting for future works to consider this question for other forms of heteroscedastic noise.

Some of our other papers considered various aspects of noise that is heteroscedastic across samples. In particular, [\[18, 19\]](#) derive the asymptotic performance of unweighted PCA and characterize the impact of heteroscedasticity. See [\[19, Sections 1.3 and 2.3\]](#) for a discussion of the connections to previous analyses of PCA for homoscedastic noise (such as [\[25, 38, 41\]](#)), and see [\[19, Section S1\]](#) for a discussion of the connections to spiked covariance models. Alternatively, [\[20, 21\]](#) consider a probabilistic PCA approach, where the noise heteroscedasticity is modeled via the statistical likelihood. The resulting method is not a weighted PCA. Instead, one must solve a challenging optimization problem, and [\[20, 21\]](#) develop several algorithms for this purpose.

A closely related model for heterogeneous data arises in the context of low-rank clutter estimation for RADAR. In this setting, the noise is homoscedastic but the clutter signal has heterogeneous strengths, i.e., the clutter covariances are a common low-rank matrix scaled by heterogeneous power factors. Maximum likelihood estimation of the common low-rank matrix and the power factors involves solving a challenging

optimization problem, and [8, 9, 46, 11] develop efficient algorithms for this purpose. The estimation performance is analyzed in [6], and [1] considers maximum a posteriori estimation. This heterogeneous signal strength model is related to the heteroscedastic noise model in the present paper through a straightforward rescaling of the data: scaling each sample by the inverse of its power factor yields the model in the present paper. Thus, the optimally weighted PCA developed here can be straightforwardly modified to apply to the heterogeneous signal strength model; see [Appendix A](#) for details.

Several recent works develop PCA variants for high-dimensional data with noise that is heteroscedastic across features. In contrast to the samplewise heteroscedasticity we consider, featurewise heteroscedasticity produces a nonuniform bias along the diagonal of the covariance matrix that skews its eigenvectors even with infinitely many samples. An approach based on spectral shrinkage with noise whitening (to make it homoscedastic) is developed in [32]; the noise is whitened by weighting the features by their inverse noise variance. Whitening both the features and samples is considered in [33], and whitening in the context of linearly transformed signals is considered in [15]. Alternatively, [54] addresses the bias in the covariance matrix by iteratively replacing its biased diagonal entries using low-rank approximation. Estimating the number of underlying principal components is another important problem in this setting, and recent works [22, 29, 31] have developed new methods for tackling this challenge under heteroscedastic noise. Estimated principal components are also often combined with estimates of associated signal variances to obtain estimates of an underlying signal matrix or covariance. For homoscedastic noise, existing works have made tremendous progress on how to estimate these signal variances to optimize various objectives, typically by applying a carefully designed shrinkage to the eigenvalues of the sample covariance matrix; see, e.g., [16] and the references therein. A few recent works address this question in the context of heteroscedastic noise: [32] derives optimal shrinkages for use with whitening, [33] derives optimal spectral denoisers, and [37] derives an optimal data-driven shrinkage.

Many works have considered weighted PCA methods in general; see [28, Section 14.2.1] for a survey of some of these works. For example, [12, Sections 5.4–5.5] discusses weighting features by inverse noise variance weights to account for featurewise heteroscedasticity. Weighting both samples and features is proposed in [10] for analyzing spectrophotometric data from scanning wavelength kinetics experiments; the weights are again inverse noise variance. Similar schemes have also been proposed in metabolomics [24] and astronomy [4, 47], to name just a few areas. Weighting data by inverse noise variance weights has been a recurring theme.

Overall, previous work on PCA for heteroscedastic noise made significant progress on various important questions. However, the important question of optimal weighting was not previously considered. This paper addresses that question for noise that is heteroscedastic across samples.

1.5 Organization of the paper

[Section 2](#) states our main result: optimal weights and performance for high-dimensional data with heteroscedastic noise. [Section 3](#) performs numerical simulations in finite dimensions, and [section 4](#) compares the asymptotic optimal weights with existing weighting schemes: inverse noise variance weighted PCA, PCA using only a single block of the data, and unweighted PCA. [Section 5](#) compares optimally weighted PCA with some additional methods. [Section 6](#) proves the main result. The optimal weights depend on the signal and noise variances. [Section 7](#) describes how estimates of these variances can be used when the true variances are unknown. [Section 8](#) illustrates optimally weighted PCA on real data coming from astronomy. Codes for reproducing the figures in this paper are available online at: <https://gitlab.com/dahong/optimally-weighted-pca-heteroscedastic-data>

For readers mostly interested in understanding the underlying theory and proofs of the main result, we suggest starting with [sections 2](#) and [6](#). For readers mostly interested in using optimally weighted PCA, we suggest starting with [sections 2](#) to [5](#), [7](#) and [8](#).

2 Main result: optimal weights and performance

We begin by making precise the notion of optimal weights and optimal performance. Consider a dataset $\mathbf{Y}$ having k underlying orthonormal components $\mathbf{u}_1, \dots, \mathbf{u}_k$, where $\mathbf{Y}$ is made of L blocks $\mathbf{Y}_1, \dots, \mathbf{Y}_L$ of samples with heteroscedastic noise. Then, given weights $w_1, \dots, w_L \geq 0$, the $\mathbf{w}$ -weighted PCA estimate of the i th component $\mathbf{u}_i$ from $\mathbf{Y}$, denoted $\hat{\mathbf{u}}_i(\mathbf{w}, \mathbf{Y})$, is

$$\hat{\mathbf{u}}_i(\mathbf{w}, \mathbf{Y}) := i\text{th leading eigenvector of the weighted sample covariance } \sum_{\ell=1}^L w_\ell \mathbf{Y}_\ell \mathbf{Y}_\ell^H. \quad (2.1)$$

A natural way to measure the performance of the estimate, i.e., how well $\hat{\mathbf{u}}_i(\mathbf{w}, \mathbf{Y})$ recovers the i th component $\mathbf{u}_i$, is by the square inner product $r_i(\mathbf{w}, \mathbf{Y})$ given by

$$r_i(\mathbf{w}, \mathbf{Y}) := |\mathbf{u}_i^H \hat{\mathbf{u}}_i(\mathbf{w}, \mathbf{Y})|^2. \quad (2.2)$$

Finally, optimal weights $\mathbf{w}_i^*(\mathbf{Y})$ and the optimal performance $r_i^*(\mathbf{Y})$ for the i th component are defined by

$$\mathbf{w}_i^*(\mathbf{Y}) \in \operatorname{argmax}_{\mathbf{w}} r_i(\mathbf{w}, \mathbf{Y}), \quad r_i^*(\mathbf{Y}) = \max_{\mathbf{w}} r_i(\mathbf{w}, \mathbf{Y}). \quad (2.3)$$

Note that the performance (2.2) depends on the underlying component $\mathbf{u}_i$. However, in practice, $\mathbf{u}_i$ is of course unknown so the optimization (2.3) cannot be done. Fortunately, as our main result below shows, the optimal weights $\mathbf{w}_i^*(\mathbf{Y})$ and optimal performance $r_i^*(\mathbf{Y})$ can be predicted when the data (a) satisfies some natural statistical assumptions, and (b) grows large in size, i.e., under the following setting.

Setting: We will assume the following setting throughout the remainder of the paper.

- (a) The noisy data blocks $\mathbf{Y}_1 \in \mathbb{C}^{d \times n_1}, \dots, \mathbf{Y}_L \in \mathbb{C}^{d \times n_L}$ are generated from the components $\mathbf{u}_1, \dots, \mathbf{u}_k \in \mathbb{C}^d$ with corresponding signal variances $\lambda_1 > \dots > \lambda_k > 0$ as follows:

$$\mathbf{Y}_\ell = \mathbf{F} \mathbf{Z}_\ell + \mathbf{E}_\ell \in \mathbb{C}^{d \times n_\ell}, \quad \text{for } \ell = 1, \dots, L, \quad (2.4)$$

where

- $\mathbf{F} := [\sqrt{\lambda_1} \mathbf{u}_1, \dots, \sqrt{\lambda_k} \mathbf{u}_k] \in \mathbb{C}^{d \times k}$ is a deterministic factor matrix common to all the blocks,
- $\mathbf{Z}_\ell \in \mathbb{C}^{k \times n_\ell}$ is a coefficient matrix with IID entries having zero mean and unit variance,
- $\mathbf{E}_\ell \in \mathbb{C}^{d \times n_\ell}$ is a noise matrix with IID entries having zero mean and variance $v_\ell > 0$,

and the noise entries further satisfy a technical condition: bounded a -th moment with $a > 4$, i.e., $\exists_{a>4}$ s.t. $\mathbb{E}[|\mathbf{E}_\ell|_{i,j}|^a] < \infty$.¹ Note that this model also includes real-valued data with real-valued coefficients and noise.

- (b) The number of features d and numbers of samples $n_1, \dots, n_L$ all grow towards infinity but with fixed aspect ratios $n_\ell/d = c_\ell > 0$. This asymptotic regime captures datasets where the number of features and samples are roughly comparable, as is common in modern big data settings.

Note that under the model (2.4), the optimal weights and performance (2.3) are random quantities so their convergence will be probabilistic. Specifically, the convergence holds with probability one, i.e., it is *almost sure convergence*, which we will denote by $\xrightarrow{\text{a.s.}}$.

We are now ready to state the main result on the optimal weights and performance.

¹This technical condition on the noise is satisfied by numerous distributions including the sub-Gaussian and sub-Exponential families [49, Propositions 2.5.2 and 2.7.1].

Theorem 2.1 (Asymptotic optimal weights and performance). *The optimal weights $\mathbf{w}_i^*(\mathbf{Y})$ and corresponding optimal performance $r_i^*(\mathbf{Y})$ converge as*

$$\mathbf{w}_i^*(\mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{\mathbf{w}}_i^* := \left(\frac{1}{v_1} \frac{1}{1 + v_1/\lambda_i}, \dots, \frac{1}{v_L} \frac{1}{1 + v_L/\lambda_i} \right) \quad \text{up to scaling}, \quad (2.5)$$

$$r_i^*(\mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{r}_i^* := \text{the unique solution } x \in (0, 1) \text{ of } \sum_{\ell=1}^L \frac{c_\ell}{v_\ell/\lambda_i} \frac{1-x}{v_\ell/\lambda_i + x} = 1, \quad (2.6)$$

except when $\sum_{\ell=1}^L c_\ell (\lambda_i/v_\ell)^2 \leq 1$, in which case $\mathbf{u}_i$ is asymptotically unrecoverable by any weighted PCA, i.e., $r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} 0$ for all weights $\mathbf{w}$.

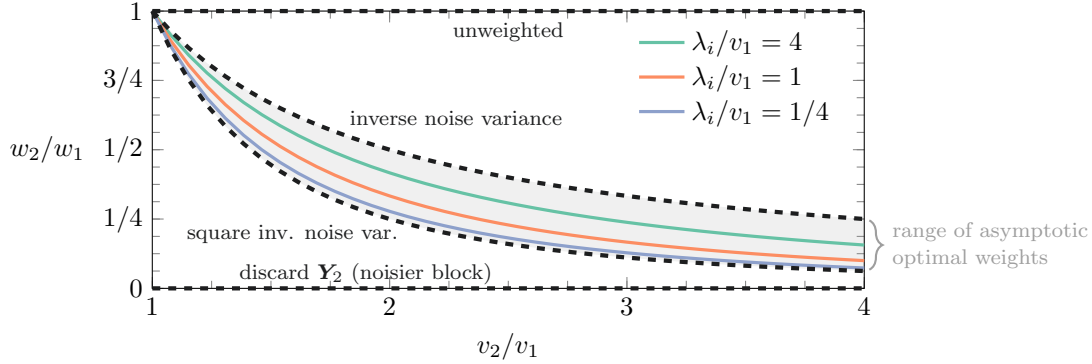


Figure 1. Relative weight w_2/w_1 given by the optimal weights (2.5) as a function of the relative noise variance v_2/v_1 for various signal-to-noise ratios λ_i/v_1 . The optimal weights downweight noisier data more aggressively than inverse noise variance weights, but also do not discard noisier data. They lie in the region between inverse and square inverse noise variance weights.

Remark 2.2 (Optimal weights downweight more than inverse noise variance weights). The optimal weights (2.5) are not the inverse noise variance weights that are commonly used. As illustrated in Figure 1, optimal weights downweight noisier data more aggressively, but never discard data. Specifically, when the signal-to-noise ratio λ_i/v_ℓ is small, the optimal weights are square inverse noise variance weights up to scale. As λ_i/v_ℓ grows, the optimal weights gradually become less aggressive and approach inverse noise variance weights.

Remark 2.3 (Using estimated signal and noise variances). The optimal weights (2.5) depend on the noise variances $\mathbf{v}$ and on the signal component variance λ_i . In some settings, these parameters are known, e.g., from calibration data. When they are unknown, one may estimate them using existing ideas and approaches, then plug them in to obtain estimated weights. As discussed in section 7, these estimated weights are also asymptotically optimal as long as the variance estimates are consistent.

Remark 2.4 (Heterogeneous signal strengths). The result may at first appear to be limited to data with homogeneous signal strengths. However, it generalizes straightforwardly to the case where the signal in each block is scaled by an associated signal strength, as arises, e.g., in RADAR applications [8, 9, 46]. Simply preweight the data to recover the model (2.4); see Appendix A for a detailed description.

Remark 2.5 (Handling potentially degenerate cases). A careful reader may note the subtle and technical point that there may exist degenerate choices of $\mathbf{w}$ for which $r_i(\mathbf{w}, \mathbf{Y})$ is not well-defined, e.g., if the i th leading eigenvector becomes undefined due to eigenvalue multiplicity. At such points, we define $r_i(\mathbf{w}, \mathbf{Y})$ by its limsup over $\mathbf{w}$. Doing so makes $r_i(\mathbf{w}, \mathbf{Y})$ upper semi-continuous in $\mathbf{w}$ and avoids degenerate situations where its maximum does not exist.

Remark 2.6 (Nonorthogonality of estimated components). Since the optimal weights (2.5) are component-specific, the components $\hat{\mathbf{u}}_1(\bar{\mathbf{w}}_1^*, \mathbf{Y}), \dots, \hat{\mathbf{u}}_k(\bar{\mathbf{w}}_k^*, \mathbf{Y})$ estimated by optimally weighted PCA may not be orthogonal in practice. In applications where orthogonality is crucial, one option is to sacrifice component-wise optimality and use a single set of weights, e.g., that just optimizes recovery of the weakest component or that optimizes some appropriate overall metric of performance. Alternatively, in many such cases, the principal subspace is of greater interest than the individual components; in these cases, one could orthogonalize the components, e.g., via Gram-Schmidt.

Remark 2.7 (Phase transition). Analogous to unweighted PCA under homoscedastic noise, optimally weighted PCA exhibits a phase transition between settings with zero asymptotic performance and those with nonzero asymptotic performance. As described in Theorem 2.1, optimally weighted PCA has nonzero asymptotic performance when $\sum_{\ell=1}^L c_\ell(\lambda_i/v_\ell)^2 > 1$ (or in other words, $\lambda_i > (\sum_{\ell=1}^L c_\ell/v_\ell^2)^{-1/2}$). Notably, if any weighting scheme has nonzero asymptotic performance, then optimally weighted PCA does too, as illustrated in subsection 4.3.

Before proving the main result (Theorem 2.1) in section 6, we provide some more intuition about it through numerical simulations in finite dimensions (section 3) and comparisons with existing weighting schemes (section 4).

3 Numerical simulation

This section performs numerical simulations in finite dimensions. Specifically, we generate $L = 2$ blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ according to the model (2.4) with

- $k = 1$ component $\mathbf{u}_1 \in \mathbb{R}^d$ uniformly drawn from the unit sphere,
- Gaussian coefficients $(\mathbf{Z}_\ell)_{ij} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1)$ and noise entries $(\mathbf{E}_\ell)_{ij} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, v_\ell)$,
- component variance $\lambda_1 = 1$ and noise variances $\mathbf{v} = (1, 3)$.

Figure 2 shows the nonasymptotic empirical distributions of the optimal weights $\mathbf{w}_1^*(\mathbf{Y})$ and the corresponding optimal performance $r_1^*(\mathbf{Y})$ from (2.3) obtained using the true underlying component $\mathbf{u}_1$. Since the weights are meaningful only up to scale, we show the optimal relative weight $w_{1,2}^*(\mathbf{Y})/w_{1,1}^*(\mathbf{Y})$, where $w_{i,\ell}^*(\mathbf{Y})$ is the ℓ th entry of $\mathbf{w}_i^*(\mathbf{Y})$. Similarly, $\bar{w}_{i,\ell}^*$ denotes the ℓ th entry of $\bar{\mathbf{w}}_i^*$.

Note first that the nonasymptotic distributions for both the optimal weights $\mathbf{w}_1^*(\mathbf{Y})$ and the optimal performance $r_1^*(\mathbf{Y})$ are generally centered around their respective theoretical limits $\bar{\mathbf{w}}_1^*$ and $\bar{r}_1^*$ from (2.5) and (2.6). Moreover, they concentrate as the data grows in size from Figure 2a to Figure 2b. This illustrates the almost sure convergence of the nonasymptotic optimal weights and performance to their limits.

Naturally, one also wonders whether the asymptotic results of Theorem 2.1 can be used to choose optimal weights or predict optimal performance for real data, which are finite-dimensional. These experiments demonstrate that the asymptotic optimal weights (2.5) and performance (2.6) can indeed be applied to choose weights that are often close to optimal for finite-dimensional data and to predict their corresponding performance.

4 Comparison with existing weighting schemes

This section compares the asymptotic optimal weights of (2.5) with existing weighting schemes. To ease discussion, we will focus on the case with only two blocks $\mathbf{Y}_1$ and $\mathbf{Y}_2$, i.e., $L = 2$, but the same insights apply

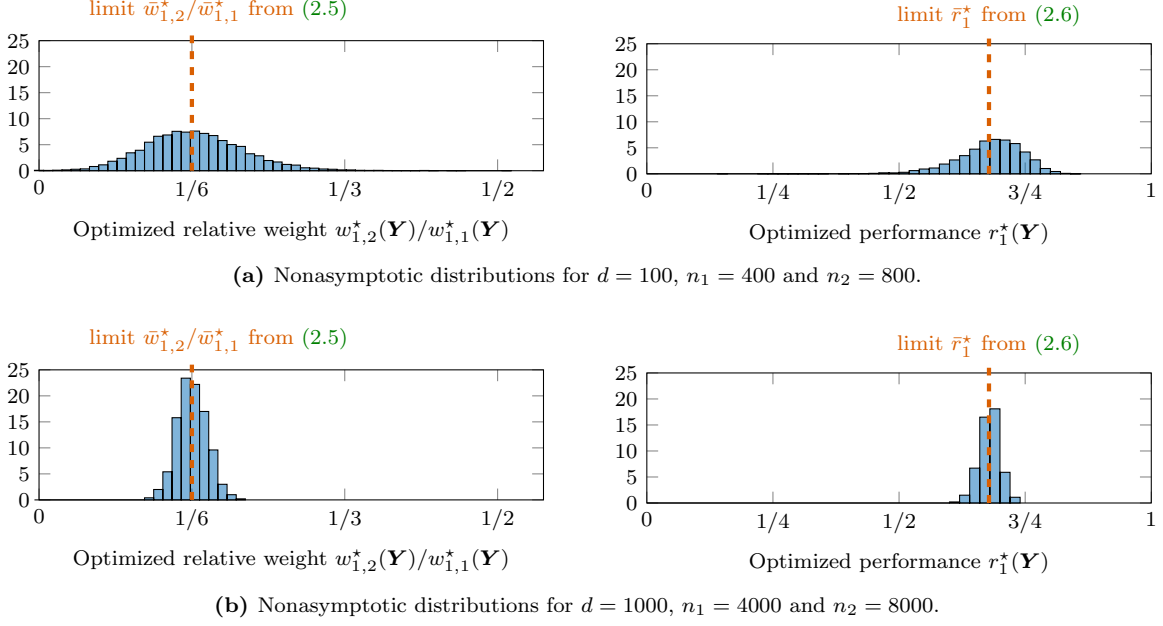


Figure 2. Nonasymptotic empirical distributions of optimal weights $\mathbf{w}_1^*(\mathbf{Y})$ and optimal performance $r_1^*(\mathbf{Y})$ from (2.3) for an illustrative example with two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ generated with noise variances $v_1 = 1$ and $v_2 = 3$, and with one underlying component having variance $\lambda_1 = 1$.

more generally. The weighting schemes considered are:

$$\text{inverse noise variance:} \quad \mathbf{w} = (1/v_1, 1/v_2), \quad (4.1)$$

$$\text{only use } \mathbf{Y}_1: \quad \mathbf{w} = (1, 0), \quad (4.2)$$

$$\text{only use } \mathbf{Y}_2: \quad \mathbf{w} = (0, 1), \quad (4.3)$$

$$\text{unweighted:} \quad \mathbf{w} = (1, 1), \quad (4.4)$$

where the weights are, as always, meaningful only up to scale. Note that using only $\mathbf{Y}_1$ or $\mathbf{Y}_2$ are both special cases of general binary weights that discard blocks of data.

4.1 Comparison of weights

This section compares how close the various weighting schemes are to the distribution of the actual empirically optimized weights $\mathbf{w}_i^*(\mathbf{Y})$ from (2.3) in an illustrative example. As in section 3, $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ are generated from the model (2.4) with Gaussian coefficients and noise, and a single signal component $\mathbf{u}_1 \in \mathbb{R}^d$ drawn uniformly from the unit sphere. The noise variances are $\mathbf{v} = (1, 3)$, and the data sizes are $d = 1000$, $n_1 = 4000$ and $n_2 = 8000$.

Figure 3 shows the nonasymptotic distribution of the empirically optimized weights (2.3) with the asymptotic optimal weights $\bar{\mathbf{w}}_1^*$ from (2.5) and the existing weights (4.1)–(4.4) overlaid. Figure 3a shows a case with a moderate signal component variance $\lambda_1 = 1$, and Figure 3b shows a case with a strong signal component $\lambda_1 = 30$. Since the weights are meaningful only up to scale, we show the relative weight $w_{1,2}^*(\mathbf{Y})/w_{1,1}^*(\mathbf{Y})$ given to the noisier block $\mathbf{Y}_2$, where $w_{i,\ell}^*(\mathbf{Y})$ is the ℓ th entry of $\mathbf{w}_i^*(\mathbf{Y})$, as in Figure 2.

We make the following observations from Figure 3:

- As before, the optimal weights are centered around the asymptotic optimal weights.
- When the signal component variance is moderate, the optimal weights do not overlap with the existing

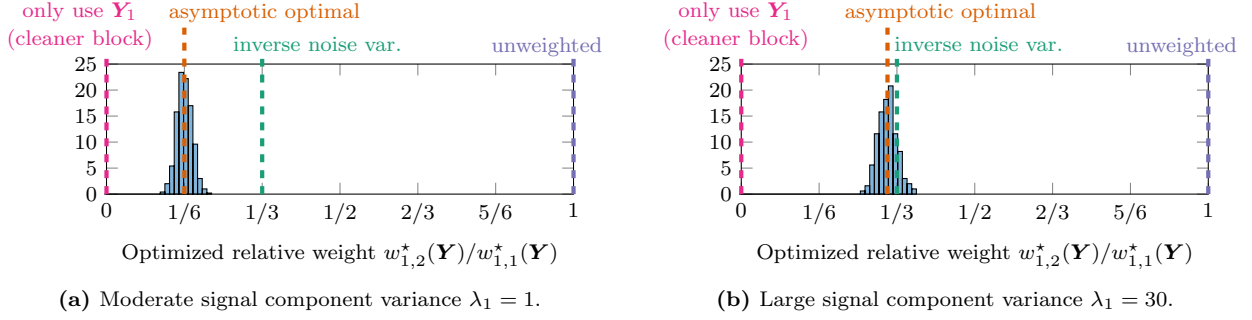


Figure 3. Comparison of weighting schemes for an illustrative example with two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ generated with noise variances $v_1 = 1$ and $v_2 = 3$, where $d = 1000$, $n_1 = 4000$ and $n_2 = 8000$. The nonasymptotic distributions of empirically optimized weights $\mathbf{w}_1^*(\mathbf{Y})$ from (2.3) are shown as histograms, with the asymptotic optimal weights $\bar{\mathbf{w}}_1^*$ from (2.5) and the existing weights (4.1)–(4.4) overlaid as lines (the weights (4.3) that only use $\mathbf{Y}_2$ do not appear since they are at infinity).

weighting schemes. They more aggressively downweight the noisier block than inverse noise variance weights but also do not discard the noisier block.

- When the signal component variance is large, the optimal weights overlap with inverse noise variance weights.

4.2 Comparison of performance

This section compares the various weighting schemes in terms of their performance $r_i(\mathbf{w}, \mathbf{Y})$ in finite dimensions. As in section 3, $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ are generated from the model (2.4) with Gaussian coefficients and noise, and a single signal component $\mathbf{u}_1 \in \mathbb{R}^d$ drawn uniformly from the unit sphere. The signal component variance is $\lambda_1 = 2$, and the data sizes are $d = 10^3$, $n_1 = 10^3$ and $n_2 = 10^4$. The first noise variance is $v_1 = 1$ and the second noise variance ranges from $v_2 = 1$ to $v_2 = 20$. Figure 4 shows the nonasymptotic distribution of performance for the asymptotic optimal weights (2.5) and the existing weights (4.1)–(4.4).

We make the following observations from Figure 4:

- Across the entire sweep, the asymptotic optimal weights are generally best.
- The performance of the asymptotic optimal weights is well predicted by the asymptotic performance from Theorem 2.1.
- When v_2 is small, there is a lot of clean data coming from $\mathbf{Y}_2$ since n_2 is fairly large. All weighting schemes do well except the scheme of using only $\mathbf{Y}_1$.
- As v_2 grows, $\mathbf{Y}_2$ becomes noisier and the methods that use $\mathbf{Y}_2$ degrade in performance. Asymptotic optimal weighting degrades the most slowly / gracefully.
- As v_2 continues to grow, all methods that use $\mathbf{Y}_2$ eventually do worse than using only $\mathbf{Y}_1$ and hit zero asymptotic performance, *except for the asymptotic optimal weighting*.
- When v_2 is large, asymptotic optimal weighting performs similarly to using only $\mathbf{Y}_1$.
- Asymptotic optimal weights naturally transition from using $\mathbf{Y}_2$ to largely ignoring $\mathbf{Y}_2$ without any tuning or manual choice of a “cutoff”.

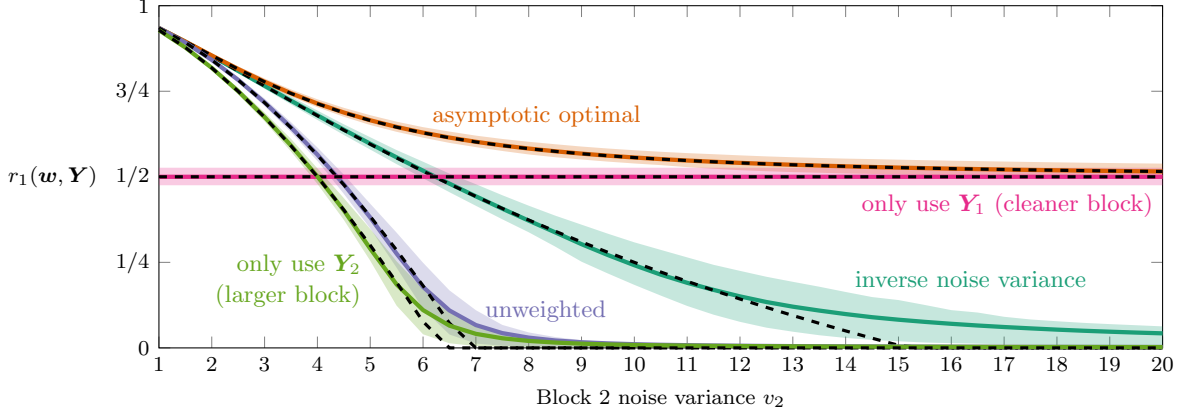


Figure 4. Performance comparison of weighting schemes (2.5) and (4.1)–(4.4) for an illustrative example with two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ with $d = 10^3$, $n_1 = 10^3$, $n_2 = 10^4$, and signal component variance $\lambda_1 = 2$. The first block has noise variance $v_1 = 1$, while the second block noise variance ranges from $v_2 = 1$ to $v_2 = 20$. For each weighting scheme, the solid colored curve is the average from 400 trials, the ribbon indicates the corresponding interquartile interval, and the dashed black curve is the asymptotic performance from Theorem 2.1 and Proposition 4.1.

- Unweighted PCA and inverse noise variance weighted PCA sometimes perform worse when given more data; in some cases using only $\mathbf{Y}_1$ or $\mathbf{Y}_2$ was better. In contrast, asymptotic optimal weighting always performs better than using only $\mathbf{Y}_1$ or only $\mathbf{Y}_2$. Moreover, it always uses all data blocks, i.e., all weights are nonzero. With optimal weighting, more data can only help, it never hurts.

Figure 4 overlaid the asymptotic performance for each weighting scheme. For optimally weighted PCA, this limit was given in (2.6). The following proposition provides an analogous result for the existing weighting schemes.

Proposition 4.1 (Asymptotic performance of the existing weighting schemes). *The weighting schemes (4.1)–(4.4) have corresponding performance converging as*

$$\text{inverse noise variance:} \quad r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} \max \left(0, \frac{c - \bar{v}^2 / \lambda_i^2}{c + \bar{v} / \lambda_i} \right), \quad (4.5)$$

$$\text{only use } \mathbf{Y}_1: \quad r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} \max \left(0, \frac{c_1 - v_1^2 / \lambda_i^2}{c_1 + v_1 / \lambda_i} \right), \quad (4.6)$$

$$\text{only use } \mathbf{Y}_2: \quad r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} \max \left(0, \frac{c_2 - v_2^2 / \lambda_i^2}{c_2 + v_2 / \lambda_i} \right), \quad (4.7)$$

$$\text{unweighted:} \quad r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} \max \left(0, \frac{A(\beta_i)}{\beta_i B'_i(\beta_i)} \right), \quad (4.8)$$

where $c := c_1 + \dots + c_L$, $\bar{v} := (p_1/v_1 + \dots + p_L/v_L)^{-1}$, $p_\ell := c_\ell/c$,

$$A(x) := 1 - \sum_{\ell=1}^L \frac{c_\ell v_\ell^2}{(x - v_\ell)^2}, \quad B_i(x) := 1 - \lambda_i \sum_{\ell=1}^L \frac{c_\ell}{x - v_\ell},$$

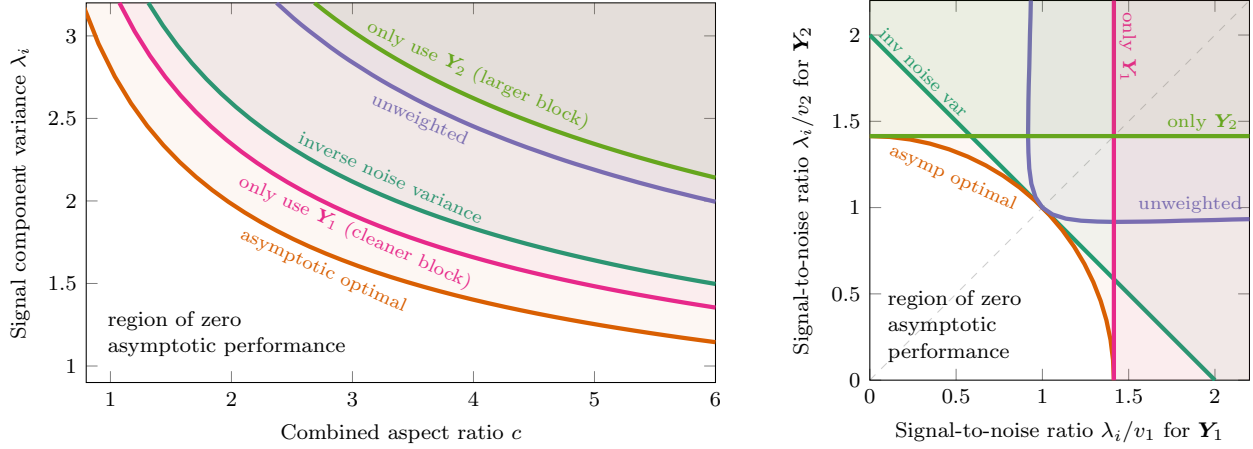
and β_i is the largest real root of the rational function B_i .

Proposition 4.1 is a by-product of Lemma 6.2 in our proof of Theorem 2.1, with some parts shown previously and some shown in this paper. Specifically, (4.6) and (4.7) are exactly the well-studied homoscedastic case [25, 26, 38, 41], since the noise is homoscedastic when using only $\mathbf{Y}_1$ or $\mathbf{Y}_2$. For unweighted PCA (4.8),

[19] derived the performance for the case where $A(\beta_i) > 0$ and conjectured the behavior for $A(\beta_i) \leq 0$. Finally, for the performance of inverse noise variance weights (4.5), closely related results were contemporaneously derived in the recent work [32].

4.3 Comparison of phase transitions

The asymptotic performances (2.6) and (4.5)–(4.8) of the various weighting schemes exhibit phase transitions between settings with zero asymptotic performance and those with nonzero asymptotic performance. Namely, each scheme has nonzero asymptotic performance for data parameters $\mathbf{c}$, $\mathbf{v}$ and λ_i in an associated regime. Figure 5 compares these regimes.



(a) With respect to combined aspect ratio c for entire dataset $\mathbf{Y}$ and signal component variance λ_i , for blocks with associated aspect ratios $\mathbf{c} = c \cdot (1/11, 10/11)$ and variances $\mathbf{v} = (1, 5)$.

(b) With respect to signal-to-noise ratios λ_i/v_1 and λ_i/v_2 , for blocks with associated aspect ratios $\mathbf{c} = (1/2, 1/2)$.

Figure 5. Comparison of phase transitions from Theorem 2.1 and Proposition 4.2 for an illustrative example of two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ with variances v_1 and v_2 and signal component variance λ_i . For each weighting scheme, asymptotic performance is zero below the phase transition and nonzero above it. In all cases, the shading goes up and to the right.

We make the following observations from Figure 5:

- None of the existing schemes dominates the rest with respect to nonzero asymptotic performance. In some cases one is better than another.
- Asymptotic optimal weighting dominates all of them. Whenever nonzero asymptotic performance is possible for one of the existing schemes, it is also possible for asymptotic optimal weighting.
- Asymptotic optimal weighting also achieves nonzero asymptotic performance in settings where all of the existing schemes have zero asymptotic performance.

For optimally weighted PCA, the condition defining the regime is given in Theorem 2.1. The following proposition gives analogous conditions for the existing weighting schemes and follows straightforwardly from Proposition 4.1.

Proposition 4.2 (Phase transitions of existing weighting schemes). *The asymptotic performance of the*

weighting schemes are nonzero if and only if, respectively,

$$\begin{aligned} \text{inverse noise variance:} & \quad c \cdot (\lambda_i/\bar{v})^2 > 1, \\ \text{only use } \mathbf{Y}_1: & \quad c_1 \cdot (\lambda_i/v_1)^2 > 1, \\ \text{only use } \mathbf{Y}_2: & \quad c_2 \cdot (\lambda_i/v_2)^2 > 1, \\ \text{unweighted:} & \quad A(\beta_i) > 0, \end{aligned}$$

where c , $\bar{v}$, A , and β_i are as in [Proposition 4.1](#).

5 Comparison with additional methods

This section compares the performance of optimally weighted PCA with additional PCA methods designed for some form of heteroscedastic noise. Specifically, we consider the following iterative methods:

- HePPCAT [\[21\]](#) is a probabilistic PCA approach that accounts for noise with samplewise heteroscedasticity by modeling the heteroscedasticity in the statistical likelihood. We used 1000 iterations.
- HeteroPCA [\[54\]](#) addresses the bias in the diagonal of the covariance matrix caused by noise with featurewise heteroscedasticity. It does so by iteratively replacing the biased entries using low-rank approximation. We used 100 iterations.

[Figure 6](#) shows the nonasymptotic distribution of performance for these methods for the setup considered in [subsection 4.2](#): $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ are generated from the model [\(2.4\)](#) with Gaussian coefficients and noise, and a single signal component $\mathbf{u}_1 \in \mathbb{R}^d$ drawn uniformly from the unit sphere. The signal component variance is $\lambda_1 = 2$, and the data sizes are $d = 10^3$, $n_1 = 10^3$ and $n_2 = 10^4$. The first noise variance is $v_1 = 1$ and the second noise variance ranges from $v_2 = 1$ to $v_2 = 20$.

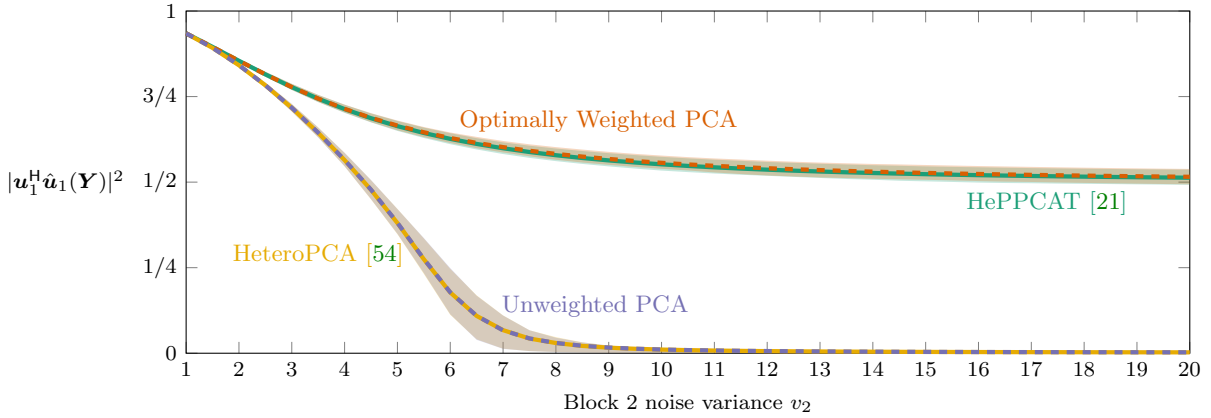


Figure 6. Performance comparison with additional methods for the illustrative example in [Figure 4](#): two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$ with $d = 10^3$, $n_1 = 10^3$, $n_2 = 10^4$, and signal component variance $\lambda_1 = 2$. The first block has noise variance $v_1 = 1$, while the second block noise variance ranges from $v_2 = 1$ to $v_2 = 20$. For each method, the colored curve is the average from 400 trials, and the ribbon indicates the corresponding interquartile interval.

We make the following observations from [Figure 6](#):

- HePPCAT accounts for the heterogeneous quality of the data blocks, and it performs very similarly to optimally weighted PCA in this case (the curves overlap). Note that HePPCAT involves solving a nonconvex optimization problem, and it currently lacks a guarantee of convergence to a global optimizer. In contrast, optimally weighted PCA can be computed simply and reliably via the well-studied singular value decomposition.

- HeteroPCA is also designed to account for heteroscedastic noise, but does so primarily for featurewise heteroscedasticity. It treats the samples uniformly and performs very similarly to unweighted PCA in this case (the curves overlap); it performs worse than optimally weighted PCA and is even eventually worse than using only $\mathbf{Y}_1$. This example highlights how samplewise and featurewise heteroscedasticity in the noise differ. They have qualitatively different impacts that seem to call for distinct approaches.

6 Proof of main result

Theorem 2.1 states that unless $\sum_{\ell=1}^L c_\ell(\lambda_i/v_\ell)^2 \leq 1$, the optimal weights $\mathbf{w}_i^*(\mathbf{Y})$ and corresponding optimal performance $r_i^*(\mathbf{Y})$ for the i th component converge almost surely to $\bar{\mathbf{w}}_i^*$ (up to scale) and $\bar{r}_i^*$ in the right hand sides of (2.5) and (2.6).

Namely, $\bar{\mathbf{w}}_i^*$ and $\bar{r}_i^*$ are the result of first optimizing (with respect to the weights $\mathbf{w}$) then taking the limit (as the data $\mathbf{Y}$ grows in size). Unfortunately, $\mathbf{w}_i^*(\mathbf{Y})$ and $r_i^*(\mathbf{Y})$ are complicated functions, making it challenging to directly analyze their limit. So, we instead first take the limit then optimize. More precisely, using aslim to denote almost sure limits and writing $\mathbf{Y}^{(d)}$ to make the limits more explicit, we first derive

$$\operatorname{argmax}_{\mathbf{w}} \bar{r}_i(\mathbf{w}) \quad \text{where} \quad \bar{r}_i(\mathbf{w}) = \text{aslim}_{d \rightarrow \infty} r_i(\mathbf{w}, \mathbf{Y}^{(d)}),$$

then use that result to obtain the result we want, i.e.,

$$\text{aslim}_{d \rightarrow \infty} \mathbf{w}_i^*(\mathbf{Y}^{(d)}) \quad \text{where} \quad \mathbf{w}_i^*(\mathbf{Y}^{(d)}) = \operatorname{argmax}_{\mathbf{w}} r_i(\mathbf{w}, \mathbf{Y}^{(d)}).$$

The following diagram illustrates the approach:

$$\begin{array}{ccc}
 r_i(\mathbf{w}, \mathbf{Y}) & \xrightarrow[\text{optimize w.r.t. } \mathbf{w}]{\substack{\text{Definition} \\ \text{(Equation (2.3))}}} & \mathbf{w}_i^*(\mathbf{Y}), r_i^*(\mathbf{Y}) \\
 \downarrow \substack{\text{Subsection 6.1} \\ \text{(Lemma 6.2)}} \text{ a.s. limit} & & \downarrow \substack{\text{a.s. limit} \\ \text{Main result} \\ \text{(Theorem 2.1)}} \\
 \bar{r}_i(\mathbf{w}) & \xrightarrow[\substack{\text{Subsection 6.2} \\ \text{(Lemma 6.3)}}]{\text{optimize w.r.t. } \mathbf{w}} & \bar{\mathbf{w}}_i^*, \bar{r}_i^*
 \end{array} \tag{6.1}$$

For this approach to work, the diagram must commute, i.e., the optimizer of the almost sure limit (which we derive) must match the almost sure limit of the optimizer (which we want). The following lemma states a suitable sufficient condition under which this happens, i.e., the maximizer of the limit is the limit of the maximizer. This lemma may be proved using techniques and results from variational analysis, e.g., [44, Chapter 7]. For convenience, [Appendix B](#) also provides an elementary, self-contained, and concise proof.

Lemma 6.1 (Diagram commutes). *Let $\mathcal{X} \subseteq \mathbb{R}^d$ be compact, $f_n : \mathcal{X} \rightarrow \mathbb{R}$ be a sequence of functions, and $f : \mathcal{X} \rightarrow \mathbb{R}$ be a function, such that on $\mathcal{X}$*

1. *each f_n has a maximum,*
2. *f_n converges uniformly to f ,*
3. *f is continuous, and*
4. *f has a unique maximizer.*

Then the maximum of f_n and the set of maximizers of f_n both converge, i.e.,

$$\max_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) \rightarrow \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}), \quad \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) \rightarrow \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}),$$

where the set convergence is with respect to the Hausdorff distance.

With [Lemma 6.1](#) in hand, it now remains to: a) derive the almost sure limit $\bar{r}_i(\mathbf{w})$ of $r_i(\mathbf{w}, \mathbf{Y})$, and show that the convergence is uniform in $\mathbf{w}$ ([Lemma 6.2](#) in [subsection 6.1](#)); and b) optimize $\bar{r}_i(\mathbf{w})$ and show that the optimizer is unique (up to scaling) except when $\sum_{\ell=1}^L c_\ell (\lambda_i/v_\ell)^2 \leq 1$ ([Lemma 6.3](#) in [subsection 6.2](#)).

6.1 Almost sure limit of the performance

This section derives the almost sure limit of the performance $r_i(\mathbf{w}, \mathbf{Y})$, where the convergence is uniform in the weights $\mathbf{w}$.

Lemma 6.2 (Almost sure limit of the performance). *For $i = 1, \dots, k$,*

$$r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{r}_i(\mathbf{w}) := \max \left(0, \frac{1}{\beta_{i,\mathbf{w}}} \frac{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})} \right), \quad (6.2)$$

where the convergence is uniform with respect to $\mathbf{w}$ on $\mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}$,

$$A_{\mathbf{w}}(x) := 1 - \sum_{\ell=1}^L \frac{c_\ell w_\ell^2 v_\ell^2}{(x - w_\ell v_\ell)^2}, \quad B_{i,\mathbf{w}}(x) := 1 - \lambda_i \sum_{\ell=1}^L \frac{c_\ell w_\ell}{x - w_\ell v_\ell}, \quad (6.3)$$

and $\beta_{i,\mathbf{w}}$ is the largest real root of $B_{i,\mathbf{w}}$.

The remainder of this subsection proves [Lemma 6.2](#). After defining some notations, we derive the limit of the singular values, then derive the limit of $r_i(\mathbf{w}, \mathbf{Y})$ in two regimes (above and below the phase transition), and finally derive the above algebraic form. There are several other ways to structure these derivations; see, e.g., [\[7\]](#). The approach we take here carefully combines the general perturbation approach of [\[5\]](#) with celebrated random matrix theoretic results on local laws [\[13, 30, 51\]](#) (reviewed in [Appendix C.1](#)) to obtain uniform convergence for the singular values and vectors. The algebraic form is derived following the approach in [\[19\]](#). Throughout the proof, we postpone some detailed calculations to the appendix.

6.1.1 Notation and preliminaries

Let $\hat{\theta}_{i,\mathbf{w}}$, $\hat{\mathbf{u}}_{i,\mathbf{w}}$, and $\hat{\mathbf{q}}_{i,\mathbf{w}}$ denote, respectively, the i th singular value, i th left singular vector and i th right singular vector of the normalized and weighted data matrix

$$\tilde{\mathbf{Y}}_{\mathbf{w}} := \frac{1}{\sqrt{n}} [\sqrt{w_1} \mathbf{Y}_1, \dots, \sqrt{w_L} \mathbf{Y}_L] = \mathbf{U} \mathbf{\Theta} \tilde{\mathbf{Q}}_{\mathbf{w}}^H + \tilde{\mathbf{E}}_{\mathbf{w}}, \quad (6.4)$$

where $\mathbf{U} := [\mathbf{u}_1, \dots, \mathbf{u}_k]$, $\mathbf{\Theta} := \text{diag}(\theta_1, \dots, \theta_k)$ has diagonal entries $\theta_i := \sqrt{\lambda_i}$, and

$$\tilde{\mathbf{Q}}_{\mathbf{w}} := \frac{1}{\sqrt{n}} [\sqrt{w_1} \mathbf{Z}_1, \dots, \sqrt{w_L} \mathbf{Z}_L]^H, \quad \tilde{\mathbf{E}}_{\mathbf{w}} := \frac{1}{\sqrt{n}} [\sqrt{w_1} \mathbf{E}_1, \dots, \sqrt{w_L} \mathbf{E}_L],$$

are normalized and weighted coefficients and noise. We indicate that these are functions of the weights via the subscript $\mathbf{w}$, and omit the dependence of the singular values and vectors on the data (for brevity). We also omit the index for d ; all limits are as $d \rightarrow \infty$ unless otherwise specified.

Note that the left singular vectors of $\tilde{\mathbf{Y}}_{\mathbf{w}}$ are exactly the eigenvectors of the weighted sample covariance, so the performance can be written as $r_i(\mathbf{w}, \mathbf{Y}) = |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2$. The noise matrix $\tilde{\mathbf{E}}_{\mathbf{w}}$ satisfies the usual random matrix theoretic conditions, and is well known to have a singular value distribution that converges weakly almost surely to a nonrandom compactly supported measure $\mu_{\mathbf{w}}$ whose Stieltjes transform

$$m_{\text{MP}}(\mathbf{w}, \zeta) := \int \frac{1}{\zeta^2 - t^2} d\mu_{\mathbf{w}}(t) \quad (6.5)$$

is the unique solution to the generalized Marchenko-Pastur equation

$$\frac{1}{m_{\text{MP}}(\mathbf{w}, \zeta)} = \zeta^2 - \sum_{\ell=1}^L \frac{p_\ell w_\ell v_\ell}{1 - w_\ell v_\ell m_{\text{MP}}(\mathbf{w}, \zeta)/c}, \quad (6.6)$$

for which $\text{Im } m_{\text{MP}}(\mathbf{w}, \zeta) < 0$ for ζ^2 in the upper half complex plane [36].

Moreover, the operator norm of the noise matrix converges to the upper-edge of $\mu_{\mathbf{w}}$ (see [Appendix C.2](#) for detailed derivation), i.e.,

$$\|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}} \xrightarrow[\mathbf{w} \in \Delta_L]{\text{a.s.}} b_{\mathbf{w}} := \sup\{\text{support of } \mu_{\mathbf{w}}\}, \quad (6.7)$$

where $\Delta_L := \{\mathbf{w} \in \mathbb{R}_{\geq 0}^L : w_1 + \dots + w_L = 1\}$ is the probability simplex, and $\xrightarrow{\text{a.s.}}$ denotes almost sure uniform convergence.

6.1.2 Limits of the singular values

Using a similar argument as [5, Section 4], one can show that any singular value of $\tilde{\mathbf{Y}}_{\mathbf{w}}$ that does not tend to a limit greater than $b_{\mathbf{w}}$ tends to $b_{\mathbf{w}}$, so we focus on singular values with limits greater than $b_{\mathbf{w}}$. Following the approach of [5, Section 4], this section studies those singular values through the matrix-valued function

$$\mathbf{M}(\mathbf{w}, \zeta) := \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} \mathbf{G}(\mathbf{w}, \zeta) \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} - \begin{bmatrix} \mathbf{\Theta}^{-1} & \mathbf{\Theta}^{-1} \end{bmatrix}, \quad (6.8)$$

where $\mathbf{G}(\mathbf{w}, \zeta)$ is the resolvent (or Green's function) defined as

$$\mathbf{G}(\mathbf{w}, \zeta) := \left(\zeta \mathbf{I}_{d+n} - \begin{bmatrix} & \tilde{\mathbf{E}}_{\mathbf{w}} \\ \tilde{\mathbf{E}}_{\mathbf{w}}^{\text{H}} & \end{bmatrix} \right)^{-1}. \quad (6.9)$$

The link to singular values is made through an extension of [5, Lemma 4.1] that incorporates the weights (see [Appendix C.3](#) for detailed derivation). It states that

$$\forall_{\zeta > \|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}}} \quad \zeta \text{ is a singular value of } \tilde{\mathbf{Y}}_{\mathbf{w}} \iff \det \mathbf{M}(\mathbf{w}, \zeta) = 0, \quad (6.10)$$

so we instead study $\mathbf{M}$ in the limit. A careful application of anisotropic local laws [13, 30, 51] (see [Appendix C.4](#) for detailed calculations) yields that for any $\tau > 0$,

$$\mathbf{M}(\mathbf{w}, \zeta) \xrightarrow[\mathbf{w}, \zeta \in \Omega(\tau)]{\text{a.s.}} \bar{\mathbf{M}}(\mathbf{w}, \zeta) := \begin{bmatrix} \varphi_{1,\mathbf{w}}(\zeta) \mathbf{I}_k & \\ & \varphi_{2,\mathbf{w}}(\zeta) \mathbf{I}_k \end{bmatrix} - \begin{bmatrix} \mathbf{\Theta}^{-1} & \mathbf{\Theta}^{-1} \end{bmatrix}, \quad (6.11)$$

where $\Omega(\tau) := \{(\mathbf{w}, \zeta) \in \Delta_L \times \mathbb{C} : (\text{Re } \zeta, \text{Im } \zeta) \in [b_{\mathbf{w}} + \tau, \infty) \times [-1, 1]\}$ and

$$\varphi_{1,\mathbf{w}}(\zeta) := \zeta m_{\text{MP}}(\mathbf{w}, \zeta) = \int \frac{\zeta}{\zeta^2 - t^2} d\mu_{\mathbf{w}}(t), \quad \varphi_{2,\mathbf{w}}(\zeta) := \sum_{\ell=1}^L \frac{p_{\ell} w_{\ell}}{\zeta - w_{\ell} v_{\ell} \varphi_{1,\mathbf{w}}(\zeta)/c}. \quad (6.12)$$

Finally, we apply [5, Lemma A.1] with (6.10) and (6.11) in the same way as [5, Section 4]; straightforward calculations (see [Appendix C.5](#)) verify that $\varphi_{1,\mathbf{w}}$ and $\varphi_{2,\mathbf{w}}$ satisfy the conditions of [5, Lemma A.1]. Moreover, noting that these arguments extend to uniform convergence in $\mathbf{w} \in \Delta_L$ yields that

$$\hat{\theta}_{i,\mathbf{w}} \xrightarrow[\mathbf{w} \in \Delta_L]{\text{a.s.}} \bar{\theta}_{i,\mathbf{w}}, \quad \text{where} \quad \bar{\theta}_{i,\mathbf{w}} := \begin{cases} \rho_{i,\mathbf{w}} & \text{if } \theta_i > \tilde{\theta}_{\mathbf{w}}, \\ b_{\mathbf{w}} & \text{otherwise,} \end{cases} \quad (6.13)$$

where $D_{\mathbf{w}}(\zeta) := \varphi_{1,\mathbf{w}}(\zeta) \varphi_{2,\mathbf{w}}(\zeta)$ for $\zeta > b_{\mathbf{w}}$, $\rho_{i,\mathbf{w}} := D_{\mathbf{w}}^{-1}(1/\theta_i^2)$, and $\tilde{\theta}_{\mathbf{w}}^2 := 1/D_{\mathbf{w}}(b_{\mathbf{w}}^+)$. Here $D_{\mathbf{w}}^{-1}$ denotes the inverse function of $D_{\mathbf{w}}$, and $f(b^+) := \lim_{\zeta \rightarrow b^+} f(\zeta)$ is the limit from above.

6.1.3 Performance above the phase transition

This section derives the limit of the performance $r_i(\mathbf{w}, \mathbf{Y})$ above the phase transition. Namely, for any $\nu > 0$, we prove uniform convergence with respect to $\mathbf{w}$ over the domain $\mathcal{W}_{>}(\nu) := \{\mathbf{w} \in \Delta_L : \hat{\theta}_{i,\mathbf{w}} > b_{\mathbf{w}} + \nu\}$.

Following the approach of [5, Section 5], we study $r_i(\mathbf{w}, \mathbf{Y}) = |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2$ through the following extension of [5, Lemma 5.1] (see Appendix C.6 for detailed derivation):

$$\forall \mathbf{w} \in \Delta_L \quad \hat{\theta}_{i,\mathbf{w}} > \|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}} \implies \left\{ \begin{array}{l} 0 = \mathbf{M}(\mathbf{w}, \hat{\theta}_{i,\mathbf{w}}) \begin{bmatrix} \Theta \tilde{\mathbf{Q}}_{\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}} \\ \Theta \mathbf{U}^H \hat{\mathbf{u}}_{i,\mathbf{w}} \end{bmatrix} \text{ and} \\ 1 = \chi_1(\mathbf{w}) + \chi_2(\mathbf{w}) + 2 \operatorname{Re} \chi_3(\mathbf{w}) \end{array} \right\}, \quad (6.14a)$$

$$1 = \chi_1(\mathbf{w}) + \chi_2(\mathbf{w}) + 2 \operatorname{Re} \chi_3(\mathbf{w}) \quad (6.14b)$$

where $\mathbf{\Gamma}_{\mathbf{w}} := (\hat{\theta}_{i,\mathbf{w}}^2 \mathbf{I}_d - \tilde{\mathbf{E}}_{\mathbf{w}} \tilde{\mathbf{E}}_{\mathbf{w}}^H)^{-1}$ and

$$\chi_1(\mathbf{w}) := \sum_{j_1, j_2=1}^k \theta_{j_1} \theta_{j_2} \cdot (\tilde{\mathbf{q}}_{j_1,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}}) \cdot (\tilde{\mathbf{q}}_{j_2,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}})^* \cdot \mathbf{u}_{j_2}^H \hat{\theta}_{i,\mathbf{w}}^2 \mathbf{\Gamma}_{\mathbf{w}}^2 \mathbf{u}_{j_1}, \quad (6.15)$$

$$\chi_2(\mathbf{w}) := \sum_{j_1, j_2=1}^k \theta_{j_1} \theta_{j_2} \cdot (\mathbf{u}_{j_1}^H \hat{\mathbf{u}}_{i,\mathbf{w}}) \cdot (\mathbf{u}_{j_2}^H \hat{\mathbf{u}}_{i,\mathbf{w}})^* \cdot \tilde{\mathbf{q}}_{j_2,\mathbf{w}}^H \tilde{\mathbf{E}}_{\mathbf{w}}^H \mathbf{\Gamma}_{\mathbf{w}}^2 \tilde{\mathbf{E}}_{\mathbf{w}} \tilde{\mathbf{q}}_{j_1,\mathbf{w}},$$

$$\chi_3(\mathbf{w}) := \sum_{j_1, j_2=1}^k \theta_{j_1} \theta_{j_2} \cdot (\mathbf{u}_{j_1}^H \hat{\mathbf{u}}_{i,\mathbf{w}}) \cdot (\tilde{\mathbf{q}}_{j_2,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}})^* \cdot \mathbf{u}_{j_2}^H \hat{\theta}_{i,\mathbf{w}} \mathbf{\Gamma}_{\mathbf{w}}^2 \tilde{\mathbf{E}}_{\mathbf{w}} \tilde{\mathbf{q}}_{j_1,\mathbf{w}}.$$

It follows from (6.13) that almost surely, eventually, for all $\mathbf{w} \in \mathcal{W}_{>}(\nu)$, $\hat{\theta}_{i,\mathbf{w}} > \|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}}$ and hence the condition of (6.14) holds. It remains to study the limits of χ_1 , χ_2 , and χ_3 .

Carefully applying (6.11) and (6.13) to (6.14a) in a similar way as [5, Section 5] yields that (see Appendix C.7 for detailed calculations), for any $\nu > 0$,

$$\chi_1(\mathbf{w}) - \theta_i^2 \left[+ \frac{\varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}})}{2\rho_{i,\mathbf{w}}} - \frac{\varphi'_{1,\mathbf{w}}(\rho_{i,\mathbf{w}})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2 \frac{\varphi_{2,\mathbf{w}}(\rho_{i,\mathbf{w}})}{\varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}})} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (6.16a)$$

$$\chi_2(\mathbf{w}) - \theta_i^2 \left[- \frac{\varphi_{2,\mathbf{w}}(\rho_{i,\mathbf{w}})}{2\rho_{i,\mathbf{w}}} - \frac{\varphi'_{2,\mathbf{w}}(\rho_{i,\mathbf{w}})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2 \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (6.16b)$$

$$\chi_3(\mathbf{w}) \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0. \quad (6.16c)$$

Finally, applying (6.16) to (6.14b) yields that (see Appendix C.8 for detailed calculations), for any $\nu > 0$,

$$r_i(\mathbf{w}, \mathbf{Y}) = |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2 \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} - \frac{2\varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}})}{\theta_i^2 D'_{\mathbf{w}}(\rho_{i,\mathbf{w}})}. \quad (6.17)$$

6.1.4 Performance below the phase transition

This section bounds the limit of the performance $r_i(\mathbf{w}, \mathbf{Y})$ below the phase transition. Namely, for any $\nu > 0$, we derive a uniform bound with respect to $\mathbf{w}$ over the domain $\mathcal{W}_{\leq}(\nu) := \Delta_L \setminus \mathcal{W}_{>}(\nu) = \{\mathbf{w} \in \Delta_L : \hat{\theta}_{i,\mathbf{w}} \leq b_{\mathbf{w}} + \nu\}$.

Following the approach of [7], we study $r_i(\mathbf{w}, \mathbf{Y}) = |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2$ by first obtaining the following deterministic bound (see Appendix C.9 for detailed calculations):

$$|\mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}|^2 \leq -\nu \cdot \operatorname{Im} \left\{ \zeta_{i,\mathbf{w}}^{-1} [\tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}}) - \tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}}) [\mathbf{M}(\mathbf{w}, \zeta_{i,\mathbf{w}})]^{-1} \tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}})]_{ii} \right\} \quad (6.18)$$

where $\zeta_{i,\mathbf{w}}^2 := \hat{\theta}_{i,\mathbf{w}}^2 + \nu$ and

$$\widetilde{\mathbf{M}}(\mathbf{w}, \zeta) := \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\mathbf{H}} \mathbf{G}(\mathbf{w}, \zeta) \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}. \quad (6.19)$$

Next, by standard calculations (see [Appendix C.10](#)), we have that almost surely, eventually, the following bounds hold for all $\mathbf{w} \in \mathcal{W}_{\leq}(\nu)$:

$$|\zeta_{i,\mathbf{w}}^{-1}| \leq \tilde{C}_1, \quad \|\widetilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}})\|_{\text{op}} \leq \tilde{C}_2, \quad \|\mathbf{M}(\mathbf{w}, \zeta_{i,\mathbf{w}})^{-1}\|_{\text{op}} \leq \tilde{C}_3 \nu^{-1/2}, \quad (6.20)$$

where $\tilde{C}_1$, $\tilde{C}_2$, and $\tilde{C}_3$ do not depend on ν or $\mathbf{w}$.

Finally, applying (6.20) to (6.18) yields that for any $\nu > 0$, almost surely, eventually,

$$\sup_{\mathbf{w} \in \mathcal{W}_{\leq}(\nu)} r_i(\mathbf{w}, \mathbf{Y}) = \sup_{\mathbf{w} \in \mathcal{W}_{\leq}(\nu)} |\mathbf{u}_i^{\mathbf{H}} \hat{\mathbf{u}}_{i,\mathbf{w}}|^2 \leq \tilde{C}_4(\nu + \nu^{1/2}), \quad (6.21)$$

where $\tilde{C}_4 := \tilde{C}_1 \max(\tilde{C}_2, \tilde{C}_2^2 \tilde{C}_3)$ does not depend on ν .

6.1.5 Uniform convergence and algebraic form of performance

Noting that $\nu > 0$ can be arbitrarily small in (6.17) and (6.21) yields uniform convergence across $\mathbf{w} \in \Delta_L$, i.e.,

$$r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow[\mathbf{w} \in \Delta_L]{\text{a.s.}} \bar{r}_i(\mathbf{w}) := \begin{cases} -\frac{2\varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}})}{\theta_i^2 D'_{\mathbf{w}}(\rho_{i,\mathbf{w}})}, & \text{if } \theta_i > \tilde{\theta}_{\mathbf{w}}, \\ 0, & \text{otherwise.} \end{cases} \quad (6.22)$$

Since $r_i(\mathbf{w}, \mathbf{Y})$ is scale invariant, i.e., $r_i(\gamma \mathbf{w}, \mathbf{Y}) = r_i(\mathbf{w}, \mathbf{Y})$ for any $\gamma > 0$, it immediately follows that the convergence is also uniform over $\mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}$ as well.

The proof concludes by deriving the algebraic description (6.2) of $\bar{r}_i(\mathbf{w})$. Following the approach of [19, Section 5.2], we change variables to

$$\psi_{\mathbf{w}}(\zeta) := \frac{c\zeta}{\varphi_{1,\mathbf{w}}(\zeta)}, \quad (6.23)$$

and observe that, analogously to [19, Section 5.3],

$$D_{\mathbf{w}}(\zeta) = \frac{1 - B_{i,\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}{\theta_i^2}, \quad \frac{D'_{\mathbf{w}}(\zeta)}{\zeta} = -\frac{2c}{\theta_i^2} \frac{B'_{i,\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}{A_{\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}, \quad (6.24)$$

$\psi_{\mathbf{w}}(b_{\mathbf{w}}^+) = \alpha_{\mathbf{w}}$ and $\psi_{\mathbf{w}}(\rho_{i,\mathbf{w}}) = \beta_{i,\mathbf{w}}$ when $\theta_i^2 > \tilde{\theta}_{\mathbf{w}}^2$, where $\alpha_{\mathbf{w}}$ and $\beta_{i,\mathbf{w}}$ are the largest real roots of $A_{\mathbf{w}}$ and $B_{i,\mathbf{w}}$, respectively. See [Appendix C.11](#) for the detailed derivations.

Even though $\psi_{\mathbf{w}}(\rho_i)$ is defined only when $\theta_i^2 > \tilde{\theta}_{\mathbf{w}}^2$, the largest real roots $\alpha_{\mathbf{w}}$ and $\beta_{i,\mathbf{w}}$ are always defined and always larger than $\max_{\ell}(w_{\ell} v_{\ell})$. Thus

$$\theta_i^2 > \tilde{\theta}_{\mathbf{w}}^2 = \frac{\theta_i^2}{1 - B_{i,\mathbf{w}}(\psi_{\mathbf{w}}(b_{\mathbf{w}}^+))} \Leftrightarrow B_{i,\mathbf{w}}(\alpha_{\mathbf{w}}) < 0 \Leftrightarrow \alpha_{\mathbf{w}} < \beta_{i,\mathbf{w}} \Leftrightarrow A_{\mathbf{w}}(\beta_{i,\mathbf{w}}) > 0, \quad (6.25)$$

where the final two equivalences hold because $A_{\mathbf{w}}(x)$ and $B_{i,\mathbf{w}}(x)$ are both strictly increasing functions for $x > \max_{\ell}(w_{\ell} v_{\ell})$ and $A_{\mathbf{w}}(\alpha_{\mathbf{w}}) = B_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) = 0$.

Finally, using (6.24) and (6.25) to rewrite (6.22) concludes the proof of [Lemma 6.2](#).

6.2 Optimization of the almost sure limit

This section optimizes $\bar{r}_i(\mathbf{w})$ and shows that the maximizer is unique (up to scaling).

Lemma 6.3 (Optimization of the almost sure limit). *The asymptotic performance $\bar{r}_i(\mathbf{w})$ is continuous and is maximized as:*

$$\{\gamma \bar{\mathbf{w}}_i^* : \gamma > 0\} = \operatorname{argmax}_{\mathbf{w} \in \mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}} \bar{r}_i(\mathbf{w}), \quad \bar{r}_i^* = \max_{\mathbf{w} \in \mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}} \bar{r}_i(\mathbf{w}), \quad (6.26)$$

except when $\sum_{\ell=1}^L c_\ell (\lambda_i / v_\ell)^2 \leq 1$, in which case $\bar{r}_i(\mathbf{w}) = 0$ for all weights $\mathbf{w} \in \mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}$.

The remainder of this subsection proves Lemma 6.3. A major challenge for the proof is that $\bar{r}_i(\mathbf{w})$ is defined implicitly via the root $\beta_{i,\mathbf{w}}$, and setting the gradient equal to zero to obtain local maxima yields a complicated nonlinear system of equations to solve. Surprisingly, the system turns out to have a simple solution that we derive by carefully exploiting the structure of the system. Moreover, we show that the solution is globally optimal, obtaining the optimal weights and their corresponding performance.

Before deriving the gradient, note that $\bar{r}_i(\mathbf{w})$ is not always differentiable everywhere due to its truncation at zero. However, it can be rewritten as

$$\bar{r}_i(\mathbf{w}) = \max(0, \tilde{r}_i(\mathbf{w})) \quad \text{where} \quad \tilde{r}_i(\mathbf{w}) := \frac{1}{\beta_{i,\mathbf{w}}} \frac{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}, \quad (6.27)$$

so the problem reduces to maximizing the differentiable function $\tilde{r}_i(\mathbf{w})$ then checking whether it is positive. Furthermore, $\tilde{r}_i(\mathbf{w})$ achieves its maximum over the feasible region $\mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}$. To see why, note that $\tilde{r}_i(\mathbf{w})$ is scale-invariant, i.e.,

$$\forall \gamma > 0 \quad \tilde{r}_i(\gamma \mathbf{w}) = \frac{1}{\beta_{i,\gamma \mathbf{w}}} \frac{A_{\gamma \mathbf{w}}(\beta_{i,\gamma \mathbf{w}})}{B'_{i,\gamma \mathbf{w}}(\beta_{i,\gamma \mathbf{w}})} = \frac{1}{\gamma \beta_{i,\mathbf{w}}} \frac{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{(1/\gamma) B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})} = \tilde{r}_i(\mathbf{w})$$

since $A_{\gamma \mathbf{w}}(x) = A_{\mathbf{w}}(x/\gamma)$, $B_{i,\gamma \mathbf{w}}(x) = B_{i,\mathbf{w}}(x/\gamma)$, and $B'_{i,\gamma \mathbf{w}}(x) = (1/\gamma) B'_{i,\mathbf{w}}(x/\gamma)$, resulting additionally in $\beta_{i,\gamma \mathbf{w}} = \gamma \beta_{i,\mathbf{w}}$. Thus, $\tilde{r}_i(\mathbf{w})$ can equivalently be maximized over the compact set $\Delta_L := \{\mathbf{w} \in \mathbb{R}_{\geq 0}^L : w_1 + \dots + w_L = 1\}$ and hence achieves its maximum.

Next, note that the feasible region $\mathbb{R}_{\geq 0}^L \setminus \{\mathbf{0}_L\}$ is not open, so we partition it into $2^L - 1$ sets according to which weights are zero. Namely, consider partitions of the form

$$\mathcal{P}_{\mathcal{L}} := \{\mathbf{w} \in \mathbb{R}_{\geq 0}^L : \forall \ell \in \mathcal{L} \ w_\ell = 0, \forall \ell \notin \mathcal{L} \ w_\ell > 0\}, \quad \text{for } \mathcal{L} \subset \{1, \dots, L\} \text{ a proper subset.}$$

Since $\tilde{r}_i(\mathbf{w})$ achieves its maximum, a maximizer exists within at least one of the partitions. Moreover, since $\tilde{r}_i(\mathbf{w})$ is differentiable, $\tilde{r}_i(\mathbf{w})$ is maximized at a critical point of a partition. It remains to identify and compare the critical points of all the partitions $\mathcal{P}_{\mathcal{L}}$.

First consider $\mathcal{P}_{\emptyset}$, i.e., the set of positive weights $w_1, \dots, w_L > 0$, and let $\tilde{w}_j := 1/w_j$. This parameterization ends up simplifying the manipulations. Differentiating $\tilde{r}_i(\mathbf{w})$ and (6.3) with respect to $\tilde{w}_j$ yields

$$\frac{\partial \tilde{r}_i(\mathbf{w})}{\partial \tilde{w}_j} = \tilde{r}_i(\mathbf{w}) \left[-\frac{1}{\beta_{i,\mathbf{w}}} \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} + \frac{1}{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})} \frac{\partial A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} - \frac{1}{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})} \frac{\partial B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} \right], \quad (6.28a)$$

$$\frac{\partial A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} = A'_{\mathbf{w}}(\beta_{i,\mathbf{w}}) \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} + 2 \frac{c_j v_j^2}{(\beta_{i,\mathbf{w}} \tilde{w}_j - v_j)^3} \beta_{i,\mathbf{w}}, \quad (6.28b)$$

$$\frac{\partial B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} = B''_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} - 2 \lambda_i \frac{c_j \tilde{w}_j}{(\beta_{i,\mathbf{w}} \tilde{w}_j - v_j)^3} \beta_{i,\mathbf{w}} + \lambda_i \frac{c_j}{(\beta_{i,\mathbf{w}} \tilde{w}_j - v_j)^2}, \quad (6.28c)$$

$$0 = \frac{\partial B_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} = B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} + \lambda_i \frac{c_j}{(\beta_{i,\mathbf{w}} \tilde{w}_j - v_j)^2} \beta_{i,\mathbf{w}}, \quad (6.28d)$$

where one must carefully account for the fact that $A_{\mathbf{w}}(x)$, $B_{i,\mathbf{w}}(x)$ and $\beta_{i,\mathbf{w}}$ are functions of $\tilde{w}_j$. The problem now is to set $\partial \tilde{r}_i(\mathbf{w})/\partial \tilde{w}_j$ equal to zero and carefully exploit the structure of the system (6.28) to solve it.

In particular, rewriting (6.28b) and (6.28c) in terms of $\partial \beta_{i,\mathbf{w}}/\partial \tilde{w}_j$ using (6.28d) yields

$$\frac{\partial A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} = \left[A'_{\mathbf{w}}(\beta_{i,\mathbf{w}}) - \frac{2B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\lambda_i} \frac{v_j^2}{\beta_{i,\mathbf{w}}\tilde{w}_j - v_j} \right] \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j}, \quad (6.29a)$$

$$\frac{\partial B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\partial \tilde{w}_j} = \left[B''_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) + 2B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) \frac{\tilde{w}_j}{\beta_{i,\mathbf{w}}\tilde{w}_j - v_j} \right] \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} - \frac{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{\beta_{i,\mathbf{w}}} \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j}. \quad (6.29b)$$

Substituting (6.29a) and (6.29b) into (6.28a) then rearranging yields

$$\frac{\partial \tilde{r}_i(\mathbf{w})}{\partial \tilde{w}_j} = \frac{2}{\lambda_i \beta_{i,\mathbf{w}}} \frac{\partial \beta_{i,\mathbf{w}}}{\partial \tilde{w}_j} \left[\lambda_i \Delta_{i,\mathbf{w}} - \frac{\lambda_i \beta_{i,\mathbf{w}} \tilde{r}_i(\mathbf{w}) \tilde{w}_j + v_j^2}{\beta_{i,\mathbf{w}} \tilde{w}_j - v_j} \right], \quad (6.30)$$

where the following term is independent of j :

$$\Delta_{i,\mathbf{w}} := \frac{1}{2} \frac{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})} \left[\frac{A'_{\mathbf{w}}(\beta_{i,\mathbf{w}})}{A_{\mathbf{w}}(\beta_{i,\mathbf{w}})} - \frac{B''_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})}{B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})} \right].$$

Since $\beta_{i,\mathbf{w}} > \max_{\ell}(w_{\ell}v_{\ell}) > 0$, it follows from (6.28d) that $\partial \beta_{i,\mathbf{w}}/\partial \tilde{w}_j \neq 0$, so it follows from (6.30) that $\partial \tilde{r}_i(\mathbf{w})/\partial \tilde{w}_j$ is zero exactly when

$$\lambda_i \Delta_{i,\mathbf{w}} = \frac{\lambda_i \beta_{i,\mathbf{w}} \tilde{r}_i(\mathbf{w}) \tilde{w}_j + v_j^2}{\beta_{i,\mathbf{w}} \tilde{w}_j - v_j}. \quad (6.31)$$

Rearranging (6.27) and substituting (6.31) yields

$$\begin{aligned} 0 &= A_{\mathbf{w}}(\beta_{i,\mathbf{w}}) - \tilde{r}_i(\mathbf{w}) \beta_{i,\mathbf{w}} B'_{i,\mathbf{w}}(\beta_{i,\mathbf{w}}) = 1 - \sum_{\ell=1}^L \frac{c_{\ell}(v_{\ell}^2 + \lambda_i \beta_{i,\mathbf{w}} \tilde{r}_i(\mathbf{w}) \tilde{w}_{\ell})}{(\beta_{i,\mathbf{w}} \tilde{w}_{\ell} - v_{\ell})^2} \\ &= 1 - \Delta_{i,\mathbf{w}} \lambda_i \sum_{\ell=1}^L \frac{c_{\ell}}{\beta_{i,\mathbf{w}} \tilde{w}_{\ell} - v_{\ell}} = 1 - \Delta_{i,\mathbf{w}} (1 - B_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})) = 1 - \Delta_{i,\mathbf{w}}, \end{aligned}$$

so $\Delta_{i,\mathbf{w}} = 1$. Substituting into (6.31) and solving for $\tilde{w}_j$ yields

$$w_j = \frac{1}{\tilde{w}_j} = \frac{(1 - \tilde{r}_i(\mathbf{w})) \beta_{i,\mathbf{w}}}{v_j(1 + v_j/\lambda_i)} = \frac{\gamma_{i,\mathbf{w}}}{v_j(1 + v_j/\lambda_i)}, \quad (6.32)$$

where the constant $\gamma_{i,\mathbf{w}} := (1 - \tilde{r}_i(\mathbf{w})) \beta_{i,\mathbf{w}}$ is: a) independent of j , b) parameterizes the ray of critical points in $\mathcal{P}_{\emptyset}$, and c) can be chosen freely, e.g., as unity yielding $\tilde{\mathbf{w}}_i^*$ from (2.5).

Solving (6.32) for $\beta_{i,\mathbf{w}}$, substituting into $0 = B_{i,\mathbf{w}}(\beta_{i,\mathbf{w}})$, and rearranging yields that the corresponding $\tilde{r}_i(\tilde{\mathbf{w}}_i^*)$ is a root of

$$R_{i,\emptyset}(x) := 1 - \sum_{\ell=1}^L \frac{c_{\ell}}{v_{\ell}/\lambda_i} \frac{1-x}{v_{\ell}/\lambda_i + x}.$$

Since $R_{i,\emptyset}(x)$ increases from negative infinity to one as x increases from $-\min_{\ell}(v_{\ell})/\lambda_i$ to one, it has exactly one real root in that domain. In particular, this root is the largest real root since $R_{i,\emptyset}(x) \geq 1$ for $x \geq 1$. Furthermore, $\tilde{r}_i(\tilde{\mathbf{w}}_i^*)$ increases continuously to one as $c := c_1 + \dots + c_L$ increases to infinity, so $\tilde{r}_i(\tilde{\mathbf{w}}_i^*)$ must be the largest real root.

Likewise, the critical points of other partitions $\mathcal{P}_{\mathcal{L}}$ are given by setting the positive weights proportional to $\tilde{\mathbf{w}}_i^*$ with the corresponding $\tilde{r}_i(\tilde{\mathbf{w}}_i^*)$ given by the largest real root of

$$R_{i,\mathcal{L}}(x) := 1 - \sum_{\ell \notin \mathcal{L}} \frac{c_{\ell}}{v_{\ell}/\lambda_i} \frac{1-x}{v_{\ell}/\lambda_i + x}.$$

For $\mathcal{L}_1 \subset \mathcal{L}_2$ a proper subset, the largest real root of $R_{i,\mathcal{L}_1}$ is greater than that of $R_{i,\mathcal{L}_2}$ since $R_{i,\mathcal{L}_1}(x) < R_{i,\mathcal{L}_2}(x)$ for any $x \in (-\min_\ell(v_\ell)/\lambda_i, 1)$. As a result, $\tilde{r}_i(\mathbf{w})$ is maximized in $\mathcal{P}_\emptyset$.

Finally, we check when $\tilde{r}_i(\bar{\mathbf{w}}_i^*)$ is positive. Recalling that $R_{i,\emptyset}(x)$ is an increasing function on $x \in (0, 1)$ and noting that $R_{i,\emptyset}(0) = 1 - \sum_{\ell=1}^L c_\ell(\lambda_i/v_\ell)^2$ yields that $\tilde{r}_i(\bar{\mathbf{w}}_i^*)$ is positive if and only if $\sum_{\ell=1}^L c_\ell(\lambda_i/v_\ell)^2 > 1$. When it is positive, the maximizers are given by the critical points above; otherwise $\tilde{r}_i(\mathbf{w}) = 0$ for all $\mathbf{w}$. This concludes the proof of [Lemma 6.3](#).

7 When signal and noise variances are unknown

The asymptotic optimal weights from the main result ([Theorem 2.1](#)) depend on both the signal component variance λ_i and the noise variances $\mathbf{v}$, but one or both of these variances may be unknown in some settings. Of course, one could estimate these variances using existing ideas and approaches, then plug them into (2.5) in [Theorem 2.1](#). The question then is, are these resulting estimated weights also asymptotically optimal? Fortunately, it is straightforward to see that the answer is yes under natural conditions on the variance estimates. The following proposition makes this statement precise; it follows straightforwardly from [Lemmas 6.2](#) and [6.3](#).

Proposition 7.1 (Asymptotic optimality with estimated variances). *Suppose $\hat{\lambda}_i(\mathbf{Y})$ and $\hat{\mathbf{v}}(\mathbf{Y})$ are consistent estimates of λ_i and $\mathbf{v}$, i.e.,*

$$\hat{\lambda}_i(\mathbf{Y}) \xrightarrow{\text{a.s.}} \lambda_i, \quad \hat{\mathbf{v}}(\mathbf{Y}) \xrightarrow{\text{a.s.}} \mathbf{v}. \quad (7.1)$$

Let $\hat{\mathbf{w}}_i^*(\mathbf{Y})$ be the estimated weights obtained by plugging $\hat{\mathbf{v}}(\mathbf{Y})$ and $\hat{\lambda}_i(\mathbf{Y})$ into (2.5) in [Theorem 2.1](#), i.e.,

$$\hat{\mathbf{w}}_i^*(\mathbf{Y}) := \left(\frac{1}{\hat{v}_1(\mathbf{Y})} \frac{1}{1 + \hat{v}_1(\mathbf{Y})/\hat{\lambda}_i(\mathbf{Y})}, \dots, \frac{1}{\hat{v}_L(\mathbf{Y})} \frac{1}{1 + \hat{v}_L(\mathbf{Y})/\hat{\lambda}_i(\mathbf{Y})} \right). \quad (7.2)$$

Then the estimated weights and their corresponding performance converge to the asymptotic optimal weights and performance, i.e.,

$$\hat{\mathbf{w}}_i^*(\mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{\mathbf{w}}_i^*, \quad r_i(\hat{\mathbf{w}}_i^*(\mathbf{Y}), \mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{r}_i^*. \quad (7.3)$$

Namely, the estimated weights are asymptotically optimal.

Thus, optimal weighting only needs consistent estimates of λ_i and $\mathbf{v}$ that may be obtained using any one of various existing approaches; which one is most appropriate will depend on the specific application. Here we consider a simple pair of estimators as an illustrative example.

Example 7.2 (Variance estimators). As an illustrative example, consider the following simple estimators for the signal and noise variances

$$\hat{\lambda}_i(\mathbf{Y}) := \Xi \left(\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \hat{\mathbf{v}}(\mathbf{Y})); \left[\sum_{\ell=1}^L \frac{p_\ell}{\hat{v}_\ell(\mathbf{Y})} \right]^{-1} \right), \quad \hat{\mathbf{v}}(\mathbf{Y}) := \left(\frac{\|\mathbf{Y}_1\|_{\text{F}}^2}{dn_1}, \dots, \frac{\|\mathbf{Y}_L\|_{\text{F}}^2}{dn_L} \right), \quad (7.4)$$

where

$$\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \mathbf{v}) := \text{ith leading eigenvalue of } \sum_{\ell=1}^L \left(\frac{1/v_\ell}{n_1/v_1 + \dots + n_L/v_L} \right) \mathbf{Y}_\ell \mathbf{Y}_\ell^H, \quad (7.5)$$

$$\Xi(\lambda; v) := \text{the larger root of the quadratic polynomial } (x + v/c)(x + v) - \lambda x. \quad (7.6)$$

It is straightforward to verify with standard techniques (see [Appendix D](#) for details) that these estimators are consistent as long as $c > (\bar{v}/\lambda_i)^2$, i.e., when inverse noise variance weighting is above the phase transition.

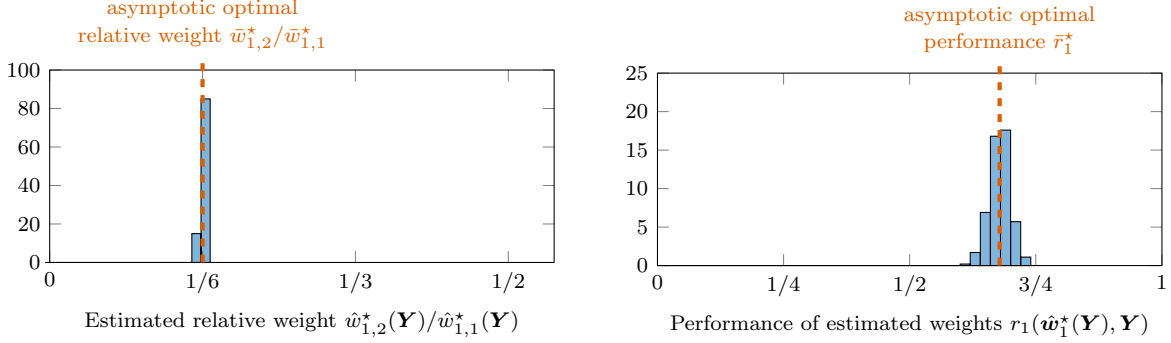


Figure 7. Nonasymptotic empirical distributions of estimated weights $\hat{\mathbf{w}}_1^*(\mathbf{Y})$ from (7.2) and corresponding performance $r_1(\hat{\mathbf{w}}_1^*(\mathbf{Y}), \mathbf{Y})$ for an illustrative example with two blocks of data $\mathbf{Y}_1 \in \mathbb{R}^{d \times n_1}$ and $\mathbf{Y}_2 \in \mathbb{R}^{d \times n_2}$, where the estimated weights are computed using the signal and noise variance estimators (7.4) from Example 7.2. The data blocks are generated with noise variances $v_1 = 1$ and $v_2 = 3$, one underlying component having variance $\lambda_1 = 1$, and dimensions $d = 1000$, $n_1 = 4000$ and $n_2 = 8000$.

Thus, by Proposition 7.1, the resulting estimated weights and their corresponding performance converge to the asymptotic optimal weights and performance.

Figure 7 illustrates the nonasymptotic behavior of these estimated weights in numerical simulations. The data is generated as in section 3, with dimensionality $d = 1000$ and block sizes $\mathbf{n} = (4000, 8000)$. The estimated weights and their performance generally concentrate around the asymptotic optimal weights and performance. Moreover, the estimated weights achieve performance closely matching the nonasymptotic empirically optimized weights in Figure 2b.

8 Illustration on real data from astronomy

This section illustrates optimally weighted PCA on real data from astronomy. In particular, we consider quasar spectra from the Sloan Digital Sky Survey (SDSS) Data Release 16 [2] using the associated DR16Q quasar catalog [35]. Each spectrum is a vector of flux measurements across wavelengths for a particular quasar, and comes with associated noise variance estimates across wavelengths. The noise is heteroscedastic across both quasars and wavelengths, but here we focus on a subset that is somewhat homoscedastic across wavelengths. Appendix E describes the details of the subset selected and the preprocessing performed (filtering, interpolation, centering, normalization).

The resulting data has $d = 281$ wavelengths measured for $n^{(\text{full})} = 10459$ spectra, yielding a data matrix $\mathbf{Y}^{(\text{full})} \in \mathbb{R}^{d \times n^{(\text{full})}}$ with a vector of associated variances $\mathbf{v}^{(\text{full})} \in \mathbb{R}_{\geq 0}^{n^{(\text{full})}}$. Figure 8 shows plots of components $\mathbf{u}_1, \dots, \mathbf{u}_5 \in \mathbb{R}^d$ computed via an unweighted PCA on the 5000 samples from $\mathbf{Y}^{(\text{full})}$ with smallest variances. We regard these components as “ground truth.”

To evaluate the various weighting schemes, we formed a test dataset $\mathbf{Y} \in \mathbb{R}^{d \times n}$ containing $n = 5000$ samples by combining the 3000 samples with the smallest variances (a subset of the 5000 samples used to produce ground truth) and the 2000 samples with the largest variances. This provides a dataset with noise heteroscedasticity across samples, shown as a heatmap in Figure 9a. Figure 9b shows a few sample spectra from the dataset; note that they have a common shape but have varying levels of noise.

We computed the leading singular vectors $\hat{\mathbf{u}}_1(\mathbf{w}_1, \mathbf{Y}), \dots, \hat{\mathbf{u}}_5(\mathbf{w}_5, \mathbf{Y})$ via unweighted PCA, inverse variance weighted PCA, and optimally weighted PCA; the optimally weighted PCA has component-specific weights, so its weights are not the same across the components.² We used the provided noise variances,

²While the components $\hat{\mathbf{u}}_1(\hat{\mathbf{w}}_1^*(\mathbf{Y}), \mathbf{Y}), \dots, \hat{\mathbf{u}}_5(\hat{\mathbf{w}}_1^*(\mathbf{Y}), \mathbf{Y})$ from optimally weighted PCA are not guaranteed to be orthogonal in general (due to the component-specific weighting), they were approximately orthogonal here. In particular, $\max_{i \neq j} |\hat{\mathbf{u}}_i(\hat{\mathbf{w}}_i^*(\mathbf{Y}), \mathbf{Y})^H \hat{\mathbf{u}}_j(\hat{\mathbf{w}}_j^*(\mathbf{Y}), \mathbf{Y})|^2 \approx 0.00013$.

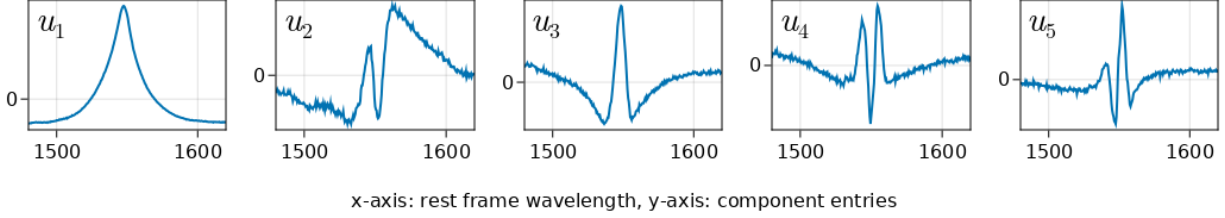


Figure 8. Ground truth components computed from 5000 samples with smallest variances.

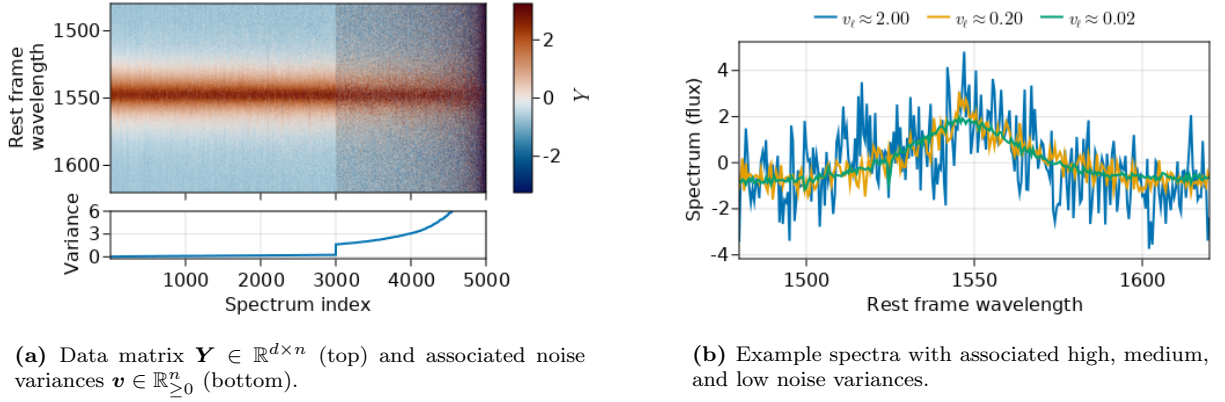


Figure 9. Quasar spectra dataset and example spectra.

and estimated the signal variances using the estimator from [Example 7.2](#) (with the provided noise variances substituted). [Table 1](#) shows the recovery of the ground truth singular vectors calculated from the “clean” samples above.

Table 1. Recoveries $r_i(\mathbf{w}, \mathbf{Y})$ for unweighted PCA ($w_\ell = 1$), inverse variance weighted PCA ($w_\ell = 1/v_\ell$), and optimally weighted PCA. Higher is better, best value (up to rounding) is shown in bold.

Component	1	2	3	4	5
Unweighted PCA	0.003	0.307	0.009	0.004	0.018
Inverse variance weighted PCA	1.000	0.903	0.915	0.817	0.811
Optimally weighted PCA	1.000	0.920	0.934	0.884	0.880

The following observations apply to this data and also summarize some of the main themes of this paper:

- unweighted PCA performs poorly for heteroscedastic data,
- inverse variance weighted PCA performs much better than unweighted PCA, and
- optimally weighted PCA performs even better than inverse variance weighted PCA.

Similar comparisons for these leading components occurred for many of the other test datasets we tried. The comparisons were less consistent for components 6 and on; optimal weights were sometimes better and inverse noise variance weights were sometimes better. This inconsistent behavior is potentially due to the data not matching the model closely enough or perhaps the ground truth needing to be chosen differently. A detailed investigation of this phenomenon is beyond our present scope.

Overall, this example with quasar spectra coming from astronomy illustrates the potential for optimally weighted PCA to improve the recovery of underlying principal components from real data that has heteroscedastic noise.

9 Conclusion

This paper derived asymptotic optimal weights for weighted PCA when the data is high-dimensional with noise that is heteroscedastic across samples. The optimal weights are a simple function of the signal and noise variances, and are not the inverse noise variance weights commonly used in practice. Numerical simulations illustrated that the asymptotic optimal weights are often close to the optimal weights in finite dimensions when the dimensions are large enough. Comparisons of the asymptotic optimal weights with existing weighting schemes illustrated that the asymptotic optimal weights: a) are generally closer to the optimal weights in finite dimensions, b) appropriately combine all the data to achieve the best performance, and c) achieve nonzero asymptotic performance in the widest range of settings. Additional simulations illustrated that optimally weighted PCA compares favorably with other PCA methods designed for some form of heteroscedastic noise. Finally, we briefly discussed how one can use estimated signal and noise variances when the true variances are unknown, and illustrated the optimal weights on real data from astronomy.

Overall, optimally weighted PCA is a simple, principled, and promising method for estimating underlying principal components from high-dimensional data with noise that is heteroscedastic across samples. However, many open questions remain. While the asymptotic optimal performance (2.6) is often close to the performance of optimally weighted PCA in finite dimensions, it would be useful to also derive higher-order asymptotics for the performance as well as nonasymptotic characterizations. These results would provide refined estimates of the performance. Another interesting variation of the asymptotic regime is to allow the number of blocks L to grow with d , potentially with $n_1, \dots, n_L = O(1)$. This regime may better capture datasets where each sample has its own associated noise variance (e.g., like the astronomy data in section 8) or where the block structure is unknown. While the proof of Theorem 2.1 does not seem to readily generalize to such cases, we conjecture that the optimal weights will still have the same form and that estimates of the signal and noise variances can still be used when the true variances are unknown. Another interesting direction is to study whether optimally weighted PCA is optimal across not just weighted PCA but also across more general classes of methods, e.g., by deriving fundamental bounds on the achievable performance. Finally, it would be interesting to study how optimally weighted PCA might be combined with other techniques to handle settings where the noise is not just heteroscedastic across samples but also across features.

Acknowledgements

The authors thank Raj Rao Nadakuditi, Romain Couillet and Edgar Dobriban for helpful discussions regarding the singular values and vectors of low rank perturbations of large random matrices. The authors also thank Dejiao Zhang for helpful comments about the form of the optimal weights.

The example quasar spectra were provided by the Sloan Digital Sky Survey [2, 35]. Funding for the Sloan Digital Sky Survey IV has been provided by the Alfred P. Sloan Foundation, the U.S. Department of Energy Office of Science, and the Participating Institutions.

SDSS-IV acknowledges support and resources from the Center for High Performance Computing at the University of Utah. The SDSS website is www.sdss.org.

SDSS-IV is managed by the Astrophysical Research Consortium for the Participating Institutions of the SDSS Collaboration including the Brazilian Participation Group, the Carnegie Institution for Science, Carnegie Mellon University, Center for Astrophysics — Harvard & Smithsonian, the Chilean Participation Group, the French Participation Group, Instituto de Astrofísica de Canarias, The Johns Hopkins University, Kavli Institute for the Physics and Mathematics of the Universe (IPMU) / University of Tokyo, the Korean Participation Group, Lawrence Berkeley National Laboratory, Leibniz Institut für Astrophysik Potsdam

(AIP), Max-Planck-Institut für Astronomie (MPIA Heidelberg), Max-Planck-Institut für Astrophysik (MPA Garching), Max-Planck-Institut für Extraterrestrische Physik (MPE), National Astronomical Observatories of China, New Mexico State University, New York University, University of Notre Dame, Observatório Nacional / MCTI, The Ohio State University, Pennsylvania State University, Shanghai Astronomical Observatory, United Kingdom Participation Group, Universidad Nacional Autónoma de México, University of Arizona, University of Colorado Boulder, University of Oxford, University of Portsmouth, University of Utah, University of Virginia, University of Washington, University of Wisconsin, Vanderbilt University, and Yale University.

References

- [1] R. B. ABDALLAH, A. BRELOY, M. N. E. KORSO, AND D. LAUTRU, *Bayesian signal subspace estimation with compound gaussian sources*, Signal Processing, 167 (2020), pp. 1–15, <https://doi.org/10.1016/j.sigpro.2019.107310>.
- [2] R. AHUMADA, C. A. PRIETO, A. ALMEIDA, F. ANDERS, ET AL., *The 16th data release of the Sloan Digital Sky Surveys: First release from the APOGEE-2 southern survey and full release of eBOSS spectra*, The Astrophysical Journal Supplement Series, 249 (2020), pp. 1–21, <https://doi.org/10.3847/1538-4365/ab929e>.
- [3] B. A. ARDEKANI, J. KERSHAW, K. KASHIKURA, AND I. KANNO, *Activation detection in functional MRI using subspace modeling and maximum likelihood estimation*, IEEE Transactions on Medical Imaging, 18 (1999), pp. 101–114, <https://doi.org/10.1109/42.759109>.
- [4] S. BAILEY, *Principal component analysis with noisy and/or missing data*, Publications of the Astronomical Society of the Pacific, 124 (2012), pp. 1015–1023, <https://doi.org/10.1086/668105>.
- [5] F. BENAYCH-GEORGES AND R. R. NADAKUDITI, *The singular values and vectors of low rank perturbations of large rectangular random matrices*, Journal of Multivariate Analysis, 111 (2012), pp. 120–135, <https://doi.org/10.1016/j.jmva.2012.04.019>.
- [6] O. BESSON, *Bounds for a mixture of low-rank compound-gaussian and white gaussian noises*, IEEE Transactions on Signal Processing, 64 (2016), pp. 5723–5732, <https://doi.org/10.1109/tsp.2016.2603965>.
- [7] A. BLOEMENDAL, L. ERDŐS, A. KNOWLES, H.-T. YAU, AND J. YIN, *Isotropic local laws for sample covariance and generalized Wigner matrices*, Electronic Journal of Probability, 19 (2014), pp. 1–53, <https://doi.org/10.1214/ejp.v19-3054>.
- [8] A. BRELOY, G. GINOLHAC, F. PASCAL, AND P. FORSTER, *Clutter subspace estimation in low rank heterogeneous noise context*, IEEE Transactions on Signal Processing, 63 (2015), pp. 2173–2182, <https://doi.org/10.1109/tsp.2015.2403284>.
- [9] A. BRELOY, G. GINOLHAC, F. PASCAL, AND P. FORSTER, *Robust covariance matrix estimation in heterogeneous low rank context*, IEEE Transactions on Signal Processing, 64 (2016), pp. 5794–5806, <https://doi.org/10.1109/tsp.2016.2599494>.
- [10] R. N. COCHRAN AND F. H. HORNE, *Statistically weighted principal component analysis of rapid scanning wavelength kinetics experiments*, Analytical Chemistry, 49 (1977), pp. 846–853, <https://doi.org/10.1021/ac50014a045>.
- [11] A. COLLAS, F. BOUCHARD, A. BRELOY, G. GINOLHAC, C. REN, AND J.-P. OVARLEZ, *Probabilistic PCA from heteroscedastic signals: Geometric framework and application to clustering*, IEEE Transactions on Signal Processing, 69 (2021), pp. 6546–6560, <https://doi.org/10.1109/tsp.2021.3130997>.

- [12] J.-C. DEVILLE AND E. MALINVAUD, *Data analysis in official socio-economic statistics*, Journal of the Royal Statistical Society. Series A (General), 146 (1983), pp. 335–361, <https://doi.org/10.2307/2981452>.
- [13] X. DING AND F. YANG, *A necessary and sufficient condition for edge universality at the largest singular values of covariance matrices*, The Annals of Applied Probability, 28 (2018), pp. 1679–1738, <https://doi.org/10.1214/17-aap1341>.
- [14] X. DING AND F. YANG, *Spiked separable covariance matrices and principal components*, The Annals of Statistics, 49 (2021), pp. 1113–1138, <https://doi.org/10.1214/20-aos1995>.
- [15] E. DOBRIBAN, W. LEEB, AND A. SINGER, *Optimal prediction in the linearly transformed spiked model*, The Annals of Statistics, 48 (2020), pp. 491–513, <https://doi.org/10.1214/19-aos1819>.
- [16] D. DONOHO, M. GAVISH, AND I. JOHNSTONE, *Optimal shrinkage of eigenvalues in the spiked covariance model*, The Annals of Statistics, 46 (2018), <https://doi.org/10.1214/17-aos1601>.
- [17] U. ENVIRONMENTAL PROTECTION AGENCY, *Air quality system data mart*, <https://www.epa.gov/airdata>.
- [18] D. HONG, L. BALZANO, AND J. A. FESSLER, *Towards a theoretical analysis of PCA for heteroscedastic data*, in 54th Allerton Conference on Communication, Control, and Computing, IEEE, Sept. 2016, pp. 496–503, <https://doi.org/10.1109/allerton.2016.7852272>.
- [19] D. HONG, L. BALZANO, AND J. A. FESSLER, *Asymptotic performance of PCA for high-dimensional heteroscedastic data*, Journal of Multivariate Analysis, 167 (2018), pp. 435–452, <https://doi.org/10.1016/j.jmva.2018.06.002>.
- [20] D. HONG, L. BALZANO, AND J. A. FESSLER, *Probabilistic PCA for heteroscedastic data*, in 8th Intl. Workshop on Computational Advances in Multi-Sensor Adaptive Processing, IEEE, Dec. 2019, pp. 26–30, <https://doi.org/10.1109/camsap45676.2019.9022436>.
- [21] D. HONG, K. GILMAN, L. BALZANO, AND J. A. FESSLER, *HePPCAT: Probabilistic PCA for data with heteroscedastic noise*, IEEE Transactions on Signal Processing, 69 (2021), pp. 4819–4834, <https://doi.org/10.1109/tsp.2021.3104979>.
- [22] D. HONG, Y. SHENG, AND E. DOBRIBAN, *Selecting the number of components in pca via random signflips*, 2020, <http://arxiv.org/abs/2012.02985v1>.
- [23] R. A. HORN AND C. R. JOHNSON, *Matrix Analysis*, Cambridge University Press, 2 ed., 2013.
- [24] J. J. JANSEN, H. C. J. HOEFSLOOT, H. F. M. BOELEN, J. VAN DER GREEF, AND A. K. SMILDE, *Analysis of longitudinal metabolomics data*, Bioinformatics, 20 (2004), pp. 2438–2446, <https://doi.org/10.1093/bioinformatics/bth268>.
- [25] I. M. JOHNSTONE AND A. Y. LU, *On consistency and sparsity for principal components analysis in high dimensions*, Journal of the American Statistical Association, 104 (2009), pp. 682–693, <https://doi.org/10.1198/jasa.2009.0121>.
- [26] I. M. JOHNSTONE AND D. PAUL, *PCA in high dimensions: An orientation*, Proceedings of the IEEE, 106 (2018), pp. 1277–1292, <https://doi.org/10.1109/jproc.2018.2846730>.
- [27] I. M. JOHNSTONE AND D. M. TITTERINGTON, *Statistical challenges of high-dimensional data*, Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences, 367 (2009), pp. 4237–4253, <https://doi.org/10.1098/rsta.2009.0159>.
- [28] I. T. JOLLIFFE, *Principal Component Analysis*, Springer-Verlag, New York, NY, 2 ed., 2002, <https://doi.org/10.1007/b98835>.

- [29] Z. T. KE, Y. MA, AND X. LIN, *Estimation of the number of spiked eigenvalues in a covariance matrix by bulk eigenvalue matching analysis*, Journal of the American Statistical Association, (2021), pp. 1–19, <https://doi.org/10.1080/01621459.2021.1933497>.
- [30] A. KNOWLES AND J. YIN, *Anisotropic local laws for random matrices*, Probability Theory and Related Fields, 169 (2016), pp. 257–352, <https://doi.org/10.1007/s00440-016-0730-4>.
- [31] B. LANDA, T. T. C. K. ZHANG, AND Y. KLUGER, *Biwhitening reveals the rank of a count matrix*, 2021, <http://arxiv.org/abs/2103.13840v2>.
- [32] W. LEEB AND E. ROMANOV, *Optimal spectral shrinkage and PCA with heteroscedastic noise*, IEEE Transactions on Information Theory, 67 (2021), pp. 3009–3037, <https://doi.org/10.1109/tit.2021.3055075>.
- [33] W. E. LEEB, *Matrix denoising for weighted loss functions and heterogeneous signals*, SIAM Journal on Mathematics of Data Science, 3 (2021), pp. 987–1012, <https://doi.org/10.1137/20m1319577>.
- [34] J. T. LEEK, *Asymptotic conditional singular value decomposition for high-dimensional genomic data*, Biometrics, 67 (2010), pp. 344–352, <https://doi.org/10.1111/j.1541-0420.2010.01455.x>.
- [35] B. W. LYKE, A. N. HIGLEY, J. N. MCLANE, D. P. SCHURHAMMER, ET AL., *The Sloan Digital Sky Survey quasar catalog: Sixteenth data release*, The Astrophysical Journal Supplement Series, 250 (2020), pp. 1–24, <https://doi.org/10.3847/1538-4365/aba623>.
- [36] V. A. MARČENKO AND L. A. PASTUR, *Distribution of eigenvalues for some sets of random matrices*, Mathematics of the USSR-Sbornik, 1 (1967), pp. 457–483, <https://doi.org/10.1070/sm1967v001n04abeh001994>.
- [37] R. R. NADAKUDITI, *OptShrink: An algorithm for improved low-rank signal matrix denoising by optimal, data-driven singular value shrinkage*, IEEE Transactions on Information Theory, 60 (2014), pp. 3002–3018, <https://doi.org/10.1109/tit.2014.2311661>.
- [38] B. NADLER, *Finite sample approximation results for principal component analysis: A matrix perturbation approach*, The Annals of Statistics, 36 (2008), pp. 2791–2817, <https://doi.org/10.1214/08-aos618>.
- [39] S. PAPADIMITRIOU, J. SUN, AND C. FALOUTSOS, *Streaming pattern discovery in multiple time-series*, in Proc. 31st Intl. Conf. on Very Large Data Bases, ACM, 2005, pp. 697–708, <http://www.vldb.org/archives/website/2005/program/paper/thu/p697-papadimitriou.pdf>.
- [40] D. PASSEMIER, Z. LI, AND J. YAO, *On estimation of the noise variance in high dimensional probabilistic principal component analysis*, Journal of the Royal Statistical Society: Series B (Statistical Methodology), 79 (2015), pp. 51–67, <https://doi.org/10.1111/rssb.12153>.
- [41] D. PAUL, *Asymptotics of sample eigenstructure for a large dimensional spiked covariance model*, Statistica Sinica, 17 (2007), pp. 1617–1642, <http://www3.stat.sinica.edu.tw/statistica/J17N4/J17N418/J17N418.html>.
- [42] H. PEDERSEN, S. KOZERKE, S. RINGGAARD, K. NEHRKE, AND W. Y. KIM, *k-t PCA: Temporally constrained k-t BLAST reconstruction using principal component analysis*, Magnetic Resonance in Medicine, 62 (2009), pp. 706–716, <https://doi.org/10.1002/mrm.22052>.
- [43] PURPLEAIR, *Real time air quality monitoring*, <https://www2.purpleair.com>.
- [44] R. T. ROCKAFELLAR AND R. J. B. WETS, *Variational Analysis*, Springer Berlin Heidelberg, 1998, <https://doi.org/10.1007/978-3-642-02431-3>.

- [45] N. SHARMA AND K. SAROHA, *A novel dimensionality reduction method for cancer dataset using PCA and feature ranking*, in 2015 International Conference on Advances in Computing, Communications and Informatics (ICACCI), IEEE, Aug. 2015, pp. 2261–2264, <https://doi.org/10.1109/icacci.2015.7275954>.
- [46] Y. SUN, A. BRELOY, P. BABU, D. P. PALOMAR, F. PASCAL, AND G. GINOLHAC, *Low-complexity algorithms for low rank clutter parameters estimation in radar systems*, IEEE Transactions on Signal Processing, 64 (2016), pp. 1986–1998, <https://doi.org/10.1109/tsp.2015.2512535>.
- [47] O. TAMUZ, T. MAZEH, AND S. ZUCKER, *Correcting systematic effects in a large set of photometric light curves*, Monthly Notices of the Royal Astronomical Society, 356 (2005), pp. 1466–1470, <https://doi.org/10.1111/j.1365-2966.2004.08585.x>.
- [48] P. TSALMANTZA AND D. W. HOGG, *A data-driven model for spectra: Finding double redshifts in the sloan digital sky survey*, The Astrophysical Journal, 753 (2012), pp. 1–16, <https://doi.org/10.1088/0004-637x/753/2/122>.
- [49] R. VERSHYNIN, *High-Dimensional Probability*, Cambridge University Press, Sept. 2018, <https://doi.org/10.1017/9781108231596>.
- [50] G. S. WAGNER AND T. J. OWENS, *Signal detection using multi-channel seismic data*, Bulletin of the Seismological Society of America, 86 (1996), pp. 221–231, <https://doi.org/10.1785/bssa08601a0221>.
- [51] H. XI, F. YANG, AND J. YIN, *Convergence of eigenvector empirical spectral distribution of sample covariance matrices*, The Annals of Statistics, 48 (2020), pp. 953–982, <https://doi.org/10.1214/19-aos1832>.
- [52] F. YANG, *Edge universality of separable covariance matrices*, Electronic Journal of Probability, 24 (2019), pp. 1–57, <https://doi.org/10.1214/19-ejp381>.
- [53] G. YOUNG, *Maximum likelihood estimation and factor analysis*, Psychometrika, 6 (1941), pp. 49–53, <https://doi.org/10.1007/bf02288574>.
- [54] A. R. ZHANG, T. T. CAI, AND Y. WU, *Heteroskedastic PCA: Algorithm, optimality, and applications*, The Annals of Statistics, 50 (2022), pp. 53–80, <https://doi.org/10.1214/21-aos2074>.

A Extension to heterogeneous signal strengths

The model stated in [section 2](#) has heteroscedastic noise with a common signal covariance $\mathbf{F}\mathbf{F}^H$. This model may at first appear to limit the scope of [Theorem 2.1](#) to only datasets with homogeneous signal strength. However, as mentioned in [Remark 2.4](#), the result immediately generalizes to cases with heterogeneous signal strengths $\tau_1, \dots, \tau_L > 0$. Namely, suppose $\mathbf{Y}_1, \dots, \mathbf{Y}_L$ are now generated as:

$$\mathbf{Y}_\ell = \sqrt{\tau_\ell} \mathbf{F} \mathbf{Z}_\ell + \mathbf{E}_\ell \in \mathbb{C}^{d \times n_\ell}, \quad \text{for } \ell = 1, \dots, L, \quad (\text{A.1})$$

where $\mathbf{F} \in \mathbb{C}^{d \times k}$, $\mathbf{Z}_\ell \in \mathbb{C}^{k \times n_\ell}$, and $\mathbf{E}_\ell \in \mathbb{C}^{d \times n_\ell}$ are as in [\(2.4\)](#). This model arises, e.g., in low-rank clutter estimation for RADAR [\[8, 9, 46\]](#).

Under this model, we have the following straightforward corollary.

Corollary A.1 (Extension to heterogeneous signal strengths). *Let $\lambda_{i,\ell} := \lambda_i \tau_\ell$. The optimal weights $\mathbf{w}_i^*(\mathbf{Y})$ and corresponding optimal performance $r_i^*(\mathbf{Y})$ converge as*

$$\mathbf{w}_i^*(\mathbf{Y}) \xrightarrow{\text{a.s.}} \left(\frac{1}{v_1} \frac{1}{1 + v_1/\lambda_{i,1}}, \dots, \frac{1}{v_L} \frac{1}{1 + v_L/\lambda_{i,L}} \right) \quad \text{up to scaling}, \quad (\text{A.2})$$

$$r_i^*(\mathbf{Y}) \xrightarrow{\text{a.s.}} \text{the unique solution } x \in (0, 1) \text{ of } \sum_{\ell=1}^L \frac{c_\ell}{v_\ell/\lambda_{i,\ell}} \frac{1-x}{v_\ell/\lambda_{i,\ell} + x} = 1, \quad (\text{A.3})$$

except when $\sum_{\ell=1}^L c_\ell (\lambda_{i,\ell}/v_\ell)^2 \leq 1$, in which case $\mathbf{u}_i$ is asymptotically unrecoverable by any weighted PCA, i.e., $r_i(\mathbf{w}, \mathbf{Y}) \xrightarrow{\text{a.s.}} 0$ for all weights $\mathbf{w}$.

Proof of Corollary A.1. Let $\tilde{\mathbf{Y}}_\ell := \mathbf{Y}_\ell/\sqrt{\tau_\ell}$ and $\tilde{v}_\ell := v_\ell/\tau_\ell$ for $\ell = 1, \dots, L$. Note that $\hat{\mathbf{u}}_i(\mathbf{w}, \mathbf{Y}) = \hat{\mathbf{u}}_i(\mathbf{w} \odot \boldsymbol{\tau}, \tilde{\mathbf{Y}})$, so

$$\begin{aligned} \mathbf{w}_i^*(\tilde{\mathbf{Y}}) \odot \boldsymbol{\tau} &\in \underset{\mathbf{w}}{\operatorname{argmax}} r_i(\mathbf{w} \odot \boldsymbol{\tau}, \tilde{\mathbf{Y}}) = \underset{\mathbf{w}}{\operatorname{argmax}} r_i(\mathbf{w}, \mathbf{Y}), \\ r_i^*(\tilde{\mathbf{Y}}) &= \max_{\mathbf{w}} r_i(\mathbf{w} \odot \boldsymbol{\tau}, \tilde{\mathbf{Y}}) = \max_{\mathbf{w}} r_i(\mathbf{w}, \mathbf{Y}), \end{aligned}$$

where $\odot$ and $\oslash$ denote entrywise multiplication and division. Since $\tilde{\mathbf{Y}}_1, \dots, \tilde{\mathbf{Y}}_L$ obeys [\(2.4\)](#) with noise variances $\tilde{v}_1, \dots, \tilde{v}_L$, applying [Theorem 2.1](#) and simplifying yields [\(A.2\)](#) and [\(A.3\)](#). $\square$

Note that the form of the weights is the same as before: inverse noise variance weights times an SNR-dependent term that downweights low-SNR blocks. Moreover, note that data with homoscedastic noise and heterogeneous signal strengths are uniformly weighted by inverse noise variance weights but optimal weights downweight the blocks with weaker signals. When the signal strengths $\boldsymbol{\tau}$, signal variance λ_i , and noise variances $\mathbf{v}$ are unknown, estimates of these parameters may be obtained using any one of various existing approaches, and which one is most appropriate will depend on the specific application. Notably, the simple estimators in [Example 7.2](#) could still be used in this case to estimate the signal and noise variances, and the same ideas may be adapted to also estimate the signal strengths. For example, one could apply the same signal variance estimator on each block individually, then divide the resulting per-block signal variance estimates by the signal variance estimated from the data overall.

B Proof of Lemma 6.1

First, we prove convergence of the maximum, i.e., that

$$\forall \varepsilon > 0 \quad \exists N \quad \forall n > N \quad \left| \max_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) - \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) \right| < \varepsilon.$$

Let $\varepsilon > 0$ be arbitrary, and let $\mathbf{x}^*$ be the unique maximizer of f . Since $f_n \rightarrow f$ uniformly, there exists N such that

$$\forall_{n>N} \forall_{\mathbf{x} \in \mathcal{X}} |f_n(\mathbf{x}) - f(\mathbf{x})| < \varepsilon.$$

Now, let $n > N$ be arbitrary. Since f_n has a maximum, there exists $\tilde{\mathbf{x}} \in \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x})$. Noting that

$$\begin{aligned} \max_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) &= f_n(\tilde{\mathbf{x}}) < f(\tilde{\mathbf{x}}) + \varepsilon \leq f(\mathbf{x}^*) + \varepsilon = \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) + \varepsilon, \\ \max_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) &\geq f_n(\mathbf{x}^*) > f(\mathbf{x}^*) - \varepsilon = \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x}) - \varepsilon, \end{aligned}$$

yields $|\max_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x}) - \max_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x})| < \varepsilon$, establishing convergence for the maximum.

Next, we prove convergence of the set of maximizers. Let $\mathbf{x}^*$ be the unique maximizer of f . We want to show that the set of maximizers of f_n (nonempty by assumption) converges to $\{\mathbf{x}^*\}$ with respect to the Hausdorff distance, i.e., that $\sup\{\|\tilde{\mathbf{x}} - \mathbf{x}^*\|_2 : \tilde{\mathbf{x}} \in \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x})\} \rightarrow 0$, or equivalently

$$\forall_{\delta>0} \exists_N \forall_{n>N} \forall_{\tilde{\mathbf{x}} \in \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x})} \|\tilde{\mathbf{x}} - \mathbf{x}^*\|_2 < \delta.$$

Let $\delta > 0$ be arbitrary. Note that $\mathcal{K} := \{\mathbf{x} \in \mathcal{X} : \|\mathbf{x} - \mathbf{x}^*\|_2 \geq \delta\}$ is compact since $\mathcal{X}$ is compact. Since f is continuous, it then follows from the extreme value theorem that f is maximized over $\mathcal{K}$ at some point $\hat{\mathbf{x}}$. Next, since $f_n \rightarrow f$ uniformly and the maximizer $\mathbf{x}^*$ is unique, there exists N such that

$$\forall_{n>N} \forall_{\mathbf{x} \in \mathcal{X}} |f_n(\mathbf{x}) - f(\mathbf{x})| < \frac{f(\mathbf{x}^*) - f(\hat{\mathbf{x}})}{2}.$$

Now, let $n > N$ be arbitrary. Let $\tilde{\mathbf{x}} \in \operatorname{argmax}_{\mathbf{x} \in \mathcal{X}} f_n(\mathbf{x})$. Then

$$\begin{aligned} f(\mathbf{x}^*) - f(\tilde{\mathbf{x}}) &= \underbrace{f(\mathbf{x}^*) - f_n(\mathbf{x}^*)}_{< (f(\mathbf{x}^*) - f(\hat{\mathbf{x}}))/2} + \underbrace{f_n(\mathbf{x}^*) - f_n(\tilde{\mathbf{x}})}_{\leq 0} + \underbrace{f_n(\tilde{\mathbf{x}}) - f(\tilde{\mathbf{x}})}_{< (f(\mathbf{x}^*) - f(\hat{\mathbf{x}}))/2} < f(\mathbf{x}^*) - f(\hat{\mathbf{x}}), \end{aligned}$$

so $f(\tilde{\mathbf{x}}) > f(\hat{\mathbf{x}}) = \max_{\mathbf{x} \in \mathcal{K}} f(\mathbf{x})$. Thus $\tilde{\mathbf{x}} \notin \mathcal{K}$, and so $\|\tilde{\mathbf{x}} - \mathbf{x}^*\|_2 < \delta$, concluding the proof.

C Detailed calculations from the proof of Lemma 6.2

This section provides detailed calculations from the proof of Lemma 6.2 in subsection 6.1.

C.1 Anisotropic local law

This section states an *anisotropic local law* for $\mathbf{G}(\mathbf{w}, \zeta)$ in a form that will be convenient for the rest of the proof. Roughly speaking, it says that the random resolvent matrix $\mathbf{G}(\mathbf{w}, \zeta)$ is well approximated in the limit by the deterministic matrix

$$\mathbf{\Pi}(\mathbf{w}, \zeta) := \begin{bmatrix} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \mathbf{I}_p & \\ & (\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{V} \mathbf{W})^{-1} \end{bmatrix}, \quad (\text{C.1})$$

where $\mathbf{V} := \text{blockdiag}(v_1 \mathbf{I}_{n_1}, \dots, v_L \mathbf{I}_{n_L})$, $\mathbf{W} := \text{blockdiag}(w_1 \mathbf{I}_{n_1}, \dots, w_L \mathbf{I}_{n_L})$, and $m_{\text{MP}}(\mathbf{w}, \zeta)$ is the Stieltjes transform defined in (6.5). The local law is stated over two domains: S_{edge} (around the edge $b_{\mathbf{w}}$ of the support of $\mu_{\mathbf{w}}$) and S_{out} (outside the support of $\mu_{\mathbf{w}}$). It follows straightforwardly from [14, 51]. Similar local laws have also been proved in [7, 30, 52] under various different assumptions.

Lemma C.1 (Anisotropic local laws). *Under the model assumptions, there exists a constant $\kappa > 0$ such that for any constants $\tau, \delta > 0$ and any deterministic sequences of unit vectors $\boldsymbol{\xi}_n, \boldsymbol{\nu}_n \in \mathbb{C}^{d+n}$,*

$$\sup_{\mathbf{w} \in \Delta_L} \sup_{\zeta \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa)} |\boldsymbol{\xi}_n^H [\mathbf{G}(\mathbf{w}, \zeta) - \mathbf{\Pi}(\mathbf{w}, \zeta)] \boldsymbol{\nu}_n| \xrightarrow{\text{a.s.}} 0, \quad (\text{C.2})$$

$$\sup_{\mathbf{w} \in \Delta_L} \sup_{\zeta \in S_{\text{out}}(\mathbf{w}, \tau)} |\boldsymbol{\xi}_n^H [\mathbf{G}(\mathbf{w}, \zeta) - \mathbf{\Pi}(\mathbf{w}, \zeta)] \boldsymbol{\nu}_n| \xrightarrow{\text{a.s.}} 0, \quad (\text{C.3})$$

where

$$S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa) := \{\zeta \in \mathbb{C} : (\text{Re } \zeta, \text{Im } \zeta) \in [b_{\mathbf{w}} - \kappa, b_{\mathbf{w}} + \tau) \times [n^{-1+\delta}, \tau]\}, \quad (\text{C.4})$$

$$S_{\text{out}}(\mathbf{w}, \tau) := \{\zeta \in \mathbb{C} : (\text{Re } \zeta, \text{Im } \zeta) \in [b_{\mathbf{w}} + \tau, \infty) \times [-1, 1]\}. \quad (\text{C.5})$$

Proof of Lemma C.1. Let $\bar{\mathbf{E}} := \mathbf{E}\mathbf{V}^{-1/2}$, i.e., $\bar{\mathbf{E}}$ is whitened noise. It follows immediately from the moment condition on $\mathbf{E}$ that there exist constants $a > 4$ and $C_0 > 0$ such that

$$\max_{i,j} \mathbb{E}|\bar{\mathbf{E}}_{ij}|^a \leq C_0. \quad (\text{C.6})$$

The local laws [14, 30, 51] are stated for matrices with truncated entries, so we define

$$\hat{\mathbf{E}}_{\mathbf{w}} := \hat{\mathbf{E}}\mathbf{V}^{1/2}\mathbf{W}^{1/2} \quad \text{where} \quad \hat{\mathbf{E}}_{ij} := \begin{cases} \bar{\mathbf{E}}_{ij} & \text{if } |\bar{\mathbf{E}}_{ij}| \leq \phi_n n^{\varepsilon_0}, \\ 0 & \text{otherwise,} \end{cases} \quad (\text{C.7})$$

where $\phi_n := n^{2/a}$ and $\varepsilon_0 \in (0, 1/2 - 2/a)$ is a constant to be chosen later. Likewise, let $\hat{\mathbf{G}}(\mathbf{w}, \zeta)$ be the analogue of $\mathbf{G}(\mathbf{w}, \zeta)$ with $\hat{\mathbf{E}}_{\mathbf{w}}$ used in place of $\tilde{\mathbf{E}}_{\mathbf{w}}$.

Note that it follows from (C.6) and integration by parts that

$$|\mathbb{E}\hat{\mathbf{E}}_{ij}| \leq n^{-3/2-\varepsilon_0}, \quad \mathbb{E}|\hat{\mathbf{E}}_{ij}|^2 \leq n^{-1-\varepsilon_0}, \quad \mathbb{E}|\hat{\mathbf{E}}_{ij}|^3 = O(1), \quad \mathbb{E}|\hat{\mathbf{E}}_{ij}|^4 = O(1). \quad (\text{C.8})$$

Moreover, the empirical spectral distribution of $\mathbf{V}^{1/2}\mathbf{W}^{1/2}$ satisfies [30, Example 2.8], which leads to the edge regularity condition

$$\forall \mathbf{w} \in \Delta_L \quad \min_{\ell} [1 - c^{-1} m_{\text{MP}}(\mathbf{w}, b_{\mathbf{w}}) \cdot v_{\ell} w_{\ell}] > 0. \quad (\text{C.9})$$

Under (C.8), $\max_{i,j} |\hat{\mathbf{E}}_{ij}| \leq \phi_n n^{\varepsilon_0}$, and the edge regularity condition (C.9), the following local laws follow immediately from [14, 51]: for any $\mathbf{w} \in \Delta_L$, there exists a constant $\kappa > 0$ such that for any constants $\tau, \delta > 0$, any deterministic unit vectors $\boldsymbol{\xi}, \boldsymbol{\nu} \in \mathbb{C}^{d+n}$, and any $\varepsilon_1, C_1 > 0$, eventually (i.e., for d large enough), with probability at least $1 - n^{-C_1}$ we have

$$\sup_{\zeta \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa)} |\boldsymbol{\xi}^H [\hat{\mathbf{G}}(\mathbf{w}, \zeta) - \boldsymbol{\Pi}(\mathbf{w}, \zeta)] \boldsymbol{\nu}| \leq \phi_n n^{-1/2+\varepsilon_0+\varepsilon_1} + \frac{n^{\varepsilon_1}}{\sqrt{n \text{Im } \zeta}}, \quad (\text{C.10})$$

$$\sup_{\zeta \in S_{\text{out}}(\mathbf{w}, \tau)} |\boldsymbol{\xi}^H [\hat{\mathbf{G}}(\mathbf{w}, \zeta) - \boldsymbol{\Pi}(\mathbf{w}, \zeta)] \boldsymbol{\nu}| \leq \phi_n n^{-1/2+\varepsilon_0+\varepsilon_1}. \quad (\text{C.11})$$

Namely, (C.10) follows from [51, Theorem 3.4] and (C.11) follows from [14, Theorem S.3.12].

Using a standard ε -net argument and taking a union bound with respect to $\mathbf{w}$ yields (C.10) and (C.11) uniformly in $\mathbf{w} \in \Delta_L$. Moreover, choosing ε_0 and ε_1 small enough yields $\phi_n n^{-1/2+\varepsilon_0+\varepsilon_1} \rightarrow 0$ and $n^{\varepsilon_1}/\sqrt{n \text{Im } \zeta} \leq n^{-\delta/2+\varepsilon_1} \rightarrow 0$ for all $\zeta \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa)$. Thus, applying the Borel-Cantelli lemma yields analogues of (C.2) and (C.3) with $\hat{\mathbf{E}}$ in place of $\bar{\mathbf{E}}$. The proof concludes by observing that³

$$\begin{aligned} \Pr(\hat{\mathbf{E}} \neq \bar{\mathbf{E}} \text{ i.o.}) &= \lim_{k \rightarrow \infty} \Pr\left(\bigcup_{d=2^k}^{\infty} \bigcup_{i=1}^d \bigcup_{j=1}^{n(d)} \{|\bar{\mathbf{E}}_{ij}| \geq \phi_n n^{\varepsilon_0}\}\right) \\ &= \lim_{k \rightarrow \infty} \Pr\left(\bigcup_{t=k}^{\infty} \bigcup_{d \in [2^t, 2^{t+1})} \bigcup_{i=1}^d \bigcup_{j=1}^{n(d)} \{|\bar{\mathbf{E}}_{ij}| \geq \phi_n n^{\varepsilon_0}\}\right) \\ &\leq C_0 \lim_{k \rightarrow \infty} \sum_{t=k}^{\infty} (2^{t+1})^2 (2^{t(2/a+\varepsilon_0)})^{-a} \leq 4C_0 \lim_{k \rightarrow \infty} \sum_{t=k}^{\infty} 2^{-ta\varepsilon_0} = 0, \end{aligned} \quad (\text{C.12})$$

i.e. $\hat{\mathbf{E}} = \bar{\mathbf{E}}$ eventually, almost surely. □

³In the derivation, we regard $n \equiv n(d)$ as a sequence indexed by d , which has the same order as d .

C.2 Detailed derivation of (6.7)

It follows immediately from [51, Theorem 3.7] that

$$\Pr \left\{ \sup_{\mathbf{w} \in \Delta_L} |\|\hat{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}} - b_{\mathbf{w}}| \rightarrow 0 \right\} = 1, \quad (\text{C.13})$$

where $\hat{\mathbf{E}}_{\mathbf{w}}$ is the truncated weighted noise defined in (C.7). Combining (C.13) with (C.12) (i.e., the fact that $\hat{\mathbf{E}}_{\mathbf{w}} = \tilde{\mathbf{E}}_{\mathbf{w}}$ eventually, almost surely) yields (6.7).

C.3 Detailed derivation of (6.10)

Equation (6.10) modifies [5, Lemma 4.1] to account for the weights and is proved in essentially the same way. Namely, recall that by [23, Theorem 7.3.3], ζ is a singular value of $\tilde{\mathbf{Y}}_{\mathbf{w}}$ if and only if it is a root of the characteristic polynomial

$$0 = \det \left(\zeta \mathbf{I}_{d+n} - \begin{bmatrix} \tilde{\mathbf{Y}}_{\mathbf{w}} & \\ & \tilde{\mathbf{Y}}_{\mathbf{w}}^{\text{H}} \end{bmatrix} \right) \quad (\text{C.14})$$

$$= \det \left(\underbrace{\zeta \mathbf{I}_{d+n} - \begin{bmatrix} \tilde{\mathbf{E}}_{\mathbf{w}} & \\ & \tilde{\mathbf{E}}_{\mathbf{w}}^{\text{H}} \end{bmatrix}}_{\mathbf{A}} - \underbrace{\begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}}_{\mathbf{B}} \underbrace{\begin{bmatrix} \boldsymbol{\Theta} & \\ & \boldsymbol{\Theta} \end{bmatrix}}_{\mathbf{D}} \underbrace{\begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}}}_{\mathbf{C}} \right) \quad (\text{C.15})$$

$$= \det(\mathbf{A}) \det(\mathbf{D}) \det(-\mathbf{M}(\mathbf{w}, \zeta)), \quad (\text{C.16})$$

where (C.15) is a convenient form of the matrix, and (C.16) follows from the identity

$$\det(\mathbf{A} - \mathbf{BDC}) = \det(\mathbf{A}) \det(\mathbf{D}) \det(\mathbf{D}^{-1} - \mathbf{CA}^{-1}\mathbf{B}), \quad (\text{C.17})$$

for invertible matrices $\mathbf{A}$ and $\mathbf{D}$. Note that $\mathbf{A}$ above is invertible because $\zeta > \|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}}$ (indeed, one only needs that ζ is not a singular value of $\tilde{\mathbf{E}}_{\mathbf{w}}$). As a result, (C.16) is zero if and only if $\det \mathbf{M}(\mathbf{w}, \zeta) = 0$, yielding (6.10).

C.4 Detailed derivation of (6.11)

Let $\tau > 0$ be arbitrary, and decompose $\mathbf{M}$ as

$$\begin{aligned} \mathbf{M}(\mathbf{w}, \zeta) &= \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} [\mathbf{G}(\mathbf{w}, \zeta) - \boldsymbol{\Pi}(\mathbf{w}, \zeta)] \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} \\ &\quad + \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} \boldsymbol{\Pi}(\mathbf{w}, \zeta) \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} - \begin{bmatrix} \boldsymbol{\Theta}^{-1} & \\ & \boldsymbol{\Theta}^{-1} \end{bmatrix}. \end{aligned} \quad (\text{C.18})$$

We will show that the first term vanishes and the final two terms converge to $\bar{\mathbf{M}}(\mathbf{w}, \zeta)$.

For the first term of (C.18), applying (C.3) from Lemma C.1 yields

$$\left\| \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} [\mathbf{G}(\mathbf{w}, \zeta) - \boldsymbol{\Pi}(\mathbf{w}, \zeta)] \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} \right\|_{\text{op}} \xrightarrow[\substack{\text{a.s.} \\ (\mathbf{w}, \zeta) \in \Omega(\tau)}]{\Rightarrow} 0, \quad (\text{C.19})$$

where we use the observation that $\Omega(\tau) = \{(\mathbf{w}, \zeta) \in \Delta_L \times \mathbb{C} : \zeta \in S_{\text{out}}(\mathbf{w}, \tau)\}$.

For the second term of (C.18), note that

$$\begin{aligned} &\begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} \boldsymbol{\Pi}(\mathbf{w}, \zeta) \begin{bmatrix} \mathbf{U} & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} \\ &= \begin{bmatrix} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \mathbf{I}_k & \\ & \tilde{\mathbf{Q}}_{\mathbf{w}}^{\text{H}} (\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{VW})^{-1} \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}. \end{aligned} \quad (\text{C.20})$$

With the regularity condition (C.9), it follows immediately from [13, Lemma A.6] that

$$\inf_{\mathbf{w} \in \Delta_L} \inf_{\zeta \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa) \cup S_{\text{out}}(\mathbf{w}, \tau)} \min_{i=1, \dots, n} \left[1 - c^{-1} m_{\text{MP}}(\mathbf{w}, \zeta) \cdot (\mathbf{VW})_{ii} \right] > 0. \quad (\text{C.21})$$

Note that [13, Lemma A.6] states the result for fixed $\mathbf{w}$; taking the minimum over $\mathbf{w}$ in the compact set Δ_L yields (C.21). Thus, $(\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{VW})^{-1}$ is a diagonal matrix with bounded entries. Then using a standard ε -net argument together with the law of large numbers yields that

$$\begin{aligned} & \tilde{\mathbf{Q}}_{\mathbf{w}}^H (\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{VW})^{-1} \tilde{\mathbf{Q}}_{\mathbf{w}} \\ & - \frac{1}{n} \text{tr} \left[\mathbf{W}^{1/2} (\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{VW})^{-1} \mathbf{W}^{1/2} \right] \cdot \mathbf{I}_k \xrightarrow[\substack{\text{a.s.} \\ (\mathbf{w}, \zeta) \in \Omega(\tau)}]{\text{a.s.}} \mathbf{0}_{k \times k}. \end{aligned} \quad (\text{C.22})$$

Note also that

$$\frac{1}{n} \text{tr} \left[\mathbf{W}^{1/2} (\zeta \mathbf{I}_n - c^{-1} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) \cdot \mathbf{VW})^{-1} \mathbf{W}^{1/2} \right] = \sum_{\ell=1}^L \frac{p_{\ell} w_{\ell}}{\zeta - w_{\ell} v_{\ell} \zeta m_{\text{MP}}(\mathbf{w}, \zeta) / c}. \quad (\text{C.23})$$

Combining (C.18)–(C.20), (C.22), and (C.23) with $\zeta m_{\text{MP}}(\mathbf{w}, \zeta) = \varphi_{1, \mathbf{w}}(\zeta)$ yields (6.11).

C.5 Verification of properties for $\varphi_{1, \mathbf{w}}$ and $\varphi_{2, \mathbf{w}}$

We verify that, for any $\mathbf{w} \in \Delta_L$, the functions $\varphi_{1, \mathbf{w}}$ and $\varphi_{2, \mathbf{w}}$ have the following properties:

$$\forall_{\zeta > b_{\mathbf{w}}} \varphi_{1, \mathbf{w}}(\zeta) > 0, \quad \varphi_{1, \mathbf{w}}(\zeta) \rightarrow 0 \text{ as } |\zeta| \rightarrow \infty, \quad \forall_{\zeta \notin \text{supp}(\mu_{\mathbf{w}})} \varphi_{1, \mathbf{w}}(\zeta) \in \mathbb{R} \Leftrightarrow \zeta \in \mathbb{R}, \quad (\text{C.24})$$

$$\forall_{\zeta > b_{\mathbf{w}}} \varphi_{2, \mathbf{w}}(\zeta) > 0, \quad \varphi_{2, \mathbf{w}}(\zeta) \rightarrow 0 \text{ as } |\zeta| \rightarrow \infty, \quad \forall_{\zeta \notin \text{supp}(\mu_{\mathbf{w}})} \varphi_{2, \mathbf{w}}(\zeta) \in \mathbb{R} \Leftrightarrow \zeta \in \mathbb{R}, \quad (\text{C.25})$$

where $\text{supp}(\mu_{\mathbf{w}})$ denotes the support of $\mu_{\mathbf{w}}$. First, we verify the properties in (C.24) for $\varphi_{1, \mathbf{w}}$:

- (a) For any $\zeta > b_{\mathbf{w}}$, the integrand in (6.12) is positive and bounded away from zero since the support of $\mu_{\mathbf{w}}$ lies between zero and $b_{\mathbf{w}}$.

Thus, $\forall_{\zeta > b_{\mathbf{w}}} \varphi_{1, \mathbf{w}}(\zeta) > 0$.

- (b) As $|\zeta| \rightarrow \infty$, the integrand in (6.12) goes to zero uniformly in t .

Thus, $\varphi_{1, \mathbf{w}}(\zeta) \rightarrow 0$ as $|\zeta| \rightarrow \infty$.

- (c) For $\zeta \notin \text{supp}(\mu_{\mathbf{w}})$, the imaginary part of $\varphi_{1, \mathbf{w}}(\zeta)$ is

$$\text{Im} \varphi_{1, \mathbf{w}}(\zeta) = \int \text{Im} \left(\frac{\zeta}{\zeta^2 - t^2} \right) d\mu_{\mathbf{w}}(t) = - \text{Im} \zeta \underbrace{\int \frac{|\zeta|^2 + t^2}{|\zeta^2 - t^2|^2} d\mu_{\mathbf{w}}(t)}_{>0}.$$

Thus, $\forall_{\zeta \notin \text{supp}(\mu_{\mathbf{w}})} \varphi_{1, \mathbf{w}}(\zeta) \in \mathbb{R} \Leftrightarrow \zeta \in \mathbb{R}$.

Next we verify the properties in (C.25) for $\varphi_{2, \mathbf{w}}$:

- (a) For any $\zeta > b_{\mathbf{w}}$, $m_{\text{MP}}(\mathbf{w}, \zeta) \leq m_{\text{MP}}(\mathbf{w}, b_{\mathbf{w}})$ since the integrand in (6.5) is decreasing in ζ . So, it follows from (C.9) that for $\ell = 1, \dots, L$,

$$1 - c^{-1} m_{\text{MP}}(\mathbf{w}, \zeta) \cdot v_{\ell} w_{\ell} \geq 1 - c^{-1} m_{\text{MP}}(\mathbf{w}, b_{\mathbf{w}}) \cdot v_{\ell} w_{\ell} > 0,$$

so the denominator of each summand in the definition of $\varphi_{2, \mathbf{w}}$ in (6.12) is also positive, i.e., $\zeta - w_{\ell} v_{\ell} \varphi_{1, \mathbf{w}}(\zeta) / c = \zeta [1 - c^{-1} m_{\text{MP}}(\mathbf{w}, \zeta) \cdot v_{\ell} w_{\ell}] > 0$.

Thus, $\forall_{\zeta > b_{\mathbf{w}}} \varphi_{2, \mathbf{w}}(\zeta) > 0$.

(b) As $|\zeta| \rightarrow \infty$, $|\zeta - w_\ell v_\ell \varphi_{1,w}(\zeta)/c| \rightarrow \infty$ for each $\ell \in \{1, \dots, L\}$ since $\varphi_{1,w}(\zeta) \rightarrow 0$ as shown above.

Thus, $\varphi_{2,w}(\zeta) \rightarrow 0$ as $|\zeta| \rightarrow \infty$.

(c) As shown above, for any $\zeta \notin \text{supp}(\mu_w)$, $\text{Im } \varphi_{1,w}(\zeta)$ is zero if $\text{Im } \zeta$ is zero and has the opposite sign of $\text{Im } \zeta$ otherwise. As a result,

$$\text{Im}\{\zeta - w_\ell v_\ell \varphi_{1,w}(\zeta)/c\} = \text{Im } \zeta - (w_\ell v_\ell / c) \text{Im } \varphi_{1,w}(\zeta)$$

is zero if $\text{Im } \zeta$ is zero and has the same sign as $\text{Im } \zeta$ otherwise.

Thus, $\forall_{\zeta \notin \text{supp}(\mu_w)} \varphi_{2,w}(\zeta) \in \mathbb{R} \Leftrightarrow \zeta \in \mathbb{R}$.

As a result, $\varphi_{1,w}$ and $\varphi_{2,w}$ satisfy the needed conditions of [5, Lemma A.1].

C.6 Detailed derivation of (6.14)

Equation (6.14) modifies [5, Lemma 5.1] to account for the weights and is proved in essentially the same way. Namely, let w be such that $\hat{\theta}_{i,w} > \|\tilde{E}_w\|_{\text{op}}$ and let $\tilde{X}_w := U\Theta\tilde{Q}_w^H$ be the weighted and normalized signal.

To derive (6.14a), we first use the block matrix inverse [23, Equation (0.7.3.1)] to write the resolvent $G(w, \zeta)$ from (6.9) in the following block matrix form:

$$G(w, \zeta) = \begin{bmatrix} \zeta(\zeta^2 I_d - \tilde{E}_w \tilde{E}_w^H)^{-1} & (\zeta^2 I_d - \tilde{E}_w \tilde{E}_w^H)^{-1} \tilde{E}_w \\ \tilde{E}_w^H (\zeta^2 I_d - \tilde{E}_w \tilde{E}_w^H)^{-1} & \zeta(\zeta^2 I_n - \tilde{E}_w^H \tilde{E}_w)^{-1} \end{bmatrix}. \quad (\text{C.26})$$

Then (6.14a) follows by substituting (6.8) and (C.26) then factoring. Namely,

$$M(w, \hat{\theta}_{i,w}) \begin{bmatrix} \Theta \tilde{Q}_w^H \hat{q}_{i,w} \\ \Theta U^H \hat{u}_{i,w} \end{bmatrix} \quad (\text{C.27})$$

$$= \begin{bmatrix} U^H ((\hat{\theta}_{i,w}^2 I_d - \tilde{E}_w \tilde{E}_w^H)^{-1} (\hat{\theta}_{i,w} \tilde{X}_w \hat{q}_{i,w} + \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w}) - \hat{u}_{i,w}) \\ \tilde{Q}_w^H ((\hat{\theta}_{i,w}^2 I_n - \tilde{E}_w^H \tilde{E}_w)^{-1} (\tilde{E}_w^H \tilde{X}_w \hat{q}_{i,w} + \hat{\theta}_{i,w} \tilde{X}_w^H \hat{u}_{i,w}) - \hat{q}_{i,w}) \end{bmatrix} \quad (\text{C.28})$$

$$= \begin{bmatrix} U^H (\hat{u}_{i,w} - \hat{u}_{i,w}) \\ \tilde{Q}_w^H (\hat{q}_{i,w} - \hat{q}_{i,w}) \end{bmatrix} = 0, \quad (\text{C.29})$$

where (C.28) uses the identity $\tilde{E}_w^H (\hat{\theta}_{i,w}^2 I_d - \tilde{E}_w \tilde{E}_w^H)^{-1} = (\hat{\theta}_{i,w}^2 I_n - \tilde{E}_w^H \tilde{E}_w)^{-1} \tilde{E}_w^H$, and (C.29) follows by substituting $\tilde{X}_w = \tilde{Y}_w - \tilde{E}_w$ and using the singular vector identities

$$\tilde{Y}_w \hat{q}_{i,w} = \hat{\theta}_{i,w} \hat{u}_{i,w}, \quad \tilde{Y}_w^H \hat{u}_{i,w} = \hat{\theta}_{i,w} \hat{q}_{i,w}.$$

To derive (6.14b), combine the identity $\hat{u}_{i,w} = \Gamma_w (\hat{\theta}_{i,w} \tilde{X}_w \hat{q}_{i,w} + \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w})$ used to obtain (C.29) with the fact that $\|\hat{u}_{i,w}\|_2 = 1$ and expand as

$$\begin{aligned} 1 &= \hat{u}_{i,w}^H \hat{u}_{i,w} \\ &= (\hat{\theta}_{i,w} \tilde{X}_w \hat{q}_{i,w} + \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w})^H \Gamma_w^2 (\hat{\theta}_{i,w} \tilde{X}_w \hat{q}_{i,w} + \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w}) \\ &= \chi_1(w) + \chi_2(w) + 2 \text{Re } \chi_3(w), \end{aligned}$$

where the outer terms are

$$\chi_1(w) := \hat{q}_{i,w}^H \tilde{X}_w^H \hat{\theta}_{i,w}^2 \Gamma_w^2 \tilde{X}_w \hat{q}_{i,w}, \quad \chi_2(w) := \hat{u}_{i,w}^H \tilde{X}_w \tilde{E}_w^H \Gamma_w^2 \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w}, \quad (\text{C.30})$$

and the cross term is

$$\chi_3(w) := \hat{q}_{i,w}^H \tilde{X}_w^H \hat{\theta}_{i,w} \Gamma_w^2 \tilde{E}_w \tilde{X}_w^H \hat{u}_{i,w}. \quad (\text{C.31})$$

Expanding $\tilde{X}_w = U\Theta\tilde{Q}_w^H = \theta_1 u_1 \tilde{q}_{1,w}^H + \dots + \theta_k u_k \tilde{q}_{k,w}^H$ in (C.30) and (C.31) then simplifying yields (6.15).

C.7 Detailed derivation of (6.16)

To derive (6.16), it is helpful to first establish that several terms appearing in (6.15) are bounded. Namely, almost surely, eventually, the following upper bounds hold for all $\mathbf{w} \in \mathcal{W}_{>}(\nu)$:

$$\hat{\theta}_{i,\mathbf{w}} \leq \|\tilde{\mathbf{Y}}_{\mathbf{w}}\|_{\text{op}} = \|\tilde{\mathbf{Y}}_{1_L} \mathbf{W}^{1/2}\|_{\text{op}} \leq \|\tilde{\mathbf{Y}}_{1_L}\|_{\text{op}} \|\mathbf{W}^{1/2}\|_{\text{op}} \leq \|\tilde{\mathbf{Y}}_{1_L}\|_{\text{op}} < 2 \cdot \bar{\theta}_{1,1_L}, \quad (\text{C.32a})$$

$$\|\Gamma_{\mathbf{w}}\|_{\text{op}} = \|(\hat{\theta}_{i,\mathbf{w}}^2 \mathbf{I}_d - \tilde{\mathbf{E}}_{\mathbf{w}} \tilde{\mathbf{E}}_{\mathbf{w}}^H)^{-1}\|_{\text{op}} < 2 \cdot \frac{1}{\bar{\theta}_{i,\mathbf{w}}^2 - b_{\mathbf{w}}^2} < 2 \cdot \frac{1}{\nu^2}, \quad (\text{C.32b})$$

$$\|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}} = \|\tilde{\mathbf{E}}_{1_L} \mathbf{W}^{1/2}\|_{\text{op}} \leq \|\tilde{\mathbf{E}}_{1_L}\|_{\text{op}} \|\mathbf{W}^{1/2}\|_{\text{op}} \leq \|\tilde{\mathbf{E}}_{1_L}\|_{\text{op}} < 2 \cdot b_{1_L}, \quad (\text{C.32c})$$

$$\|\tilde{\mathbf{q}}_{j,\mathbf{w}}\|_2 = \|\mathbf{W}^{1/2} \tilde{\mathbf{q}}_{j,1_L}\|_2 \leq \|\mathbf{W}^{1/2}\|_{\text{op}} \|\tilde{\mathbf{q}}_{j,1_L}\|_2 < 2, \quad (\text{C.32d})$$

$$|\tilde{\mathbf{q}}_{j,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}}| \leq \|\tilde{\mathbf{q}}_{j,\mathbf{w}}\|_2 \|\hat{\mathbf{q}}_{i,\mathbf{w}}\|_2 < 2, \quad (\text{C.32e})$$

$$|\mathbf{u}_j^H \hat{\mathbf{u}}_{i,\mathbf{w}}| \leq \|\mathbf{u}_j\|_2 \|\hat{\mathbf{u}}_{i,\mathbf{w}}\|_2 = 1 \quad (\text{C.32f})$$

where

- (C.32a) follows from the identity $\tilde{\mathbf{Y}}_{\mathbf{w}} = \tilde{\mathbf{Y}}_{1_L} \mathbf{W}^{1/2}$, submultiplicativity of the operator norm, the fact that $\|\mathbf{W}^{1/2}\|_{\text{op}} = \|\mathbf{w}\|_{\infty}^{1/2} \leq 1$ for $\mathbf{w} \in \Delta_L$, and (6.13);
- (C.32b) follows from (6.7) and (6.13), and the fact that $\bar{\theta}_{i,\mathbf{w}}^2 > (b_{\mathbf{w}} + \nu)^2 > b_{\mathbf{w}}^2 + \nu^2$ for $\mathbf{w} \in \mathcal{W}_{>}(\nu)$;
- (C.32c) follows from the identity $\tilde{\mathbf{E}}_{\mathbf{w}} = \tilde{\mathbf{E}}_{1_L} \mathbf{W}^{1/2}$, submultiplicativity of the operator norm, the fact that $\|\mathbf{W}^{1/2}\|_{\text{op}} = \|\mathbf{w}\|_{\infty}^{1/2} \leq 1$ for $\mathbf{w} \in \Delta_L$, and (6.7);
- (C.32d) follows from the identity $\tilde{\mathbf{q}}_{j,\mathbf{w}} = \mathbf{W}^{1/2} \tilde{\mathbf{q}}_{j,1_L}$, the operator norm inequality, the fact that $\|\mathbf{W}^{1/2}\|_{\text{op}} = \|\mathbf{w}\|_{\infty}^{1/2} \leq 1$ for $\mathbf{w} \in \Delta_L$, and the fact that $\|\tilde{\mathbf{q}}_{j,1_L}\|_2^2 \xrightarrow{\text{a.s.}} 1$ by the law of large numbers;
- (C.32e) follows from the Cauchy-Schwarz inequality and (C.32d);
- (C.32f) follows from the Cauchy-Schwarz inequality.

In other words, these terms are almost surely eventually uniformly bounded.

We now derive (6.16a) and (6.16b). Note first that almost surely, eventually, for any $\mathbf{w} \in \mathcal{W}_{>}(\nu)$, $\mathbf{M}(\mathbf{w}, \cdot)$ and $\bar{\mathbf{M}}(\mathbf{w}, \cdot)$ are both Lipschitz functions for $\zeta > b_{\mathbf{w}} + \nu/2$ with Lipschitz constant $O(1/\nu^2)$. Moreover, almost surely, eventually, for all $\mathbf{w} \in \mathcal{W}_{>}(\nu)$, $\hat{\theta}_{i,\mathbf{w}} > b_{\mathbf{w}} + \nu/2$. Hence, it follows from the limits (6.11) and (6.13) that

$$\mathbf{M}(\mathbf{w}, \hat{\theta}_{i,\mathbf{w}}) \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} \bar{\mathbf{M}}(\mathbf{w}, \rho_{i,\mathbf{w}}). \quad (\text{C.33})$$

Applying this to (6.14a) yields

$$\begin{bmatrix} \xi(\mathbf{w}) \\ \eta(\mathbf{w}) \end{bmatrix} := \text{proj}_{(\ker \bar{\mathbf{M}}(\mathbf{w}, \rho_{i,\mathbf{w}}))^{\perp}} \begin{bmatrix} \Theta \tilde{\mathbf{Q}}_{\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}} \\ \Theta \mathbf{U}^H \hat{\mathbf{u}}_{i,\mathbf{w}} \end{bmatrix} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0. \quad (\text{C.34})$$

Observe next that, similar to [5, Section 5],

$$\forall \mathbf{w} \in \mathcal{W}_{>}(\nu) \quad \ker \bar{\mathbf{M}}(\mathbf{w}, \rho_{i,\mathbf{w}}) = \left\{ \begin{bmatrix} s \\ t \end{bmatrix} \in \mathbb{C}^{2k} : \begin{array}{ll} t_j = \theta_i \varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}}) s_j & \text{for } j \text{ s.t. } \theta_j = \theta_i \\ t_j = s_j = 0 & \text{for } j \text{ s.t. } \theta_j \neq \theta_i \end{array} \right\}, \quad (\text{C.35})$$

so the projection entries are

$$\begin{aligned} \begin{bmatrix} \xi_i(\mathbf{w}) \\ \eta_i(\mathbf{w}) \end{bmatrix} &= (\theta_i \varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}}) \tilde{\mathbf{q}}_{i,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}} - \mathbf{u}_i^H \hat{\mathbf{u}}_{i,\mathbf{w}}) \frac{\theta_i}{\theta_i^2 \varphi_{1,\mathbf{w}}^2(\rho_{i,\mathbf{w}}) + 1} \begin{bmatrix} \theta_i \varphi_{1,\mathbf{w}}(\rho_{i,\mathbf{w}}) \\ -1 \end{bmatrix}, \\ \forall j: j \neq i \quad \begin{bmatrix} \xi_j(\mathbf{w}) \\ \eta_j(\mathbf{w}) \end{bmatrix} &= \theta_j \begin{bmatrix} \tilde{\mathbf{q}}_{j,\mathbf{w}}^H \hat{\mathbf{q}}_{i,\mathbf{w}} \\ \mathbf{u}_j^H \hat{\mathbf{u}}_{i,\mathbf{w}} \end{bmatrix}, \end{aligned} \quad (\text{C.36})$$

Applying (C.34) to (C.36) yields

$$\sum_{j:j \neq i} |\mathbf{u}_j^H \hat{\mathbf{u}}_{i,w}|^2 + |\tilde{\mathbf{q}}_{j,w}^H \hat{\mathbf{q}}_{i,w}|^2 \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (\text{C.37})$$

$$\left| \sqrt{\frac{\varphi_{1,w}(\rho_{i,w})}{\varphi_{2,w}(\rho_{i,w})}} \tilde{\mathbf{q}}_{i,w}^H \hat{\mathbf{q}}_{i,w} - \mathbf{u}_i^H \hat{\mathbf{u}}_{i,w} \right|^2 \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (\text{C.38})$$

where we recall that $D_w(\rho_{i,w}) = \varphi_{1,w}(\rho_{i,w})\varphi_{2,w}(\rho_{i,w}) = 1/\theta_i^2$ and note that $\varphi_{1,w}(\rho_{i,w})$ and $\varphi_{2,w}(\rho_{i,w})$ are uniformly upper and lower bounded with respect to $\mathbf{w} \in \Delta_L$.

Now, by (C.37) and the bounds (C.32), it follows that

$$\sum_{j_1, j_2: (j_1, j_2) \neq (i, i)} \theta_{j_1} \theta_{j_2} (\tilde{\mathbf{q}}_{j_1,w}^H \hat{\mathbf{q}}_{i,w}) (\tilde{\mathbf{q}}_{j_2,w}^H \hat{\mathbf{q}}_{i,w})^* \mathbf{u}_{j_2}^H \hat{\theta}_{i,w}^2 \mathbf{\Gamma}_w^2 \mathbf{u}_{j_1} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (\text{C.39a})$$

$$\sum_{j_1, j_2: (j_1, j_2) \neq (i, i)} \theta_{j_1} \theta_{j_2} (\mathbf{u}_{j_1}^H \hat{\mathbf{u}}_{i,w}) (\mathbf{u}_{j_2}^H \hat{\mathbf{u}}_{i,w})^* \tilde{\mathbf{q}}_{j_2,w}^H \tilde{\mathbf{E}}_w \mathbf{\Gamma}_w^2 \tilde{\mathbf{E}}_w \tilde{\mathbf{q}}_{j_1,w} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0, \quad (\text{C.39b})$$

so it remains to analyze the summands of χ_1 and χ_2 with $(j_1, j_2) = (i, i)$. For this, note that

$$\mathbf{u}_i^H \hat{\theta}_{i,w}^2 \mathbf{\Gamma}_w^2 \mathbf{u}_i = \left(+\frac{1}{2\zeta} - \frac{1}{2} \frac{\partial}{\partial \zeta} \right) [\mathbf{M}(\mathbf{w}, \zeta)]_{i,i} \bigg|_{\zeta=\hat{\theta}_{i,w}}, \quad (\text{C.40a})$$

$$\tilde{\mathbf{q}}_{i,w}^H \tilde{\mathbf{E}}_w \mathbf{\Gamma}_w^2 \tilde{\mathbf{E}}_w \tilde{\mathbf{q}}_{i,w} = \left(-\frac{1}{2\zeta} - \frac{1}{2} \frac{\partial}{\partial \zeta} \right) [\mathbf{M}(\mathbf{w}, \zeta)]_{k+i, k+i} \bigg|_{\zeta=\hat{\theta}_{i,w}}, \quad (\text{C.40b})$$

which can be verified by substituting (6.8) and (C.26), expanding and simplifying.

Now recall that (6.11) yields, for $i = 1, \dots, k$,

$$[\mathbf{M}(\mathbf{w}, \zeta)]_{i,i} \xrightarrow[(\mathbf{w}, \zeta) \in \Omega(\tau)]{\text{a.s.}} \varphi_{1,w}(\zeta), \quad [\mathbf{M}(\mathbf{w}, \zeta)]_{k+i, k+i} \xrightarrow[(\mathbf{w}, \zeta) \in \Omega(\tau)]{\text{a.s.}} \varphi_{2,w}(\zeta). \quad (\text{C.41})$$

Moreover, almost surely, eventually, for all $(\mathbf{w}, \zeta) \in \Omega(\tau)$ these are both holomorphic functions with respect to ζ , because eventually, for all $\mathbf{w} \in \Delta_L$, $\|\tilde{\mathbf{E}}_w\|_{\text{op}} < b_w + \tau$. Thus, by a standard application of Cauchy's integral formula, their derivatives converge uniformly on any compact subset. Namely, for any compact subset $\mathcal{C} \subset \Omega(\tau)$, for all $i = 1, \dots, k$,

$$\frac{\partial}{\partial \zeta} [\mathbf{M}(\mathbf{w}, \zeta)]_{i,i} \xrightarrow[(\mathbf{w}, \zeta) \in \mathcal{C}]{\text{a.s.}} \varphi'_{1,w}(\zeta), \quad \frac{\partial}{\partial \zeta} [\mathbf{M}(\mathbf{w}, \zeta)]_{k+i, k+i} \xrightarrow[(\mathbf{w}, \zeta) \in \mathcal{C}]{\text{a.s.}} \varphi'_{2,w}(\zeta). \quad (\text{C.42})$$

Thus, applying (C.41) and (C.42) to (C.40) yields

$$\mathbf{u}_i^H \hat{\theta}_{i,w}^2 \mathbf{\Gamma}_w^2 \mathbf{u}_i \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} +\frac{\varphi_{1,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{1,w}(\rho_{i,w})}{2}, \quad (\text{C.43a})$$

$$\tilde{\mathbf{q}}_{i,w}^H \tilde{\mathbf{E}}_w \mathbf{\Gamma}_w^2 \tilde{\mathbf{E}}_w \tilde{\mathbf{q}}_{i,w} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} -\frac{\varphi_{2,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{2,w}(\rho_{i,w})}{2}, \quad (\text{C.43b})$$

where we also used that for any $\mathbf{w} \in \mathcal{W}_{>}(\nu)$, almost surely, eventually, $\mathbf{M}(\mathbf{w}, \cdot)$, $\bar{\mathbf{M}}(\mathbf{w}, \cdot)$ and their derivatives are Lipschitz functions for $\zeta > b_w + \nu/2$ with Lipschitz constants uniformly bounded across $\mathbf{w} \in \mathcal{W}_{>}(\nu)$.

Finally, combining (C.38), (C.39a), and (C.43a) yields (6.16a). Similarly, combining (C.39b) and (C.43b) yields (6.16b). Equation (6.16c) follows from a similar argument. More precisely, combining (6.8), (6.11), and (C.26) yields that

$$\mathbf{u}_{j_2}^H \hat{\theta}_{i,w} \mathbf{\Gamma}_w^2 \tilde{\mathbf{E}}_w \tilde{\mathbf{q}}_{j_1,w} = -\frac{1}{2} \frac{\partial}{\partial \zeta} [\mathbf{M}(\mathbf{w}, \zeta)]_{j_2, k+j_1} \bigg|_{\zeta=\hat{\theta}_{i,w}} \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} 0.$$

Combining this with the bounds (C.32) yields (6.16c).

C.8 Detailed derivation of (6.17)

Note that

$$\begin{aligned} \left| 1 + \frac{\theta_i^2 D'_w(\rho_{i,w})}{2\varphi_{1,w}(\rho_{i,w})} |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \right| &= \left| 1 - \theta_i^2 \left[+ \frac{\varphi_{1,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{1,w}(\rho_{i,w})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \frac{\varphi_{2,w}(\rho_{i,w})}{\varphi_{1,w}(\rho_{i,w})} \right. \\ &\quad \left. - \theta_i^2 \left[- \frac{\varphi_{2,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{2,w}(\rho_{i,w})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \right| \\ &\leq |1 - \chi_1(\mathbf{w}) - \chi_2(\mathbf{w}) - 2 \operatorname{Re} \chi_3(\mathbf{w})| \end{aligned} \quad (\text{C.44a})$$

$$+ \left| \chi_1(\mathbf{w}) - \theta_i^2 \left[+ \frac{\varphi_{1,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{1,w}(\rho_{i,w})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \frac{\varphi_{2,w}(\rho_{i,w})}{\varphi_{1,w}(\rho_{i,w})} \right| \quad (\text{C.44b})$$

$$+ \left| \chi_2(\mathbf{w}) - \theta_i^2 \left[- \frac{\varphi_{2,w}(\rho_{i,w})}{2\rho_{i,w}} - \frac{\varphi'_{2,w}(\rho_{i,w})}{2} \right] |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \right| \quad (\text{C.44c})$$

$$+ |2 \operatorname{Re} \chi_3(\mathbf{w})|, \quad (\text{C.44d})$$

where almost surely

- (C.44a) is equal to zero for all $\mathbf{w} \in \mathcal{W}_{>}(\nu)$ eventually by (6.14b),
- (C.44b) converges to zero uniformly over $\mathbf{w} \in \mathcal{W}_{>}(\nu)$ by (6.16a),
- (C.44c) converges to zero uniformly over $\mathbf{w} \in \mathcal{W}_{>}(\nu)$ by (6.16b),
- (C.44d) converges to zero uniformly over $\mathbf{w} \in \mathcal{W}_{>}(\nu)$ by (6.16c),

and we used the fact that $D'_w(\zeta) = \varphi'_{1,w}(\zeta)\varphi_{2,w}(\zeta) + \varphi_{1,w}(\zeta)\varphi'_{2,w}(\zeta)$. Thus,

$$\frac{\theta_i^2 D'_w(\rho_{i,w})}{2\varphi_{1,w}(\rho_{i,w})} |\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \xrightarrow[\mathbf{w} \in \mathcal{W}_{>}(\nu)]{\text{a.s.}} -1,$$

and (6.17) follows since $2\varphi_{1,w}(\rho_{i,w})/(\theta_i^2 D'_w(\rho_{i,w}))$ is bounded over $\mathbf{w} \in \mathcal{W}_{>}(\nu)$.

C.9 Detailed derivation of (6.18)

Consider

$$\mathbf{R}(\mathbf{w}, \zeta) := \left(\zeta \mathbf{I}_{d+n} - \begin{bmatrix} \tilde{\mathbf{Y}}_w & \tilde{\mathbf{Y}}_w \end{bmatrix} \right)^{-1} = \begin{bmatrix} \zeta \mathbf{R}_1(\mathbf{w}, \zeta) & \mathbf{R}_1(\mathbf{w}, \zeta) \tilde{\mathbf{Y}}_w \\ \tilde{\mathbf{Y}}_w^H \mathbf{R}_1(\mathbf{w}, \zeta) & \zeta \mathbf{R}_2(\mathbf{w}, \zeta) \end{bmatrix}, \quad (\text{C.45})$$

where $\mathbf{R}_1(\mathbf{w}, \zeta) := (\zeta^2 \mathbf{I}_d - \tilde{\mathbf{Y}}_w \tilde{\mathbf{Y}}_w^H)^{-1}$ and $\mathbf{R}_2(\mathbf{w}, \zeta) := (\zeta^2 \mathbf{I}_n - \tilde{\mathbf{Y}}_w^H \tilde{\mathbf{Y}}_w)^{-1}$. With the spectral decomposition of $\tilde{\mathbf{Y}}_w \tilde{\mathbf{Y}}_w^H$, it is easy to see that for $i = 1, \dots, k$,

$$|\mathbf{u}_i^H \hat{\mathbf{u}}_{i,w}|^2 \leq -\nu \cdot \operatorname{Im} \{ \mathbf{u}_i^H \mathbf{R}_1(\mathbf{w}, \zeta_{i,w}) \mathbf{u}_i \} = -\nu \cdot \operatorname{Im} \left\{ \zeta_{i,w}^{-1} \begin{bmatrix} \mathbf{u}_i \\ \mathbf{0}_n \end{bmatrix}^H \mathbf{R}(\mathbf{w}, \zeta_{i,w}) \begin{bmatrix} \mathbf{u}_i \\ \mathbf{0}_n \end{bmatrix} \right\}. \quad (\text{C.46})$$

Noting that

$$\mathbf{R}(\mathbf{w}, \zeta) = \left([G(\mathbf{w}, \zeta)]^{-1} - \begin{bmatrix} \mathbf{U} & \tilde{\mathbf{Q}}_w \end{bmatrix} \begin{bmatrix} \boldsymbol{\Theta} & \boldsymbol{\Theta} \end{bmatrix} \begin{bmatrix} \mathbf{U} & \tilde{\mathbf{Q}}_w \end{bmatrix}^H \right)^{-1},$$

and applying the Sherman-Morrison-Woodbury formula [23, Equation (0.7.4.1)], we get that

$$\begin{bmatrix} \mathbf{u}_i \\ \mathbf{0}_n \end{bmatrix}^H \mathbf{R}(\mathbf{w}, \zeta_{i,w}) \begin{bmatrix} \mathbf{u}_i \\ \mathbf{0}_n \end{bmatrix} = [\tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,w}) - \tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,w}) [\mathbf{M}(\mathbf{w}, \zeta_{i,w})]^{-1} \tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,w})]_{ii}. \quad (\text{C.47})$$

Combining (C.46) and (C.47) yields (6.18).

C.10 Detailed derivation of (6.20)

We derive the three bounds of (6.20) one by one.

1. The first bound is straightforward. Note that for all $\mathbf{w} \in \Delta_L$, $b_{\mathbf{w}} \geq \sqrt{\min_{\ell} p_{\ell} v_{\ell}}$ since $\|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}} \xrightarrow{\text{a.s.}} b_{\mathbf{w}}$ and

$$\|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{op}}^2 \geq \frac{1}{d} \|\tilde{\mathbf{E}}_{\mathbf{w}}\|_{\text{F}}^2 = \sum_{\ell=1}^L \frac{w_{\ell}}{dn} \|\mathbf{E}_{\ell}\|_{\text{F}}^2 \xrightarrow{\text{a.s.}} \sum_{\ell=1}^L p_{\ell} w_{\ell} v_{\ell} \geq C_1,$$

where $C_1 := \min_{\ell} p_{\ell} v_{\ell} > 0$, and the convergence follows from the law of large numbers. Thus, it follows from (6.13) that almost surely, eventually, for all $\mathbf{w} \in \Delta_L \supseteq \mathcal{W}_{\leq}(\nu)$,

$$|\zeta_{i,\mathbf{w}}^{-1}| \leq |\hat{\theta}_{i,\mathbf{w}}^{-1}| < 2 \cdot \bar{\theta}_{i,\mathbf{w}}^{-1} \leq 2 \cdot b_{\mathbf{w}}^{-1} \leq \tilde{C}_1, \quad (\text{C.48})$$

where $\tilde{C}_1 := 2/C_1$ does not depend on ν or $\mathbf{w}$.

2. For the second bound, observe that for $\delta = 1/2$ and τ sufficiently large, it follows from (6.13) that almost surely, eventually, for all $\mathbf{w} \in \mathcal{W}_{\leq}(\nu)$, $\zeta_{i,\mathbf{w}} \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa)$. Thus, using (C.21) yields that almost surely, eventually,

$$\sup_{\mathbf{w} \in \Delta_L} \|\mathbf{\Pi}(\mathbf{w}, \zeta_{i,\mathbf{w}})\|_{\text{op}} \leq C_2$$

where C_2 does not depend on ν or $\mathbf{w}$. Next using (C.2), yields that almost surely, eventually,

$$\begin{aligned} \|\tilde{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}})\|_{\text{op}} &\leq \|\mathbf{\Pi}(\mathbf{w}, \zeta_{i,\mathbf{w}})\|_{\text{op}} \left\| \begin{bmatrix} \mathbf{U} \\ \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} \right\|_{\text{op}}^2 \\ &\quad + \left\| \begin{bmatrix} \mathbf{U} \\ \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix}^{\text{H}} [\mathbf{G}(\mathbf{w}, \zeta_{i,\mathbf{w}}) - \mathbf{\Pi}(\mathbf{w}, \zeta_{i,\mathbf{w}})] \begin{bmatrix} \mathbf{U} \\ \tilde{\mathbf{Q}}_{\mathbf{w}} \end{bmatrix} \right\|_{\text{op}} \\ &\leq \tilde{C}_2, \end{aligned} \quad (\text{C.49})$$

where $\tilde{C}_2$ does not depend on ν or $\mathbf{w}$.

3. For the third bound, observe that

$$\|\bar{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}})^{-1}\|_{\text{op}} \leq C_3 (\text{Im } m_{\text{MP}}(\mathbf{w}, \zeta_{i,\mathbf{w}}))^{-1} \quad (\text{C.50})$$

where C_3 does not depend on ν or $\mathbf{w}$. It has been shown in [13, Lemma 3.6] that

$$\text{Im } m_{\text{MP}}(\mathbf{w}, \zeta_{i,\mathbf{w}}) \geq \frac{C_4 \nu}{\sqrt{|\hat{\theta}_{i,\mathbf{w}} - b_{\mathbf{w}}| + \nu}} \quad (\text{C.51})$$

where C_4 does not depend on ν or $\mathbf{w}$. Moreover, note that the uniform convergence (6.11) extends to the domain $\Omega_{\leq}(\tau) := \{(\mathbf{w}, \zeta) \in \Delta_L \times \mathbb{C} : \zeta \in S_{\text{edge}}(\mathbf{w}, \tau, \delta; \kappa)\}$. Combining this with (C.50) and (C.51) yields that almost surely, eventually,

$$\begin{aligned} \|\mathbf{M}(\mathbf{w}, \zeta_{i,\mathbf{w}})^{-1}\|_{\text{op}} &\leq 2 \cdot \|\bar{\mathbf{M}}(\mathbf{w}, \zeta_{i,\mathbf{w}})^{-1}\|_{\text{op}} \leq 2 \cdot \frac{C_3}{C_4 \nu} \sqrt{|\hat{\theta}_{i,\mathbf{w}} - b_{\mathbf{w}}| + \nu} \\ &\leq 2 \cdot \frac{C_3}{C_4 \nu} \sqrt{2\nu + \nu} \leq \tilde{C}_3 \nu^{-1/2}, \end{aligned} \quad (\text{C.52})$$

where $\tilde{C}_3 := 4 \cdot C_3/C_4$ does not depend on ν or $\mathbf{w}$.

C.11 Detailed derivation of (6.24)

Analogous to [19, Section 5.3], observe that $\psi_{\mathbf{w}}$ has the following properties for all $\mathbf{w} \in \Delta_L$:

(a) $0 = Q_{\mathbf{w}}(\psi_{\mathbf{w}}(\zeta), \zeta)$ for all $\zeta > b_{\mathbf{w}}$ where

$$Q_{\mathbf{w}}(s, \zeta) := \frac{c\zeta^2}{s^2} + \frac{c-1}{s} - \sum_{\ell=1}^L \frac{c_{\ell}}{s - w_{\ell}v_{\ell}}, \quad (\text{C.53})$$

(b) $\max_{\ell}(w_{\ell}v_{\ell}) < \psi_{\mathbf{w}}(\zeta) < c\zeta^2$,

(c) $0 < \psi_{\mathbf{w}}(b_{\mathbf{w}}^+) < \infty$ and $\psi'_{\mathbf{w}}(b_{\mathbf{w}}^+) = \infty$.

Expressing $D_{\mathbf{w}}$ in terms of $\psi_{\mathbf{w}}$ yields

$$D_{\mathbf{w}}(\zeta) = \varphi_{1,\mathbf{w}}(\zeta) \sum_{\ell=1}^L \frac{p_{\ell}w_{\ell}}{\zeta - w_{\ell}v_{\ell}\varphi_{1,\mathbf{w}}(\zeta)/c} = \sum_{\ell=1}^L \frac{c_{\ell}w_{\ell}}{\psi_{\mathbf{w}}(\zeta) - w_{\ell}v_{\ell}} = \frac{1 - B_{i,\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}{\theta_i^2}, \quad (\text{C.54})$$

$$\frac{D'_{\mathbf{w}}(\zeta)}{\zeta} = -\frac{c\psi'_{\mathbf{w}}(\zeta)}{\zeta} \sum_{\ell=1}^L \frac{p_{\ell}w_{\ell}}{(\psi_{\mathbf{w}}(\zeta) - w_{\ell}v_{\ell})^2} = -\frac{2c}{\theta_i^2} \frac{B'_{i,\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}{A_{\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}. \quad (\text{C.55})$$

The second equality in (C.55) follows analogously to [19, Section 5.4] by deriving the identity

$$\psi'_{\mathbf{w}}(\zeta) = \frac{2c\zeta}{A_{\mathbf{w}}(\psi_{\mathbf{w}}(\zeta))}, \quad (\text{C.56})$$

from Property (a) then simplifying.

Rearranging (C.56) then applying Property (c) yields

$$A_{\mathbf{w}}(\psi_{\mathbf{w}}(b_{\mathbf{w}}^+)) = \frac{2cb_{\mathbf{w}}}{\psi'_{\mathbf{w}}(b_{\mathbf{w}}^+)} = 0, \quad (\text{C.57})$$

so $\psi_{\mathbf{w}}(b_{\mathbf{w}}^+)$ is a root of $A_{\mathbf{w}}$. If $\theta_i > \tilde{\theta}_{\mathbf{w}}$, then $\rho_{i,\mathbf{w}} = D_{\mathbf{w}}^{-1}(1/\theta_i^2)$ and rearranging (C.54) yields

$$B_{i,\mathbf{w}}(\psi_{\mathbf{w}}(\rho_{i,\mathbf{w}})) = 1 - \theta_i^2 D_{\mathbf{w}}(\rho_{i,\mathbf{w}}) = 0, \quad (\text{C.58})$$

so $\psi_{\mathbf{w}}(\rho_{i,\mathbf{w}})$ is a root of $B_{i,\mathbf{w}}$. Recall that $\psi_{\mathbf{w}}(b_{\mathbf{w}}^+), \psi_{\mathbf{w}}(\rho_{i,\mathbf{w}}) \geq \max_{\ell}(w_{\ell}v_{\ell})$ by Property (b), and observe that both $A_{\mathbf{w}}(x)$ and $B_{i,\mathbf{w}}(x)$ monotonically increase for $x > \max_{\ell}(w_{\ell}v_{\ell})$ from negative infinity to one. Thus, each has exactly one real root larger than $\max_{\ell}(w_{\ell}v_{\ell})$, i.e., its largest real root, and so $\psi_{\mathbf{w}}(b_{\mathbf{w}}^+) = \alpha_{\mathbf{w}}$ and $\psi_{\mathbf{w}}(\rho_{i,\mathbf{w}}) = \beta_{i,\mathbf{w}}$ when $\theta_i^2 > \tilde{\theta}_{\mathbf{w}}^2$, where $\alpha_{\mathbf{w}}$ and $\beta_{i,\mathbf{w}}$ are the largest real roots of $A_{\mathbf{w}}$ and $B_{i,\mathbf{w}}$, respectively.

D Detailed verification of consistency for estimators in Example 7.2

The simple noise variance estimator $\hat{\mathbf{v}}(\mathbf{Y})$ is well-known; its consistency follows straightforwardly by noting that

$$\frac{\|\mathbf{F}\mathbf{Z}_{\ell}\|_{\mathbb{F}}^2}{dn_{\ell}} \leq \|\mathbf{F}\|_{\text{op}}^2 \frac{\|\mathbf{Z}_{\ell}\|_{\mathbb{F}}^2}{dn_{\ell}} = \lambda_1 \frac{\|\mathbf{Z}_{\ell}\|_{\mathbb{F}}^2}{kn_{\ell}} \frac{k}{d} \xrightarrow{\text{a.s.}} 0, \quad \frac{\|\mathbf{E}_{\ell}\|_{\mathbb{F}}^2}{dn_{\ell}} \xrightarrow{\text{a.s.}} v_{\ell},$$

by the law of large numbers, so $\hat{v}_{\ell}(\mathbf{Y}) = \|\mathbf{Y}_{\ell}\|_{\mathbb{F}}^2/(dn_{\ell}) = \|\mathbf{F}\mathbf{Z}_{\ell} + \mathbf{E}_{\ell}\|_{\mathbb{F}}^2/(dn_{\ell}) \xrightarrow{\text{a.s.}} v_{\ell}$.

The signal variance estimator $\hat{\lambda}_i(\mathbf{Y})$ is essentially a bias-corrected version of the inverse noise variance weighted data eigenvalue $\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \mathbf{v})$. As is well known, such data eigenvalues are typically upwardly biased

for high-dimensional data; see, e.g., [19, 26, 32, 37, 40]. The shrinkage Ξ simply corrects for the asymptotic bias.

In particular, note that it follows from (6.13) in the proof of Lemma 6.2 that

$$\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \hat{\mathbf{v}}(\mathbf{Y})) = \hat{\theta}_{i, \hat{\mathbf{w}}}^2(\mathbf{Y}) \xrightarrow{\text{a.s.}} \begin{cases} \rho_{i, \bar{\mathbf{w}}}^2 & \text{if } \theta_i > \tilde{\theta}_{\bar{\mathbf{w}}}, \\ b_{\bar{\mathbf{w}}}^2 & \text{otherwise,} \end{cases}$$

where $\hat{w}_\ell^{(\text{inv})}(\mathbf{Y}) := [\sum_{\ell=1}^L p_\ell / \hat{v}_\ell(\mathbf{Y})]^{-1} / \hat{v}_\ell(\mathbf{Y}) \xrightarrow{\text{a.s.}} \bar{w}_\ell^{(\text{inv})} := \bar{v} / v_\ell$ since $\hat{\mathbf{v}}(\mathbf{Y}) \xrightarrow{\text{a.s.}} \mathbf{v}$. To obtain the algebraic form, recall from subsection 6.1.5 that

$$\psi_{\bar{\mathbf{w}}^{(\text{inv})}}(b_{\bar{\mathbf{w}}^{(\text{inv})}}^+) = \alpha_{\bar{\mathbf{w}}^{(\text{inv})}}, \quad \psi_{\bar{\mathbf{w}}^{(\text{inv})}}(\rho_{i, \bar{\mathbf{w}}^{(\text{inv})}}) = \beta_{i, \bar{\mathbf{w}}^{(\text{inv})}}, \quad \theta_i > \tilde{\theta}_{\bar{\mathbf{w}}^{(\text{inv})}} \Leftrightarrow \alpha_{\bar{\mathbf{w}}^{(\text{inv})}} < \beta_{i, \bar{\mathbf{w}}^{(\text{inv})}},$$

note that $\alpha_{\bar{\mathbf{w}}^{(\text{inv})}} = \bar{v} + \bar{v}\sqrt{c}$ and $\beta_{i, \bar{\mathbf{w}}^{(\text{inv})}} = \bar{v} + c\lambda_i$, and apply Property (a) from Appendix C.11 to invert $\psi_{\bar{\mathbf{w}}^{(\text{inv})}}$ and conclude that when $c(\lambda_i / \bar{v})^2 > 1$

$$\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \hat{\mathbf{v}}(\mathbf{Y})) \xrightarrow{\text{a.s.}} \rho_{i, \bar{\mathbf{w}}^{(\text{inv})}}^2 = \frac{(\lambda_i + \bar{v}/c)(\lambda_i + \bar{v})}{\lambda_i}.$$

This gives the asymptotic bias for inverse noise variance weighted data eigenvalues; a closely related result was derived in [32]. It follows immediately that

$$\hat{\lambda}_i(\mathbf{Y}) = \Xi \left(\hat{\lambda}_i^{(\text{inv})}(\mathbf{Y}; \hat{\mathbf{v}}(\mathbf{Y})); \left[\sum_{\ell=1}^L \frac{p_\ell}{\hat{v}_\ell(\mathbf{Y})} \right]^{-1} \right) \xrightarrow{\text{a.s.}} \Xi \left(\frac{(\lambda_i + \bar{v}/c)(\lambda_i + \bar{v})}{\lambda_i}; \bar{v} \right) = \lambda_i,$$

and so $\hat{\lambda}_i(\mathbf{Y})$ is consistent.

E Preprocessing details for SDSS data

This section describes the details of the subset selected and the preprocessing performed on the SDSS data for illustrating optimally weighted PCA in section 8. In particular, the dataset was formed via the following steps:

1. Collect the spectra from DR16Q for which
 - SURVEY = "eboss",
 - PLATEQUALITY = "good",
 - redshift: $2.0 < Z < 2.1$,
 - BAL_PROB < 0.2 ,
 - the measured rest frame wavelengths cover the range of 1480–1620 without any missing entries (i.e., no entries in the range with IVAR = 0).
2. Form data $\mathbf{y}_j \in \mathbb{R}^d$ and variance profile $\mathbf{p}_j \in \mathbb{R}^d$ vectors for each collected spectrum via linear interpolation of FLUX and $1 \oslash \text{IVAR}$ on the grid of rest frame wavelengths LAMREST = (1480, 1480.5, ..., 1620) $\in \mathbb{R}^d$.
3. Center each spectrum $\mathbf{y}_j \leftarrow \mathbf{y}_j - (1/d)\mathbf{1}_{d \times d}\mathbf{y}_j$.
4. Normalize each spectrum so that its mean flux for rest frame wavelengths in 1525–1575 is ± 1 . Namely,
 - (a) $\sigma_j \leftarrow |\text{mean}(\mathbf{y}_j(1525 < \text{LAMREST} < 1575))|$,
 - (b) $\mathbf{y}_j \leftarrow \mathbf{y}_j / \sigma_j$,
 - (c) $\mathbf{p}_j \leftarrow \mathbf{p}_j / \sigma_j^2$.
5. Remove spectra with $\min(\mathbf{p}_j) / \max(\mathbf{p}_j) \leq 0.4$ (i.e., keep only roughly homogeneous variance profiles).
6. Compute average variances $v_j \leftarrow \text{mean}(\mathbf{p}_j)$ for each sample.