

UC Santa Barbara

UC Santa Barbara Electronic Theses and Dissertations

Title

James--Stein Type Shrinkage for Eigenvectors of Random Matrices

Permalink

<https://escholarship.org/uc/item/9x65h82s>

ISBN

9798276025056

Author

Jin, Zhuoli

Publication Date

2025-12-12

Peer reviewed|Thesis/dissertation

University of California
Santa Barbara

James–Stein Type Shrinkage for Eigenvectors of Random Matrices

A dissertation submitted in partial satisfaction
of the requirements for the degree

Doctor of Philosophy
in
Statistics and Applied Probability

by

Zhuoli Jin

Committee in charge:

Professor Alex Shkolnik, Chair
Professor Guo Yu
Professor Sui Tang

December 2025

The Dissertation of Zhuoli Jin is approved.

Professor Guo Yu

Professor Sui Tang

Professor Alex Shkolnik, Committee Chair

December 2025

James–Stein Type Shrinkage for Eigenvectors of Random Matrices

Copyright © 2025

by

Zhuoli Jin

To the memory of Mr. Jianxin Ma,
my middle school vice principal,
who first inspired my love for mathematics.
May he rest in peace.

Acknowledgements

I would like to express my deepest gratitude to my advisor, Professor Alex Shkolnik, for his invaluable guidance, patience, and insight throughout my Ph.D. studies. His mentorship has shaped the way I think about mathematics and research, and his high standards and encouragement have continuously inspired me to pursue clarity and rigor in my work. I am deeply grateful for his support, from the earliest ideas to the completion of this dissertation.

I would also like to thank my committee members, Professor Sui Tang and Professor Guo Yu, for their time, feedback, and support during my dissertation defense.

I am especially thankful to Mengye Liu, my long-time friend and colleague who has shared both my undergraduate and Ph.D. journeys in statistics. Her companionship and encouragement have been a source of strength and motivation through every stage of this process. I am also grateful to my fellow graduate students and friends in the Department of Statistics and Applied Probability at UCSB for their collaboration, insightful discussions, and friendship.

My deepest gratitude goes to my parents, whose unwavering love, understanding, and belief in me have sustained me throughout this long path. I am equally grateful to Jessi Liu for her kindness, patience, and steady encouragement, which have made this journey both meaningful and fulfilling.

Finally, I would like to thank Shanghai Restaurant, where I spent many days studying, thinking, and working through the details of this research. The quiet background noise, familiar atmosphere, and a bowl of Dandan noodles often provided just the right environment for clarity and focus.

Curriculum Vitæ

Zhuoli Jin

Education

2025	Ph.D. in Statistics and Applied Probability (Expected), University of California, Santa Barbara.
2018	M.A. in Statistics, Wake Forest University
2016	B.S. in Mathematics and Applied Mathematics, Xi'an Jiaotong University

Publications

Sun, Chengkun, Jinqian Pan, Zhuoli Jin, Russell Stevens Terry, Jiang Bian, and Jie Xu. "Beyond Skip Connection: Pooling and Unpooling Design for Elimination Singularities." AAAI 2025 Conference

Jin, Zhuoli, and Robert J. Erhardt. "Incorporating climate change projections into risk measures of index-based insurance." North American Actuarial Journal 24, no. 4 (2020): 611-625.

Abstract

James–Stein Type Shrinkage for Eigenvectors of Random Matrices

by

Zhuoli Jin

This thesis investigates the asymptotic behavior and bias properties of spike eigenvector estimators in high-dimensional statistical models. Motivated by challenges in estimating latent structures from noisy, high-dimensional data, it focuses on how shrinkage and model specification shape the geometry and accuracy of estimated eigenvectors across different asymptotic regimes.

Chapter 1 provides the theoretical background and mathematical tools, introducing connections between random matrix theory, spectral geometry, and shrinkage estimation. Chapter 2 develops and analyses a James–Stein type shrinkage framework for eigenvector estimation under rank misspecification in the factor model $Y = BX + \varepsilon$. In high-dimensional settings, empirical eigenvectors are inherently inconsistent even under correct model specification. The analysis characterizes the additional optimization bias that arises when the estimated rank K differs from true latent dimension q . Theoretical results show that when $K > q$, the shrunk eigenvectors achieve asymptotically unbiased alignment with the population eigenspace, but for $K < q$, the bias persists.

Chapter 3 extends the study to the proportional asymptotic regime where both the dimension p and sample size n grow with $p/n \rightarrow \kappa > 0$. Building on the spectral theory of Wigner matrices, it examines the semicircle law for eigenvalues and the misalignment between empirical and population eigenvectors in high-dimensional models. We further derive a James–Stein correction for spiked Wigner matrices and mitigated the spectral distortion introduced by high-dimensional noise. As a result, the adjusted estimators

achieve improved alignment with the true signal directions in the spiked Wigner model. Lastly, we demonstrate the advantage of James–Stein correction over their empirical counterparts in a stochastic block model with node preferences that we propose.

Together, these findings offer a unified theoretical perspective on eigenvector inconsistency, shrinkage correction, and asymptotic geometry in high-dimensional inference.

Contents

Curriculum Vitae	vi
Abstract	vii
1 Introduction	1
1.1 Motivation	2
1.1.1 Spectral Structure of Random Covariance Matrices	5
1.1.1.1 Noise-Only Covariance Models	5
1.1.1.2 Low-Rank Signal and Spiked Models	7
1.1.2 Stochastic Block Models and the Semicircle Law	16
1.1.3 Summary: Two Spectral Regimes in Random Graph Models . . .	19
1.2 Literature Review	20
1.3 Mathematical Preliminaries	22
1.3.1 Linear Algebra	22
1.3.2 Random Matrix Theory	30
1.3.3 High-Dimensional Geometry	37
1.4 Thesis Overview	41
2 The Quadratic Optimization Bias of Misspecified PCA and Associated James-Stein Estimators.	43
2.1 Literature Review	48
2.2 Motivation	50
2.2.1 Problem Formulation and Theoretical Limits	50
2.2.2 Discrepancy, Optimization Bias and Its Geometric Source	52
2.2.3 Decomposition and Shrinkage Estimator Framework	54
2.2.4 Application	56
2.3 Model Setup and Problem Formulation	56
2.3.1 Model Setup	57
2.3.2 Problem Formulation	59
2.3.3 $\mathcal{H}_\sharp$ Construction	60

2.4	Theoretical Results on Rank Misspecification	
	Bias	62
2.4.1	James-Stein Behavior Under Overspecified rank $K > q$	62
2.4.2	James-Stein Behavior Under Underspecified rank $K < q$	63
2.4.3	A Closer Look: Choose number of factors $K = 1$	65
2.4.4	Equivalence with a James-Stein type estimator	67
2.5	Portfolio Optimization with Misspecified Factor Models: A Simulation Study	69
2.5.1	Simulation Framework and Methodology	69
2.5.2	Empirical Results and Analysis	72
2.5.3	Interpretation and Practical Recommendations	72
2.6	Proofs for Chapter 2	75
2.6.1	Proofs for Section 2.2	75
2.6.2	Proofs for Section 2.3.1	78
2.6.3	Proofs for Section 2.3.3	79
2.6.4	Proofs for Section 2.4.1	80
2.6.5	Proofs for Section 2.4.2	83
2.6.6	Proofs for Section 2.4.3	87
3	A Stochastic Block Model With Preferences and James-Stein Estimation for Spiked Wigner Matrices	93
3.1	Literature Review	95
3.2	Motivation	97
3.3	Problem formulation	98
3.3.1	The spiked Wigner matrix model	100
3.3.2	Spiked eigenvector inconsistency	102
3.4	James-Stein type estimator	104
3.4.1	A concentration of measure result	106
3.5	Numerical Simulation	107
3.5.1	Stochastic block model with preferences	107
3.5.2	A numerical study on preference recovery	111
3.6	Proofs for Chapter 3	113
3.6.1	Proof of Theorem 3.3.1	114
3.6.2	Proofs of Theorems 3.3.4 and 3.4.1	115
3.6.3	Proofs of Section 3.3	118
	Bibliography	121

Chapter 1

Introduction

This chapter provides the empirical and conceptual motivation for this thesis by examining two distinct spectral phenomena in high-dimensional data. We begin with covariance and Gram matrices governed by the Marchenko-Pastur law, where we analyze both pure noise models and spiked populations with low-rank signals. We then explore network data through stochastic block models, where community structure produces eigenvalue spikes separated from a semicircular bulk. These concrete examples demonstrate the prevalence of spiked spectral structures across diverse data types, motivating the need for improved eigenvector estimation methods that will be developed in subsequent chapters.

Our literature review synthesizes key contributions across three asymptotic regimes: classical fixed-dimension theory, random matrix theory with $p/n \rightarrow c$, and the high-dimension, low-sample-size setting ($p \uparrow \infty$, n fixed). This discussion situates our work within the broader theoretical landscape.

We then introduce essential mathematical tools from linear algebra, random matrix theory, and high-dimensional geometry, including eigenvalue perturbation theory, the Marčenko–Pastur law, the geometry of the unit sphere $\mathbb{S}^{p-1}$, and connections between uniform measure and Wiener measure. These concepts form the backbone of our analysis

in later chapters. Finally, we preview the main contributions of Chapters 2 and 3, which examine the asymptotic behavior of eigenvectors under pure noise and the optimization bias arising from rank-misspecified shrinkage estimators, respectively.

1.1 Motivation

With the rapid advancement of data collection technologies, modern statistics frequently encounters high-dimensional scenarios where the number of variables p is similar to or even exceeds the number of observations n . This "large- p , small- n " paradigm has become particularly prominent in fields like genomics, where gene expression studies routinely involve thousands of genes p measured across only dozens of samples n . These difficulties are best understood in the context of a broader phase diagram of multivariate regimes.

Depending on the asymptotic relationship between the number of features p and the sample size n , three distinct statistical worlds emerge:

- **Traditional Statistics** ($p \ll n$): classical multivariate methods such as OLS, PCA, and MLE operate effectively under the assumption of abundant data;
- **Random Matrix Theory (RMT)** ($p/n \rightarrow c \in (0, \infty)$): spectral distortions become predictable and analyzable using tools like the Marčenko–Pastur law and spiked models;
- **High-Dimensional Low-Sample-Size (HDLSS)** ($p \gg n$): classical estimators break down, necessitating regularization, shrinkage, and sparsity-aware techniques.

To visualize this taxonomy, we refer to the well-known "horseshoe" diagram shown in Figure 1.1, which summarizes the spectrum of high-dimensional regimes and the corresponding methodological landscapes.

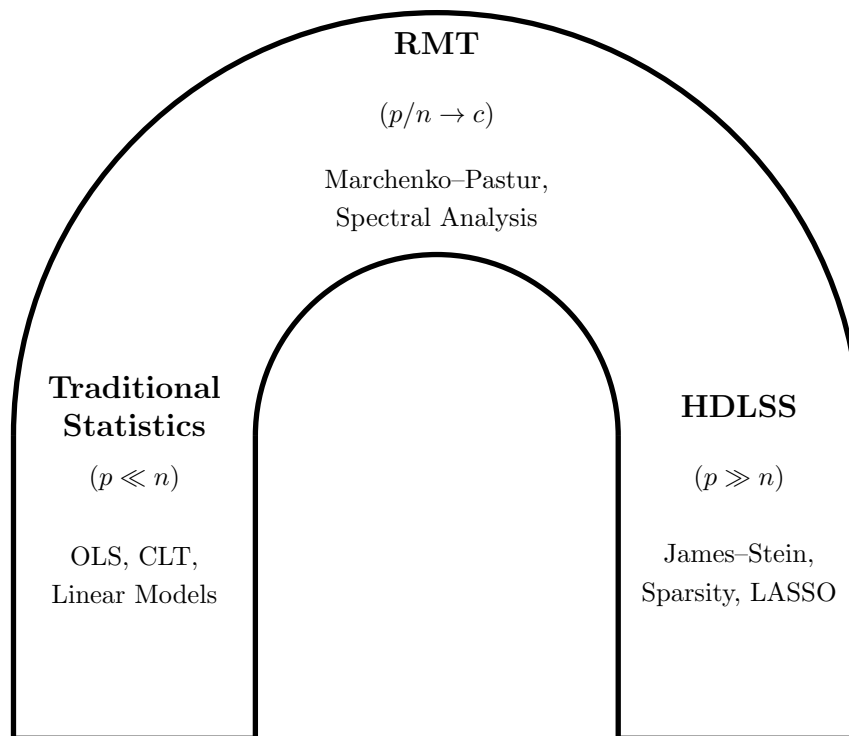


Figure 1.1: The horseshoe diagram illustrating the transition from classical statistics to random matrix theory to high-dimensional regime.

As highlighted by [1], in their work on efficient quadratic regularization for expression arrays, this extreme dimensionality ($p \gg n$) poses significant challenges for traditional covariance estimation methods. Similar challenges arise in finance, signal processing, and machine learning applications. Traditional multivariate statistical methods, designed under the assumption that sample size n far exceeds variable count p , often fail in these modern high-dimensional settings, necessitating the development of specialized regularization techniques and novel statistical approaches.

The sample covariance matrix serves as a cornerstone of multivariate analysis, but its reliability breaks down in high-dimensional settings. When the number of features p approaches or exceeds the sample size n , the sample covariance matrix exhibits severe spectral distortions: its largest eigenvalues become systematically biased upward, while smaller eigenvalues shrink disproportionately. These deviations from the popula-

tion covariance structure lead to misaligned eigenvectors and compromise the validity of subsequent analyses. Such distortions particularly impair dimensionality reduction techniques like PCA, and introduce significant errors in downstream applications including classification, clustering, and risk estimation. This phenomenon has been extensively documented in the statistical literature ([2][3][4][5]), prompting the development of regularized covariance estimators and alternative approaches for high-dimensional data.

To address these challenges, researchers have developed regularized covariance estimators, with James–Stein-type shrinkage methods emerging as a prominent solution. These estimators blend data-driven estimates with structured targets (e.g., identity matrices), effectively reducing estimation risk when p/n is non-negligible. Shrinkage toward low-rank factor structures has proven particularly valuable, offering both interpretability and strong empirical performance [3].

A critical limitation, however, lies in their dependence on correctly specified target structures—particularly the factor dimension q . In practice, q is unknown and must be estimated; when the assumed dimension $K \neq q$, the estimator’s theoretical guarantees often fail. This misspecification not only degrades the covariance estimation accuracy but also propagates errors into optimization-dependent applications.

The problem becomes acute in minimum-variance portfolio optimization, where covariance matrix inversion amplifies estimation errors as mentioned by [6]. While existing work has characterized this issue under idealized conditions [7], the behavior of shrinkage estimators under structural misspecification remains poorly understood, especially regarding eigenvector divergence from population truth.

1.1.1 Spectral Structure of Random Covariance Matrices

1.1.1.1 Noise-Only Covariance Models

We start with a $p \times n$ matrix Y with each element being independently and identically drawn from a mean 0 variance ν^2 distribution. Denote the sample covariance matrix as $S = \frac{1}{n}YY^\top$, and the eigenvalues of S as $\lambda_1(S), \lambda_2(S), \dots$. By law of large numbers, we have $\forall \epsilon > 0$,

$$\lim_{n \rightarrow \infty} P \left(\max_{1 \leq i \leq p} |\lambda_i(S) - \nu^2| > \epsilon \right) = 0, \quad (1.1)$$

that implies, when we fix the number of dimensions, all eigenvalues of sample covariance matrix converges to population variance ν^2 with probability 1. Figure 1.2 displays the histogram of eigenvalues of S , and we are able to see as $n \rightarrow \infty$ with p fixed, the eigenvalues are more and more centered at population variance ν^2 (equals to 1 here). Next, we consider setting, where we set $p \rightarrow \infty$. Unlike the previous traditional case,

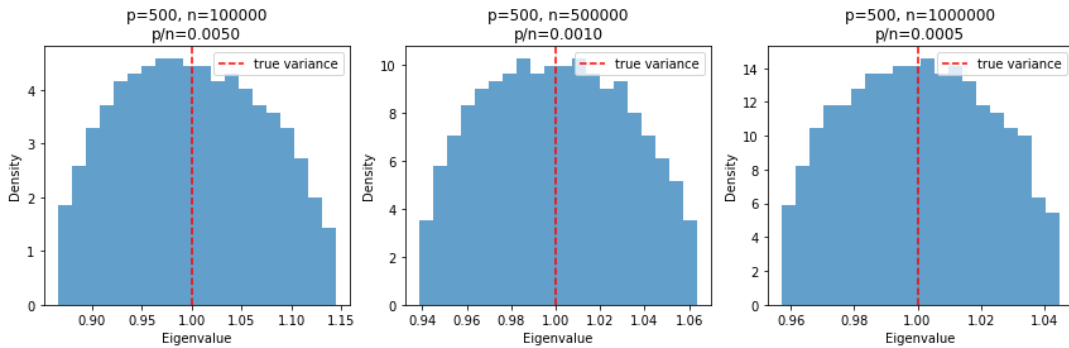
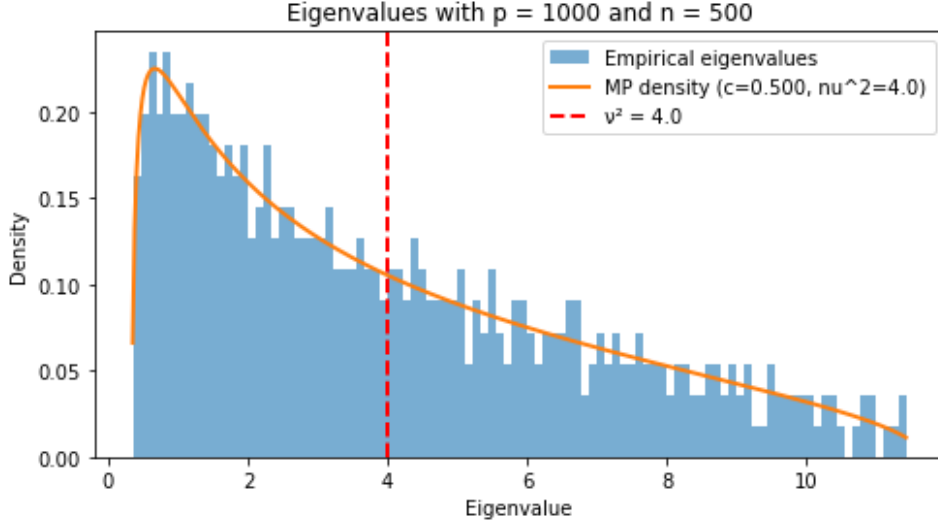


Figure 1.2: Histogram of $\frac{1}{n}YY^\top$

the eigenvalues now behave differently. From Figure 1.3, we can see that the sample eigenvalues are no longer centered at population variance, but follows a right skewed distribution (orange curve) named as Marchenko–Pastur law (MP law).

This empirical spectrum actually reflects a much deeper and universal phenomenon: whenever $S = \frac{1}{n}YY^\top$ is a Gram matrix formed from n i.i.d observations in $\mathbb{R}^p$, its

Figure 1.3: Histogram of $\frac{1}{n}YY^\top$

eigenvalues converge to the Marchenko–Pastur (MP) distribution as long as both $p, n \rightarrow \infty$ with $p/n \rightarrow c \in (0, \infty)$. Moreover, this limiting behavior is independent of normality assumption. The details of MP law will be introduced in Section 1.3.

Intuitively, the MP shape arises because $\frac{1}{n}YY^\top$ aggregates n random outer products of the form $Y_{\cdot j}Y_{\cdot j}^\top$. When p is the same order as n , these rank-one contributions overlap in a highly irregular way, and results in eigenvalues spread over a non-degenerate interval $[(1 - \sqrt{c})^2, (1 + \sqrt{c})^2]$ for some constant c .

On the other hand, the bulk eigenvectors become asymptotically Haar-distributed on the unit sphere. It reflects that for pure noise, all orthogonal directions carry comparable variance, and high-dimensional randomness forces the sample eigenvectors to “diffuse” uniformly across the sphere. This property is one of the main reasons why shrinkage becomes necessary in high dimensions, and motivates the studies in both Chapter 2 and Chapter 3.

1.1.1.2 Low-Rank Signal and Spiked Models

While the pure noise setting reveals the universal behavior of Gram matrices, the situation changes fundamentally once low-rank structure is injected into the data. Many high-dimensional datasets contain a small number of dominant factors. These structures manifest as spikes, i.e., eigenvalues that protrude above the noise bulk and whose associated eigenvectors carry meaningful information about the signal. In this subsection, we will reveal some examples from the simplest to to more complex ones.

(a) **Single Factor Model.** we consider a single-factor model

$$Y = \beta X^\top + E \in \mathbb{R}^{p \times n},$$

where $\beta \in \mathbb{R}^p$ is a fixed factor loading vector, $X \in \mathbb{R}^n$ is a vector of i.i.d. factor realizations, and E is additive noise with i.i.d. entries. Under the regime $p/n \rightarrow \infty$, the expected sample covariance becomes

$$\mathbb{E} \left[\frac{1}{n} Y Y^\top \right] = \beta \beta^\top + \sigma^2 I_p,$$

which is a rank-one spiked model. The leading eigenvalue is $\|\beta\|^2 + p\sigma^2$, and the corresponding eigenvector aligns with the signal direction $\beta/\|\beta\|$.

(b) **Dense Graph with Self-Loops via Random Adjacency Matrix.** We let $A \in \{0, 1\}^{p \times p}$ be the adjacency matrix of a directed graph with n nodes and p edges. Each element of A follows i.i.d. Bernoulli(q) for q independent of n and p , i.e., the resulting graph from A is dense and can contain self-loops. In the context of social networks, this model corresponds to a scenario where individuals form connections uniformly at random with a fixed probability q . For dense graph, most

possible connections exist, leading to a highly interconnected community. Such dense patterns can occur in small, well-mixed groups or in platforms with high incentivization for connection.

Then we define the $p \times p$ symmetric matrix $G = \frac{1}{n}AA^\top$. Each entry G_{ij} measures the normalized number of common out-neighbors shared by nodes i and j . The resulting G then encodes second-order similarity between individuals, where diagonal entries indicate individual activity levels, and off-diagonal entries quantify shared influence targets. When G_{ij} is high, it implies that individuals i and j share many common out-neighbors as they tend to follow or interact with the same set of people. This perspective is especially useful in uncovering latent community structures, role equivalence, or functional similarity in large-scale social systems.

From model setup, we have,

$$\mathbb{E}((AA^\top)_{ij}) = \begin{cases} pq, & i = j \\ pq^2, & i \neq j \end{cases}$$

So

$$\mathbb{E}(G) = \frac{p}{n}(q - q^2)I_p + \frac{pq^2}{n}\mathbf{1}_p\mathbf{1}_p^\top$$

and the leading eigenvalue of $\mathbb{E}(G)$ is:

$$\lambda_1 = \frac{p}{n}(q - q^2) + \frac{p^2q^2}{n},$$

with eigenvector $v_1 = \frac{1}{\sqrt{p}}\mathbf{1}_p$, and the remaining eigenvalues are $\frac{p}{n}(q - q^2)$.

- (c) **Sparse Directed Graph via Random Adjacency Matrix.** Now we consider the sparse random graph as an extension to previous example. Let $A \in \{0, 1\}^{p \times p}$ be a sparse binary matrix where each entry A_{ij} is an independent Bernoulli random

variable with success probability $\sqrt{n/p}$. This defines a directed Erdős–Rényi graph $G = (V, E)$ with $|V| = p$ nodes where directed edges occur independently with probability $\sqrt{n/p}$. The parameterization $n/p \ll 1$ ensures the graph remains sparse yet connected, capturing the “small-world” property observed in real-world social networks [8].

In this setting, each entry of the adjacency matrix $A \in \{0, 1\}^{p \times p}$, including the diagonal, is drawn independently as $A_{ij} \sim \text{Bernoulli}(\sqrt{n/p})$. Therefore, the expected adjacency matrix becomes

$$\mathbb{E}[A] = \sqrt{n/p} \cdot \mathbf{1}_p \mathbf{1}_p^\top,$$

indicating that every pair of nodes, including self-pairs, is equally likely to be connected by a directed edge. As a result, the expected in-degree and out-degree of each node is $\sqrt{np}$, and the matrix is numerically rank one.

To understand the spectral structure of $G = \frac{1}{n} A A^\top$, we now present the detailed study of its expected value as well as the behavior of its leading eigenvalue, the corresponding principal eigenvector, and the second eigenvector.

Recall $A_{ij} \sim \text{Bernoulli}(\sqrt{n/p})$, we have

$$\mathbb{E}(A A^\top)_{ij} = \mathbb{E}(A_i^\top A_j) = \sum_{k=1}^p \mathbb{E}(A_{ik} A_{jk}) = \begin{cases} \sqrt{np}, & i = j \\ n, & i \neq j \end{cases} \quad (1.2)$$

so

$$\mathbb{E} \left(\frac{1}{n} A A^\top \right) = \frac{p}{n} (\sqrt{n/p} - n/p) I_p + \frac{p}{n} \frac{n}{p} \mathbf{1}_p \mathbf{1}_p^\top = (\sqrt{p/n} - 1) I_p + \mathbf{1}_p \mathbf{1}_p^\top. \quad (1.3)$$

The spectral shift technique directly implies the leading eigenvalue of $\mathbb{E}(A)$ is

$$\lambda_1 = \sqrt{p/n} - 1 + p, \quad (1.4)$$

and the remaining eigenvalues approximately equal to $\sqrt{p/n} - 1$. Furthermore, the leading eigenvector of $\mathbb{E}(G)$ satisfies $v_1 \propto \mathbf{1}_p$, so

$$v_1 = z = \frac{1}{\sqrt{p}} \mathbf{1}_p. \quad (1.5)$$

It naturally implies we can apply James-Stein shrinkage on leading eigenvector of G as v_1 be the target. Also, v_1 is a point on $\mathbb{S}^{p-1}$ which points in the direction where all coordinates grow equally, and represents the equal-weight exposure direction across all assets.

In high-dimensional regime, the directions orthogonal to the principal component exhibit no preferred orientation. Formally, the bulk eigenvectors (for example, the second eigenvector in our model) converge to the Haar measure on $\mathbb{S}^{p-1}$ in distribution.

To verify the theoretical structure of the covariance matrix and its spectral behavior, we generate a sequence of random adjacency (binary) matrix $A \in \mathbb{R}^{p \times p}$, where each entry is drawn independently from a Bernoulli distribution with success probability $s = \sqrt{n/p}$. We set the parameters $p = 1000$, $n = 128$, and average the resulting $G = \frac{1}{n} AA^\top$ over 100 simulations.

The top panels in Fig1.4 displays the full spectrum of the matrix G . The lead eigenvalue is clearly separated from the rest, reaching approximately 10000, while other bulk eigenvalues stay very small. This aligns with the theoretical result as in Eq 1.4. The bottom left panel shows the the entry distribution of the leading

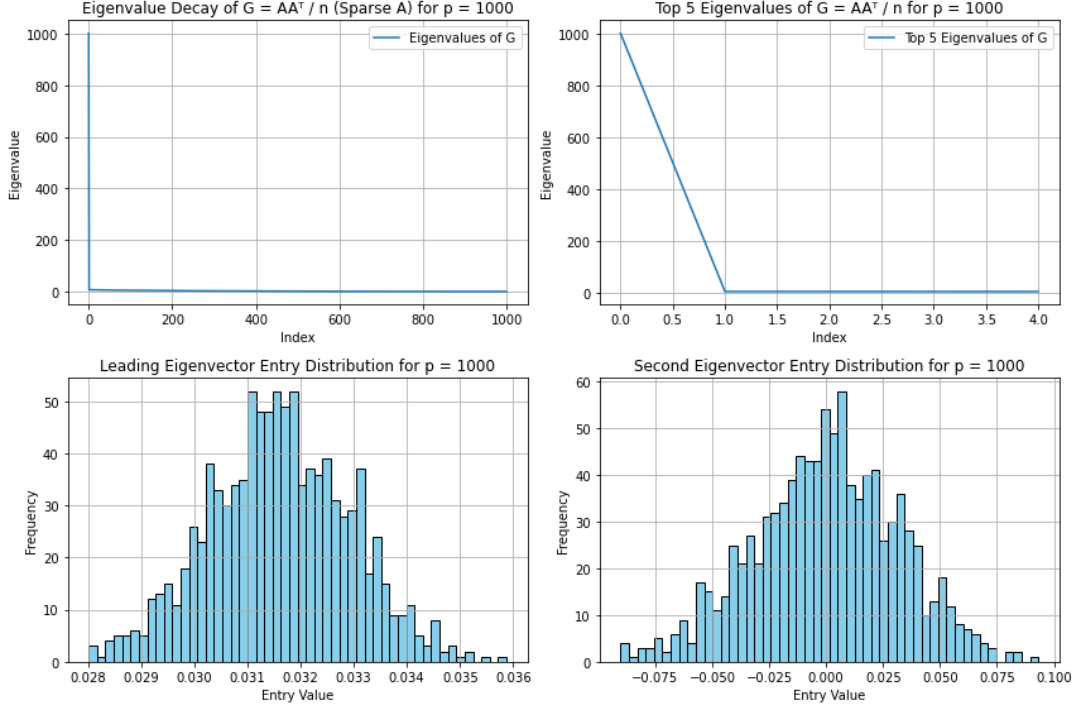
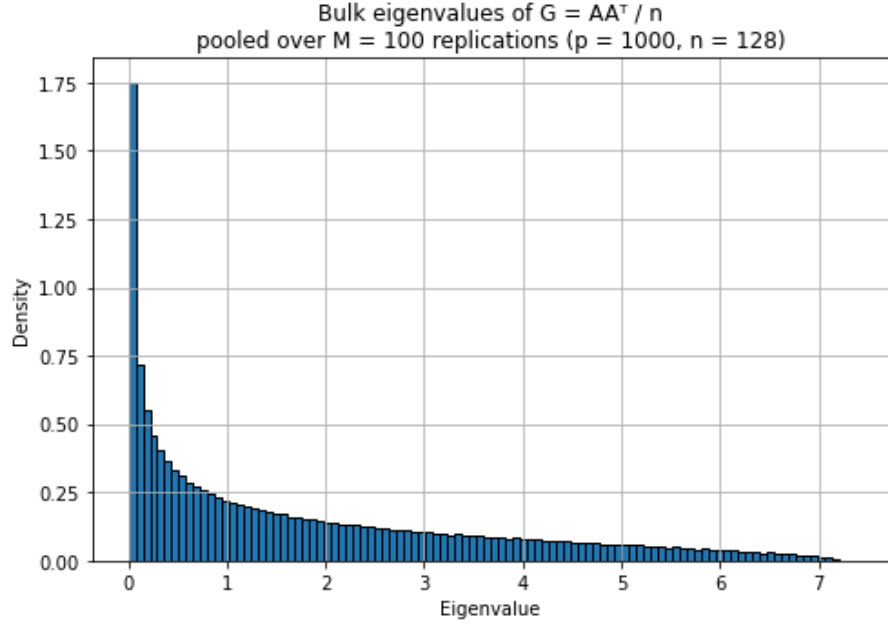


Figure 1.4: Top panels: eigenvalues of G ; Bottom panels: histogram of leading eigenvector and second eigenvector of G .

eigenvector. All entries are tightly concentrated around a positive constant near $1/\sqrt{p}$. Although the theoretical leading vector of $E(G)$ is exactly constant, the leading eigenvector of the empirical matrix G contains small fluctuations around that direction due to random sampling noise, and results in an approximately normal distribution centered at $1/\sqrt{p}$. The bottom right panel illustrates the histogram of the entries of second eigenvector. It is approximately normally distributed with mean 0 and standard deviation 0.01, i.e. $1/\sqrt{p}$, and matches the theoretical prediction that non-leading eigenvectors in a spiked model converge in distribution to the Haar measure on the unit sphere $\mathbb{S}^{p-1}$.

Furthermore, we show the histogram of bulk eigenvalues of G in Figure 1.5. The eigenvalues form a highly right-skewed distribution with a substantial concentration near zero and a long positive tail, which displays Marchenko-Pastur type limits. It

Figure 1.5: Bulk eigenvalues of A .

comes from the fact that G is gram matrix.

- (d) **Random Graph via Incidence Matrix and Edge Covariance.** For a graph with n nodes and p edges, the incidence matrix $Y \in \mathbb{R}^{p \times n}$ is defined such that each row corresponds to an edge, i.e.

$$Y_{ik} = \begin{cases} 1 & \text{if node } k \text{ is incident with edge } e_i \\ 0 & \text{otherwise} \end{cases} \quad (1.6)$$

A more rigorous definition is of the following:

Definition 1.1.1 (Random Incidence Matrix). *Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a probability space. We define the random matrix $Y \in \{0, 1\}^{p \times n}$ via:*

- (a) *For each edge $e_i, i \in \{1, 2, \dots, p\}$, independently sample a node pair uniformly from all pairs on n nodes.*

(b) Define the i -th row of Y by

$$Y_{ik} = \mathbf{1}_{\{k=K_{i,1}\}} + \mathbf{1}_{\{k=K_{i,2}\}}. \quad (1.7)$$

Correspondingly, we can define $A = YY^\top$. The i, j -th element of A equals to

$$A_{ij} = \sum_{k=1}^n Y_{ik}Y_{jk} = \#\{\text{nodes jointly connected by edge } i \text{ and edge } j\}, \quad (1.8)$$

The diagonal elements of A equal to

$$A_{ii} = \sum_{k=1}^n Y_{ik}^2 = \frac{2}{n}, \quad (1.9)$$

as there are exact two 1s in each row.

In financial modeling, the incidence matrix Y of the (undirected) graph can be interpreted as a network of financial relationships, where n represents the number of assets, and p represents the number of contracts. Each row of Y corresponds to a financial relationship, and the non-zero entries in that row identify the two assets involved. The matrix A represents a normalized exposure interaction matrix with A_{ii} being the normalized number of assets involved via incident i , and A_{ij} measures the structural overlap between two financial relationships.

For diagonal terms of A , we have

$$\mathbb{E}(A_{ii}) = 2 \quad (1.10)$$

follows from Eq 1.9. For off-diagonal terms, we have

$$\begin{aligned}
\mathbb{E}(A_{ij}) &= \sum_{k=1}^n \mathbb{E} [\mathbf{1}_{\{k \in \{K_{i,1}, K_{i,2}\}\}} \mathbf{1}_{\{k \in \{K_{j,1}, K_{j,2}\}\}}] \\
&= \sum_{k=1}^n \mathbb{P}(k \in \{K_{i,1}, K_{i,2}\}) \mathbb{P}(k \in \{K_{j,1}, K_{j,2}\}) \\
&= n \left(\frac{2}{n} \frac{2}{n} \right) = \frac{4}{n}.
\end{aligned} \tag{1.11}$$

Therefore, we can write

$$\mathbb{E}(A) = \left(2 - \frac{4}{n}\right) I_p + \frac{4}{n} \mathbf{1}_p \mathbf{1}_p^\top. \tag{1.12}$$

By spectral shift technique, we have the leading eigenvalue of $\mathbb{E}(A)$ is

$$\lambda_1 = 2 + \frac{4(p-1)}{n}, \tag{1.13}$$

and the remaining eigenvalues λ equal to $2 - \frac{4}{n}$. Therefore it implies the leading eigenvector of $\mathbb{E}(A)$ is essentially $z = \frac{1}{\sqrt{p}} \mathbf{1}_p$, and the remaining eigenvectors converges to the Haar measure on $\mathbb{S}^{p-1}$ in distribution. A larger eigenvalue suggests that many edges tend to be incident to the same set of vertices, and is encoded by the structure of the corresponding eigenvector.

Now we consider the James-Stein shrinkage on the leading eigenvector of $\frac{1}{n} Y Y^\top$. Let v denote the leading eigenvector of $\frac{1}{n} Y Y^\top$, then the James-Stein estimator (before normalization) for v is

$$v_{JS} = \psi^2 v + z_{\perp v} \frac{z^\top v (1 - \psi^2)}{|z - z_v|}, \tag{1.14}$$

where $z_v = \langle v, z \rangle v$, $z_{\perp v} = \frac{z - z_v}{|z - z_v|}$, and $\psi^2 = \frac{s_{1,p}^2 - l_p^2}{s_{1,p}^2}$ for $s_{1,p}^2$ being the leading eigenvalue

of A and l_p^2 being the average non-zero bulk eigenvalues of A .

- (e) **Structured Incidence Observation Matrix.** Now we consider a modification of the above model. Let E be a $p \times n$ edge incidence matrix constructed as above, i.e. each row of E consists of two 1s and $(n - 2)$ 0s. For $\beta = (1, \dots, 1)^\top \in \mathbb{R}^p$, and $X \in \mathbb{R}^n$ with each element of X be i.i.d. *Bernoulli*(q) variable. We let

$$Y = \beta X^\top + E,$$

and we call Y the structured incidence observation matrix.

We have

$$A = YY^\top = \beta X^\top X \beta^\top + \beta X^\top E^\top + EX \beta^\top + EE^\top. \quad (1.15)$$

The first term satisfies

$$\beta X^\top X \beta^\top = n \sum_{i=1}^n X_i^2 \mathbf{1}_p \mathbf{1}_p^\top = n \sum_{i=1}^n X_i \mathbf{1}_p \mathbf{1}_p^\top, \quad (1.16)$$

so the expectation of the leading eigenvalue of $\beta X^\top X \beta^\top$ is npq , with $z = \mathbf{1}/\sqrt{p}$ being the corresponding eigenvector.

As for the cross terms $(\beta X^\top E^\top + EX \beta^\top)$, the leading eigenvalue is calculated as

$$\langle \beta, EX \rangle + \sqrt{p} \|EX\|. \quad (1.17)$$

Note that

$$\langle \beta, EX \rangle = \sum_{j=1}^n X_j \sum_{i=1}^p E_{ij}, \quad (1.18)$$

and for each j , $\sum_{i=1}^p E_{ij}$ is the degree of node j , denoted as d_j , so

$$\mathbb{E}(\langle \beta, EX \rangle) = \sum_{j=1}^n \mathbb{E}(X_j) \mathbb{E}(d_j) = q \cdot n \cdot \frac{2p}{n} = 2pq. \quad (1.19)$$

Note that each $(EX)_i = X_{j_1} + X_{j_2}$, with j_1, j_2 being two random columns in the i -th row, so $(EX)_i \sim \text{Bin}(2, q)$, and

$$\mathbb{E}(\|EX\|^2) = \sum_{i=1}^p \mathbb{E}((EX)_i^2) = p(2q + 2q^2), \quad (1.20)$$

so

$$\mathbb{E}(\|EX\|) = \sqrt{2pq(1+q)}, \quad (1.21)$$

and the leading eigenvalue for the cross-term is $2pq + p\sqrt{2q(1+q)}$. Also note that

$$\mathbb{E}(X^\top E^\top) = \sum_{j=1}^n \mathbb{E}(X_j) \mathbb{E}(E_{j,\cdot})^\top = q \sum_{j=1}^n \mathbb{E}(E_{j,\cdot})^\top \propto q \mathbf{1}_p, \quad (1.22)$$

so the leading eigenvector for the cross-term is also $z = \mathbf{1}_p / \sqrt{p}$.

Therefore the leading eigenvalue of $A = YY^\top$ satisfies:

$$\mathbb{E}(\lambda_1) = (n+2)pq + p\sqrt{2q(1+q)} + 2 + \frac{4(p-1)}{n}, \quad (1.23)$$

and the corresponding eigenvector is $z = \mathbf{1}_p / \sqrt{p}$.

1.1.2 Stochastic Block Models and the Semicircle Law

In contrast to previous examples, we now focus on a fundamentally different source of low-rank structure called community structure in random networks. The stochastic block model (SBM) is one of the most important generative models of this kind with

community structure [9].

We begin with the simplest case of an SBM consisting of two latent communities. Let the vertex set be V with $|V| = N$, and let it be partitioned into

$$V = V_1 \cup V_2, \quad |V_1| = \lfloor N/2 \rfloor, \quad |V_2| = N - |V_1|.$$

Let $A \in \{0, 1\}^{N \times N}$ denote the adjacency matrix of the undirected graph. In the classical two-block SBM, the entries of A are generated independently according to

$$\mathbb{P}(A_{ij} = 1 \mid i \in V_a, j \in V_b) = \begin{cases} p_N, & a = b, \\ q_N, & a \neq b, \end{cases} \quad a, b \in \{1, 2\},$$

where p_N denotes the within-community connection probability and q_N denotes the between-community probability. This formulation captures the essential phenomenon that vertices belonging to the same community are more likely to be connected than vertices belonging to different communities with choice of $p_N > q_N$.

We are interested in the spectral structure of the population mean matrix

$$M = \mathbb{E}[A].$$

For the balanced two-block model described above, let $\mathbf{1} = (1, \dots, 1)^\top \in \mathbb{R}^N$ denote the all-ones vector and $s = (1, \dots, 1, -1, \dots, -1)^\top$ denote the community indicator vector that assigns +1 to vertices in V_1 and -1 to those in V_2 . Then M can be written compactly as a rank-two matrix:

$$M = \frac{p_N + q_N}{2} \mathbf{1}\mathbf{1}^\top + \frac{p_N - q_N}{2} ss^\top. \quad (1.24)$$

The top two eigenvalues of M is then

$$\sigma_1^2 = N \frac{p_N + q_N}{2}, \quad \sigma_2^2 = N \frac{p_N - q_N}{2},$$

with corresponding eigenvectors $\mathbf{1}$ and s respectively. All remaining $N - 2$ eigenvalues are zero. The leading eigenpair $(\sigma_1^2, \mathbf{1})$ represents the global density component of the network, while the second eigenpair (σ_2^2, s) captures the community contrast between V_1 and V_2 .

Hence the first eigenvector encodes the overall density of the network, while the second eigenvector separates the two communities.

In the random realization A , the matrix can be decomposed as

$$A = M + W,$$

where $W = A - M$ represents random fluctuations around the mean structure. The matrix W is approximately a centered Wigner matrix: its entries have mean zero, finite variance, and are weakly dependent only through symmetry. As $N \rightarrow \infty$ with suitable scaling of p_N and q_N (for instance $p_N, q_N = O(1/N)$), the empirical spectral distribution of the rescaled noise matrix

$$W_N = \frac{W}{\sqrt{N \operatorname{Var}(A_{ij})}}$$

converges almost surely to the semicircle law supported on $[-2, 2]$ (see Fig 1.6).

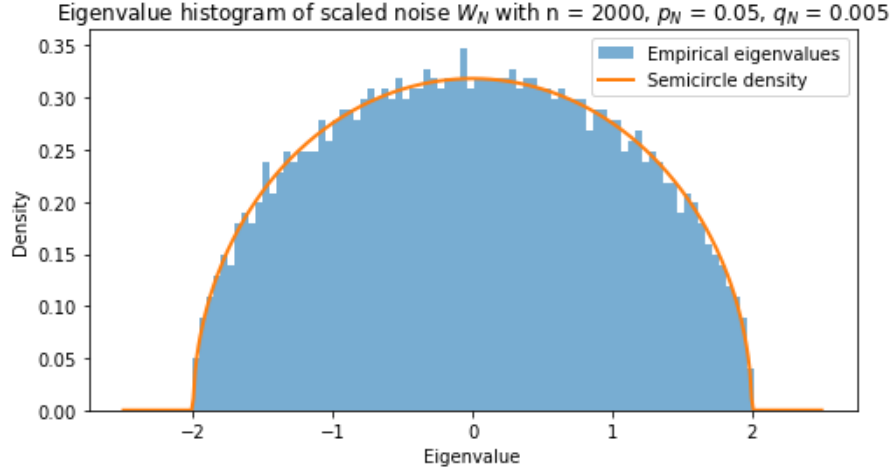


Figure 1.6: Histogram of eigenvalues of W_N with theoretical semicircle density.

1.1.3 Summary: Two Spectral Regimes in Random Graph Models

The examples presented above fall into two distinct universality classes, and each is determined by different limiting laws. We can understand their spectral behaviors through two regimes:

1. The first examples in Section 1.1.1.2 consist of Gram matrices. All matrices A that we discussed take the form $A = \frac{1}{n}YY^\top$ for some matrix Y with columns being independent observations. Despite the different interpretations of these examples, these matrices share the same underlying geometry: positive semi-definite Gram matrices with bulk eigenvalues follow MP law. Low rank structure in $\mathbb{E}(YY^\top)$ produces spikes in the spectrum, and eigenvectors orthogonal to the signal converge in distribution to Haar measure on the sphere. These properties motivate the eigenvector shrinkage ideas developed later in the thesis.
2. The SBM with preference example in Section 1.1.2 exhibits an entirely different limiting behavior. After appropriate normalization, the noise component of such

matrices converges to a Wigner ensemble, and its empirical spectrum follows the semicircle law. In this regime the graph encodes latent community structure, represented by a low-rank mean matrix whose eigenvectors determine the population blocks. The observed adjacency inherits this structure through two spike eigenvalues that separate from the semicircular bulk of noise. The resulting spectral picture is qualitatively different from the MP case. We further analyzed how this spiked structure impacts the recovery of node preferences. Beyond the spectral characterization itself, our numerical study demonstrated that the James–Stein shrinkage estimator achieves substantially more accurate preference recovery than PCA, reducing the relative estimation error by roughly 43% in all experiments.

Together, these two cases provide a unified perspective on the examples in this chapter. They illustrate how structurally different data-generating mechanisms can produce common structural feature: a low-rank signal component that produces spike eigenvalues separated from a noise bulk. This common geometry motivates the analysis of eigenvector stability and shrinkage in later chapters.

1.2 Literature Review

The asymptotic behavior of eigenvalues and eigenvectors has been a central topic in multivariate statistics since the mid-20th century. In traditional statistics, [10] showed that the limiting distribution of eigenvectors of sample covariance matrix is the conditional Haar invariant distribution when p fixed and $n \rightarrow \infty$.

However, such asymptotics become unreliable in modern high-dimensional regimes where p is large relative to n . To address this, random matrix theory (RMT) has emerged as a key framework for analyzing sample covariance matrices in the limit where both p and n diverge with $p/n \rightarrow c \in (0, \infty)$. [2] showed that, when $n/p \rightarrow \gamma >$

1, the leading eigenvalue of a spiked covariance model approaches the Tracy–Widom law of order 1. [11] stated that, when $n, p \rightarrow \infty$ and $p/n \rightarrow \gamma$ fast enough, and $\Sigma = \text{diag}(\ell_1, \dots, \ell_M, 1, 1, \dots, 1)$. The v -th sample eigenvector can be written as $\nu_v = (\nu_{v,A}^\top, \nu_{v,B}^\top)^\top$, and central limit theorem of $\nu_{v,A}/\|\nu_{v,A}\|$ holds while $\nu_{v,B}/\|\nu_{v,B}\|$ is uniformly distributed on S^{p-M-1} .

In contrast to both classical and RMT-based asymptotic frameworks, this thesis considers an alternative regime where the sample size n is fixed and the dimension $p \rightarrow \infty$. [12] stated that, for a model with m spikes, the accuracy of m eigenvector estimations decreases as the bias in corresponding eigenvalues increases. [13] states that, the limiting distribution of inner products between the sample and population eigenvectors exists, and the rest of eigenvectors are strongly inconsistent with their population counterpart.

In high-dimensional statistical settings, the inherent instability of eigenvectors has prompted the development of various regularization techniques. Among these, shrinkage estimators have emerged as particularly effective solutions, demonstrating significant improvements in estimation quality by effectively reducing variance while maintaining reasonable bias levels.

This phenomenon stems from the extreme sensitivity of sample covariance matrix eigenvectors in high dimensions, where even minor data perturbations can cause substantial directional changes. Shrinkage methods address this challenge by systematically adjusting estimates toward a target structure (such as the identity matrix or factor models). Theoretical analyses and empirical studies consistently show these approaches are especially valuable when the number of variables p is comparable to sample size n , offering reliable solutions for high-dimensional inference problems.

Shrinkage estimation, originally introduced in the context of mean estimation by [14], has become a cornerstone in modern high-dimensional statistics. The key idea is to trade unbiasedness for reduced variance by shrinking noisy empirical estimates toward a

structured target. In the context of covariance estimation, [3] proposed a linear shrinkage estimator that combines the sample covariance matrix with a multiple of the identity, and later extended this to more general targets, including low-rank factor-based structures [15].

Such James–Stein-type estimators have been shown to achieve improved performance in high-dimensional regimes, especially when the dimension-to-sample-size ratio p/n is non-negligible. Notably, [14] proved that for $p \geq 3$, their shrinkage estimator attains strictly lower risk than unbiased estimators. Theoretical results often assume that the shrinkage target is correctly specified, particularly in terms of its rank or underlying factor structure. For example, [16] proposed a James–Stein estimator $\hat{\Sigma}_{\sharp}$ on sample covariance matrix such that the discrepancy induced from $\hat{\Sigma}$ converges as $p \uparrow \infty$.

However, relatively few works have studied the behavior of shrinkage estimators under structural misspecification—for example, when the assumed number of latent factors K differs from the true rank q . This thesis addresses this gap by analyzing the asymptotic impact of rank misspecification on optimization bias and discrepancy.

1.3 Mathematical Preliminaries

This section reviews the key mathematical concepts and tools that will be used throughout the thesis. We organize the discussion into three components: linear algebraic foundations, random matrix theory, and high-dimensional geometry.

1.3.1 Linear Algebra

The analysis of high-dimensional estimators hinges on a geometric perspective, where data and parameters are viewed as vectors in Euclidean space. This subsection establishes the linear algebraic vocabulary: orthogonality, projections, matrix decompositions

and etc., all essential for describing the behavior of estimators. Crucially, we focus on concepts that characterize *low-dimensional structure* within high-dimensional spaces, such as eigenvalues and singular values, which will later govern the statistical performance of shrinkage estimators.

We begin by recalling fundamental structures that define geometry and distance.

Let $\mathbb{R}^p$ denote the p -dimensional Euclidean space, and $\mathbb{R}^{p \times q}$ the space of real $p \times q$ matrices. We start with some definitions on orthogonal matrices and inner-product space.

Definition 1.3.1 (Orthogonal Matrix [17]). *A matrix $Q \in \mathbb{R}^{n \times n}$ is called **orthogonal** if it satisfies*

$$Q^\top Q = QQ^\top = I_n.$$

Definition 1.3.2 (Projection Matrix [18]). *Let $Q \in \mathbb{R}^{n \times k}$ have orthonormal columns, i.e., $Q^\top Q = I_k$. Then the matrix*

$$P = QQ^\top$$

is called the orthogonal projection onto the column space of Q . It satisfies $P^2 = P$ and $P^\top = P$.

Definition 1.3.3 (Inner product [19]). *Let $\mathcal{X}$ be a linear space over $\mathbb{K}$. A sesquilinear form $\langle \cdot, \cdot \rangle : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{K}$ is called an **inner product** if it satisfies*

$$(i) \quad \forall x, y \in \mathcal{X}, \langle x, y \rangle = \overline{\langle y, x \rangle}$$

$$(ii) \quad \forall x \in \mathcal{X}, \langle x, x \rangle \geq 0, \text{ and } \langle x, x \rangle = 0 \Leftrightarrow x = \mathbf{0}.$$

Remark 1.3.4. *We call $(\mathcal{X}, \langle \cdot, \cdot \rangle)$ an **inner product space**.*

Definition 1.3.5 (Orthogonality [19]). *Let $\mathcal{X}$ be a linear space over $\mathbb{K}$ equipped with $\langle \cdot, \cdot \rangle$. $\forall x, y \in \mathcal{X}$, we say x and y are orthogonal, if*

$$\langle x, y \rangle = 0,$$

and denote as $x \perp y$. For $M \subset \mathcal{X}$ be a non-empty subset and $x \in \mathcal{X}$, if $\forall y \in M$, we have $x \perp y$, then $x \perp M$. Furthermore, we denote the set

$$M^\perp = \{x \in \mathcal{X} : x \perp M\}$$

be the **orthogonal complement** of M .

Definition 1.3.6 (Norm [19]). Let $\mathcal{X}$ be a linear space over $\mathbb{K}$. A mapping $\|\cdot\| : \mathcal{X} \rightarrow \mathbb{R}$ is called a **norm**, if it satisfies

$$(i) \quad \forall x \in \mathcal{X}, \|x\| \geq 0, \text{ and } \|x\| = 0 \leftrightarrow x = \mathbf{0}$$

$$(ii) \quad \forall a \in \mathbb{K}, x \in \mathcal{X}, \|ax\| = |a| \|x\|$$

$$(iii) \quad \forall x, y \in \mathcal{X}, \|x + y\| \leq \|x\| + \|y\|.$$

Definition 1.3.7 (Vector norm [19]). For $\mathbb{K} = \mathbb{R}^n$ in Definition 1.3.6, we say the corresponding $\|\cdot\|$ is a **vector norm**.

Remark 1.3.8. For $x \in \mathcal{X}$, the norm $\|x\| = \sqrt{\langle x, x \rangle}$ is a special case of a vector norm. In finite-dimensional spaces (e.g., $\mathbb{R}^n$ or $\mathbb{C}^n$), other common vector norms include:

$$\|x\|_1 := \sum_{i=1}^n |x_i|, \quad \|x\|_\infty := \max_{1 \leq i \leq n} |x_i|, \quad \|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{1/p},$$

and are called $\ell_1, \ell_\infty, \ell_p$ norm, respectively.

Remark 1.3.9. The columns (and rows) of an orthogonal matrix form an orthonormal basis of $\mathbb{R}^n$. Orthogonal matrices preserve inner products and Euclidean norms. That is, for any $x, y \in \mathbb{R}^n$, we have

$$\langle Qx, Qy \rangle = \langle x, y \rangle, \quad \text{and} \quad \|Qx\|_2 = \|x\|_2.$$

Theorem 1.3.10 (Holder's inequality [20]). *Let $\mathcal{X}$ be a linear space over $\mathbb{R}^n$. For any $x, y \in \mathcal{X}$ and $\frac{1}{p} + \frac{1}{q} = 1$,*

$$|\langle x, y \rangle| \leq \|x\|_p \|y\|_q. \quad (1.25)$$

While vector norms measure the magnitude of data, *matrix norms* are indispensable for quantifying the effect of linear transformations (e.g., applying an estimator) or the size of perturbations in high-dimensional settings.

Definition 1.3.11 (Matrix norm [19]). *For a matrix $A \in \mathbb{C}^{m \times n}$, the **induced matrix norm** (or operator norm) is defined by:*

$$\|A\| := \sup_{x \neq 0} \frac{\|Ax\|}{\|x\|} = \sup_{\|x\|=1} \|Ax\|,$$

where $\|\cdot\|$ on the right-hand side is a vector norm.

Remark 1.3.12. *Frobenius norm is a widely used matrix norm. This norm can be defined in various ways:*

$$\|A\|_F = \sqrt{\sum_i^m \sum_j^n |a_{ij}|^2} = \sqrt{\text{trace}(A^* A)} = \sqrt{\sum_{i=1}^{\min\{m,n\}} \sigma_i^2} \quad (1.26)$$

for $A = \{a_{ij}\} \in \mathbb{R}^{m \times n}$ and σ_i the singular value of A . Furthermore, for T be orthogonal transformation, then

$$\|A\|_F = \|TA\|_F. \quad (1.27)$$

We now move to eigenvalue decomposition, which is the key element in the entire thesis.

Definition 1.3.13 (Eigenvalue and eigenvector [17]). *Let $A \in \mathbb{R}^{n \times n}$, if there exists*

$\lambda_0 \in \mathbb{R}$, and a non-zero vector $a \in \mathbb{R}^n$, such that

$$Aa = \lambda_0 a,$$

then we say λ_0 is an **eigenvalue** of matrix A with a the corresponding **eigenvector**.

Remark 1.3.14. If $A \in \mathbb{R}^{n \times n}$ has n distinct eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$ with corresponding eigenvectors $a_1, a_2, \dots, a_n$, then $U^{-1}AU = \text{diag}(\lambda_1, \dots, \lambda_n)$, with $U = (a_1, \dots, a_n)$.

The singular value decomposition (SVD) is arguably the most powerful tool in high-dimensional data analysis. It provides a complete geometric description of any data matrix Y by decomposing it into orthogonal components ordered by their importance (singular values). This decomposition is the foundation for Principal Component Analysis (PCA) and directly reveals the signal-to-noise structure in random matrix models.

Theorem 1.3.15 (Singular Value Decomposition [18]). Let $A \in \mathbb{R}^{p \times q}$. Then there exist orthogonal matrices $U \in \mathbb{R}^{p \times p}$ and $V \in \mathbb{R}^{q \times q}$, and a diagonal matrix $\Sigma \in \mathbb{R}^{p \times q}$ with non-negative entries $\sigma_1 \geq \sigma_2 \geq \dots \geq 0$ such that

$$A = U\Sigma V^\top.$$

Remark 1.3.16. The nonzero entries σ_i of Σ are called the **singular values** of A . They are the square roots of the nonzero eigenvalues of both $A^\top A$ and $AA^\top$. The columns of U and V are referred to as the left and right singular vectors of A , respectively.

Remark 1.3.17. Given the SVD $A = U\Sigma V^\top$, the Moore–Penrose pseudoinverse is given by

$$A^\dagger = V\Sigma^\dagger U^\top,$$

where $\Sigma^\dagger$ is obtained by taking the reciprocal of each nonzero entry in Σ and transposing the resulting matrix.

We now present the spectral theorem, which characterizes when a real matrix can be diagonalized via an orthogonal similarity transformation.

Theorem 1.3.18 (Spectral Theorem [18]). *Let $A \in \mathbb{R}^{n \times n}$ be a symmetric matrix, i.e., $A = A^\top$. Then there exists an orthogonal matrix $Q \in \mathbb{R}^{n \times n}$ and a diagonal matrix $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$ such that*

$$A = Q\Lambda Q^\top.$$

The columns of Q are orthonormal eigenvectors of A , and the diagonal entries of Λ are the corresponding real eigenvalues.

Remark 1.3.19. *Note that Λ is a diagonal matrix in Theorem 1.3.18, we can think of A as sum of n rank-one matrices.*

Remark 1.3.20. *The spectral theorem guarantees that all real symmetric matrices are orthogonally diagonalizable. In particular, the eigenvectors corresponding to distinct eigenvalues are orthogonal, and a full orthonormal basis of eigenvectors always exists.*

Not every real matrix is diagonalizable. In particular, non-square matrices do not admit eigenvalue decompositions. To generalize the spectral representation of real matrices beyond the square case, we introduce the singular value decomposition (SVD), which exists for all real matrices.

Definition 1.3.21 (Moore–Penrose Pseudoinverse [21]). *Let $A \in \mathbb{R}^{p \times q}$. The **Moore–Penrose inverse** of A , denoted $A^\dagger \in \mathbb{R}^{q \times p}$, is the unique matrix satisfying the following four properties:*

1. $AA^\dagger A = A$

2. $A^\dagger A A^\dagger = A^\dagger$
3. $(A A^\dagger)^\top = A A^\dagger$
4. $(A^\dagger A)^\top = A^\dagger A$.

Remark 1.3.22. *The Moore-Penrose inverse gives the minimum-norm solution to the least-squares problem $\min_x \|Ax - b\|_2$. That is, for $b \in \mathbb{R}^p$, the solution $x^* = A^\dagger b$ satisfies $\|x^*\| = \min\{\|x\| : Ax = b\}$ when such x exists.*

A central technique in high-dimensional statistics is to avoid working with the large $(p \times p)$ sample covariance matrix S directly. Instead, we often analyze its smaller $(n \times n)$ dual matrix $L = \frac{1}{p} Y^\top Y$. The following theorem, which is fundamental to our analysis, establishes an explicit equivalence between the eigenstructures of S and L .

Theorem 1.3.23. *$Lu = \ell^2 u$ for $\ell^2 = s^2 n/p > 0$ and $u = Y^\top v/(\sqrt{n}s)$ if and only if $Sv = s^2 v$ where $v = Yu/(\sqrt{p}\ell)$ and $s^2 = \ell^2 p/n > 0$.*

Proof. Multiple $Lu = \ell^2 u$ by Y from both sides, we have

$$\frac{Y Y^\top Y u}{p} = \ell^2 Y u \Rightarrow S Y u = \left(\frac{\ell^2 p}{n} \right) Y u. \quad (1.28)$$

Note that $(Yu)^\top (Yu) = (u^\top L u)p = \ell^2 p$, we have $v = Yu/(\sqrt{p}\ell)$ with $\|v\| = 1$. Divide the right hand side of Eq 1.28 yields $Sv = s^2 v$ for $s^2 = \ell^2 p/n$. The converse is identical. $\square$

In statistical applications, we rarely observe the true population matrix (e.g., the covariance matrix Σ). Instead, we observe a noisy sample version $\hat{\Sigma}$. A critical question is: how do the eigenvalues and eigenvectors of $\hat{\Sigma}$ relate to those of Σ ? Matrix perturbation theory provides the answers. Weyl's inequality controls the stability of eigenvalues, while the Davis-Kahan theorem bounds the rotation of eigenspaces.

Theorem 1.3.24 (Weyl's Inequality [18]). *Let $A, E \in \mathbb{R}^{n \times n}$ be symmetric matrices. Let $\lambda_i(\cdot)$ denote the i -th largest eigenvalue of a symmetric matrix. Then for all $1 \leq i \leq n$, we have*

$$|\lambda_i(A + E) - \lambda_i(A)| \leq \|E\|_2,$$

where $\|\cdot\|_2$ denotes the spectral norm (i.e., the largest singular value).

Example 1.3.25. *Let $M \in \mathbb{R}^{p \times n}$ with $p \leq n$, the singular values of M , $\sigma_i(M)$ are the p positive eigenvalues of*

$$\begin{pmatrix} 0 & M \\ M & 0 \end{pmatrix}.$$

Weyl's inequality gives the bound for the perturbation in the singular values of a matrix M due to an additive perturbation E :

$$\|\sigma_i(M + E) - \sigma_i(M)\| \leq \sigma_1(E) = \|E\|_2. \quad (1.29)$$

Remark 1.3.26. *Weyl's inequality quantifies how much eigenvalues can change when a symmetric matrix is slightly modified. The result shows that if the modification (measured by its spectral norm) is small, then all of the matrix's eigenvalues will only shift by similarly small amounts - essentially making eigenvalues stable under small perturbations.*

Theorem 1.3.27 (Davis–Kahan Sin Θ Theorem [22]). *Let $A, \tilde{A} \in \mathbb{R}^{n \times n}$ be symmetric matrices. Let $U, \tilde{U} \in \mathbb{R}^{n \times k}$ contain orthonormal columns spanning the invariant subspaces of A and $\tilde{A}$ associated with eigenvalues in intervals $\mathcal{I}$ and $\tilde{\mathcal{I}}$, respectively. Suppose the distance between these intervals is at least $\delta > 0$. Then,*

$$\|\sin \Theta(U, \tilde{U})\|_2 \leq \frac{\|A - \tilde{A}\|_2}{\delta},$$

where $\sin \Theta(U, \tilde{U})$ is the diagonal matrix of principal angles between $\text{col}(U)$ and $\text{col}(\tilde{U})$.

Example 1.3.28. Consider the covariance model

$$\Sigma = I_p + \lambda v v^\top, \quad (1.30)$$

the leading eigenvector of Σ is v , and the gap between the leading eigenvalue and the other eigenvalues is λ . Given n i.i.d. samples $X_1, \dots, X_n$ drawn from $N(0, \Sigma)$, the sample covariance matrix $\hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n X_i X_i^\top$ with leading eigenvector $\hat{v}$.

Davis-Kahan yields

$$\|\sin \Theta(v, \hat{v})\| \leq \frac{\|\Sigma - \hat{\Sigma}\|_2}{\lambda} \lesssim (1 + \lambda) \sqrt{p/n}. \quad (1.31)$$

The last part in the above equation comes from Bernstein inequality. Therefore, if $\lambda \gg \sqrt{p/n}$, the sample leading eigenvector $\hat{v}$ is close to population eigenvector v , and if $\lambda \ll \sqrt{p/n}$, $\hat{v}$ is a random vector.

Remark 1.3.29. When symmetric matrices are perturbed, the Davis-Kahan theorem tells us how much their eigenspaces can rotate. The amount of rotation depends on both the size of the perturbation (measured by the spectral norm) and how well-separated the eigenvalues are.

1.3.2 Random Matrix Theory

The mathematical framework of random matrix theory (RMT) has become indispensable for characterizing the spectral behavior of sample covariance matrices in high dimensions. Traditional RMT results derive from an asymptotic regime where dimensionality p and sample size n jointly increase with $p/n \rightarrow c \in (0, \infty)$, yielding critical

insights that have reshaped statistical understanding of PCA limitations and regularized estimation. Although our work primarily investigates the alternative asymptotic setting where $p \rightarrow \infty$ with n fixed or bounded, classical RMT principles continue to provide crucial perspective on the structural challenges of high-dimensional covariance estimation. Here we summarize foundational RMT results that simultaneously motivate and contrast with our fixed-sample asymptotic analysis.

The foundational ideas of RMT trace back to [23], who introduced random matrices in the 1990s to study the energy levels of heavy atomic nuclei. His semicircle law described the limiting spectral distribution (LSD) of eigenvalues of large symmetric random matrices. Later developments such as the Marchenko–Pastur law for sample covariance matrices and the Tracy–Widom distribution for the largest eigenvalue established RMT as a central tool in understanding high-dimensional stochastic systems.

In modern statistics and machine learning, RMT plays a pivotal role in high-dimensional inference, especially when the number of variables p is comparable to or larger than the number of observations n . In such regimes, traditional asymptotic tools (e.g., law of large numbers, central limit theorem) fail to describe the empirical behavior of eigenvalues and eigenvectors. Instead, RMT provides non-classical limits and precise spectral characterizations. In this section, we present several results that hold for a broad class of random matrix models. These theorems are particularly useful in high-dimensional asymptotics, where universal behaviors often emerge across different modeling assumptions.

One cornerstone of random matrix theory concerns the spectral behavior of Wigner matrix.

Definition 1.3.30 (Wigner Matrix [24]). *Let $\{Y_i\}_{1 \leq i}$ and $\{Z_{ij}\}_{1 \leq i < j}$ be two real-valued families of zero mean, i.i.d. random variables. Furthermore, suppose that $\mathbb{E}(Z_{12}^2) = 1$*

and for each $k \in \mathbb{N}$,

$$\max(\mathbb{E}|Z_{12}^k|, \mathbb{E}|Y_1|^k) < \infty.$$

Consider a $n \times n$ symmetric matrix M_n whose entries are given by:

$$\begin{cases} M_n(i, i) = Y_i, \\ M_n(i, j) = Z_{ij} = M_n(j, i), \quad \text{if } i < j. \end{cases}$$

The matrix M_n is known as a **real symmetric Wigner matrix**.

For the normalized Wigner matrix $X_n = M_n/\sqrt{n}$, its n eigenvalues $\lambda_1(X_n) \leq \dots \leq \lambda_n(X_n)$ form the empirical spectral distribution (ESD):

$$\mu_n := \frac{1}{n} \sum_{j=1}^n \delta_{\lambda_j(X_n)},$$

with $\delta_{\lambda_j(X_n)}(x)$ being the indicator function $\mathbf{1}_{\lambda_j(X_n) \leq x}$. We have the semicircle law:

Theorem 1.3.31 (Semicircle Law [24]). *Let $\{X_n\}_{n=1}^\infty$ be a sequence of scaled Wigner matrices, then μ_n converges weakly, in probability to the semicircle distribution*

$$\mu(x)dx = \frac{1}{2\pi} \sqrt{(4-x^2)_+} dx.$$

And the Stieltjes transform of μ_n outside of the support $[-2, 2]$ is:

$$s_n(z) = \int_{\mathbb{R}} \frac{1}{x-z} d\mu_n = \frac{1}{n} \text{tr}(X_n - zI)^{-1} = \frac{-z + \sqrt{z^2 - 4}}{2}.$$

The semicircle law represents the universal limiting behavior of symmetric noise matrices and forms one of the two fundamental spectral limits in random matrix theory. The second one is the Marchenko-Pastur law, which gives limiting behavior of Gram

matrices.

Definition 1.3.32 (Marchenko–Pastur Law [25]). *Let $Y \in \mathbb{R}^{p \times n}$ have i.i.d. entries with mean zero and variance one, and define the sample covariance matrix $S = \frac{1}{n}YY^\top$. Suppose $p/n \rightarrow c \in (0, \infty)$ as $n \rightarrow \infty$. Then the empirical spectral distribution μ_S of S converges almost surely to the **Marchenko–Pastur distribution** $\mu_{MP,c}$, which consists of:*

- *An absolutely continuous part supported on $[a, b]$, with density*

$$\frac{d\mu_{MP,c}(x)}{dx} = \frac{1}{2\pi cx} \sqrt{(b-x)(x-a)} \cdot \mathbf{1}_{[a,b]}(x),$$

where $a = (1 - \sqrt{c})^2$, $b = (1 + \sqrt{c})^2$;

- *And a point mass at 0 of size $1 - \frac{1}{c}$, if $c > 1$:*

$$\mu_{MP,c}(\{0\}) = 1 - \frac{1}{c}.$$

Remark 1.3.33. *The Marchenko–Pastur law describes the asymptotic eigenvalue density for sample covariance matrices in the null (no-signal) case. Notably, the spectrum becomes non-degenerate and spread over a compact support even when the population covariance is the identity.*

This distortion of the eigenvalue spectrum under the null hypothesis has profound implications for signal detection. Even a weak signal must overcome this inherent noise dispersion to be identifiable. This is formalized in the context of the *spiked covariance model*, which generalizes the MP law to settings with low-rank signals.

To illustrate the application of Marchenko–Pastur law, we take n independent obser-

vations from a linear model for $y \in \mathbb{R}^p$,

$$y = \mu + \beta x + \epsilon, \quad (1.32)$$

where $x, \epsilon \in \mathbb{R}$ are mean-zero and independent latent variables, and μ, β are deterministic unknowns, and $\text{var}(x) = 1$, $\text{var}(\epsilon) = \sigma^2 I_p$. The $p \times p$ population covariance matrix Σ takes the form

$$\Sigma = \beta\beta^\top + \sigma^2 I_p, \quad (1.33)$$

and the $p \times p$ sample covariance matrix takes the form

$$S = \frac{YY^\top}{n}, \quad (1.34)$$

where Y is $p \times n$ with columns being i.i.d observations of y . We denote the leading eigenvector of Σ as b , and $b = \beta / \|\beta\|$ with leading eigenvalue $\lambda^2 = \|\beta\|^2 + \sigma^2$. We also denote the leading eigenvector of S as h with the largest eigenvalue s^2 . The relationship between λ^2 and s^2 is stated in Paul [11]:

$$\begin{aligned} s^2 &\sim \lambda^2 \left(1 + \frac{\sigma^2 \phi}{\|\beta\|^2} \right) \\ &= \|\beta\|^2 + \sigma^2 + \sigma^2 \psi + \frac{\sigma^4 \phi}{\|\beta\|^2}, \end{aligned} \quad (1.35)$$

where $\phi = p/n$. Denote $\gamma = s^2 - \sigma^2(1 + \psi)$, We have $\|\beta\|^2 = \frac{\gamma + \sqrt{\gamma^2 - 4\sigma^4 \phi}}{2}$.

The above Y is a special model in which a low-rank signal is embedded in isotropic noise, i.e. the covariance of Y is a spiked covariance model.

Definition 1.3.34 (Spiked Covariance Model [2]). *Let $B \in \mathbb{R}^{p \times q}$ be a matrix with orthonormal columns, and let $\Lambda^2 \in \mathbb{R}^{q \times q}$ be a diagonal matrix with positive entries. The*

covariance matrix

$$\Sigma = B\Lambda B^\top + \sigma^2 I_p$$

is called a **spiked covariance model**, where the low-rank component $B\Lambda^2 B^\top$ represents a structured signal and $\sigma^2 I_p$ is isotropic noise.

Remark 1.3.35. The spectrum of Σ contains q eigenvalues (spikes) strictly greater than σ^2 , and $p - q$ eigenvalues equal to σ^2 . In the high-dimensional regime, the corresponding sample covariance matrix exhibits phase transitions: when a spike exceeds a critical threshold, its sample eigenvalue separates from the bulk, and the empirical eigenvector aligns with the true signal direction.

The BBP phase transition provides a precise characterization of eigenvalue separation and eigenvector alignment in spiked covariance models. The following theorem formalizes this behavior.

Theorem 1.3.36 (BBP Phase Transition [26]). Consider the spiked covariance model $\Sigma = I_p + \sum_{i=1}^r \theta_i v_i v_i^\top$, where $\theta_1 \geq \dots \geq \theta_r > 0$ are the spike strengths and $v_{i=1}^r$ are orthonormal vectors in $\mathbb{R}^p$. Let $X \in \mathbb{R}^{p \times n}$ be a data matrix whose columns are i.i.d. $N(0, \Sigma)$, and define the sample covariance matrix $S_n = \frac{1}{n} X X^\top$. Denote by $\ell_1 \geq \dots \geq \ell_p$ the eigenvalues of S_n , and by $\hat{v}_1, \dots, \hat{v}_r$ the eigenvectors corresponding to $\ell_1, \dots, \ell_r$.

Assume the high-dimensional asymptotic regime where $p, n \rightarrow \infty$ such that $p/n \rightarrow c \in (0, \infty)$. Then for each $i = 1, \dots, r$,

$$\ell_i \xrightarrow{\text{a.s.}} \begin{cases} (1 + \sqrt{c})^2, & \theta_i \leq \sqrt{c}, \\ \theta_i + \frac{c\theta_i}{\theta_i - 1}, & \theta_i > \sqrt{c}. \end{cases}$$

Remark 1.3.37. The critical threshold $\theta_{\text{crit}} = \sqrt{c}$ delineates two distinct regimes:

- **Subcritical regime** ($\theta_i \leq \sqrt{c}$): The signal is too weak relative to the noise and dimensionality. The sample eigenvalue ℓ_i converges to the edge of the Marchenko-Pastur bulk.
- **Supercritical regime** ($\theta_i > \sqrt{c}$): The signal is strong enough to be detectable. The sample eigenvalue ℓ_i separates from the bulk, converging to a value strictly larger than $(1 + \sqrt{c})^2$.

This phase transition has profound implications for principal component analysis in high dimensions: it quantifies the minimal signal strength required for reliable signal recovery and explains the systematic bias in eigenvalues observed in finite samples.

We now move on to the angle between eigenvectors b and h . In extension of [11], and denote $\rho^2 = \frac{\lambda^2}{\|\beta\|^2}$,

$$\begin{aligned}
\langle h, b \rangle^2 &\sim 1 - \frac{\phi \nu^2}{s^2} \rho^4 \\
&\sim 1 - \frac{\phi \nu^2 \rho^2}{\|\beta\|^2 + \phi \nu^2} \\
&\sim \frac{\|\beta\|^2 - \phi \nu^4 / \|\beta\|^2}{\|\beta\|^2 + \phi \nu^2} \\
&\sim \frac{1 - \frac{\phi \nu^4}{\|\beta\|^4}}{1 + \frac{\phi \nu^2}{\|\beta\|^2}}.
\end{aligned} \tag{1.36}$$

One can find similar results in [27]. To summarize,

	Phase Transition	Limit
Strong Inconsistency	$ \beta ^2 \leq \sqrt{\phi} \gamma^2$	$\langle h, b \rangle^2 \rightarrow 0$
Inconsistency	$\sqrt{\phi} \gamma^2 < \beta ^2 < +\infty$	$\lim_{p \rightarrow \infty} \langle h, b \rangle^2 \in (0, 1)$
Consistency	$\phi \gamma^2 / \beta ^2 \rightarrow 0$	$\langle h, b \rangle^2 \rightarrow 1$

The spherical law of cosines introduced by [28] reads as

$$\langle h, z \rangle^2 = \langle h, b \rangle^2 \langle b, z \rangle^2 \tag{1.37}$$

for $z = \frac{1}{\sqrt{p}}(1, 1, \dots, 1) \in \mathbb{R}^p$.

These results quantify the two key challenges in high-dimensional PCA: eigenvalue inflation ($s^2 > \lambda^2$) and eigenvector shrinkage ($(\langle h, b \rangle)^2 < 1$). This systematic bias and rotation motivate the development of shrinkage estimators that can correct for these effects, a central theme of this thesis. The geometric relationships extend beyond the signal direction itself; the angle between the estimated eigenvector h and any other vector z in the space can be derived accordingly.

1.3.3 High-Dimensional Geometry

In high-dimensional probability and random matrix theory, many asymptotic phenomena can be understood in terms of the geometry of the unit sphere. The distribution of eigenvectors of high-dimensional random matrices often converges to a uniform distribution over the sphere, and inner products between random vectors exhibit concentration around zero. To build a rigorous foundation for analyzing eigenvector behavior in high dimensions, we first formalize the notion of the unit sphere and its associated uniform (Haar) measure, which plays a central role in describing the asymptotic distribution of eigenvectors.

Definition 1.3.38 (Unit Sphere). *The Euclidean **unit sphere** in $\mathbb{R}^p$, denoted $\mathbb{S}^{p-1}$, is defined as*

$$\mathbb{S}^{p-1} := \{x \in \mathbb{R}^p : \|x\| = 1\}.$$

*Similarly, the **infinite dimensional unit sphere** is defined as*

$$S^\infty = \{x \in \mathbb{R}^\infty : \|x\| = 1\}.$$

Remark 1.3.39. *Since $(x_1, x_2, \dots, x_n, 0) \mapsto (x_1, x_2, \dots, x_n)$, we have $S^\infty = \bigcup_{n \geq 1} S^{n-1}$.*

Definition 1.3.40 (Uniform Measure on the Sphere [29]). *The uniform probability measure on $\mathbb{S}^{p-1}$ is the unique probability measure that is invariant under the action of the orthogonal group $\mathcal{O}(p)$. That is, if $u \sim \text{Unif}(\mathbb{S}^{p-1})$, then for any orthogonal matrix $Q \in \mathbb{R}^{p \times p}$,*

$$Qu \sim \text{Unif}(\mathbb{S}^{p-1}).$$

*This measure is also known as the **Haar measure** on the sphere.*

Remark 1.3.41. *A random vector $u \sim \text{Unif}(\mathbb{S}^{p-1})$ can be generated by drawing $z \sim \mathcal{N}(0, I_p)$ and setting*

$$u = \frac{z}{\|z\|_2}.$$

This is because the standard multivariate normal distribution is spherically symmetric and invariant under orthogonal transformations.

Remark 1.3.42. *In high dimensions, two independent vectors $u, v \sim \text{Unif}(\mathbb{S}^{p-1})$ satisfy*

$$\langle u, v \rangle \rightarrow 0, \quad \text{for } p \uparrow \infty,$$

with high probability. This reflects the fact that random unit vectors are almost orthogonal in high dimensions.

Remark 1.3.43. *In this thesis, we will encounter several instances where the empirical eigenvectors of a random matrix behave asymptotically like vectors drawn uniformly from $\mathbb{S}^{p-1}$. The Haar measure serves as the limiting distribution of eigenvectors in both the null Wishart case and certain spiked models.*

An alternative way of defining the uniform measure μ_p on $\mathbb{S}^{p-1}$ that we use to determine the existence of "limiting measure" of μ_p .

Definition 1.3.44 ([30]). *If $x \in \mathbb{R}^p \setminus \{0\}$, the polar coordinates of x are*

$$r = |x| \in (0, \infty), \quad x' = \frac{x}{|x|} \in S^{p-1}$$

The introduction of polar coordinates not only provides a geometric intuition for points in Euclidean space but also lays the groundwork for a measure-theoretic treatment of spherical geometry. By decomposing vectors into their radial and angular components, we can describe the Lebesgue measure on $\mathbb{R}^p$ in terms of a product measure over radius and sphere. This formulation is crucial for defining the uniform surface measure on S^{p-1} in a way that is compatible with integration and limiting procedures. We now make this correspondence precise via the map Φ , which establishes a bijective link between $\mathbb{R}^p \setminus \{0\}$ and $(0, \infty) \times S^{p-1}$.

Definition 1.3.45 ([30]). $\Phi(x) = (r, x')$ is a continuous bijection from $\mathbb{R}^p \setminus \{0\}$ to $(0, \infty) \times S^{p-1}$ whose inverse is $\Phi^{-1}(r, x') = rx'$

Definition 1.3.46 ([30]). Denote by m^* the Borel measure on $(0, \infty) \times S^{p-1}$ induced by Φ from Lebesgue measure on $\mathbb{R}^n$, i.e. $m^*(E) = m(\Phi^{-1}(E))$. Define measure $\rho = \rho_p$ on $(0, \infty)$ by $\rho(E) = \int_E r^{p-1} dr$.

Theorem 1.3.47 ([30]). *There is a unique Borel measure $\sigma = \sigma_{p-1}$ on S^{p-1} such that $m^* = \rho \times \sigma$. If f is Borel measurable on $\mathbb{R}^p$ and $f \geq 0$ or $f \in L^1(m)$, then*

$$\int_{\mathbb{R}^p} f(x) dx = \int_0^\infty \int_{S^{p-1}} f(rx') r^{p-1} d\sigma(x') dr$$

Definition 1.3.48 (Uniform measure on S^{p-1} [30]). $\mu_p = \frac{\sigma}{\sigma(S^{p-1})}$ is the **uniform measure** on S^{p-1} , where $\sigma(S^{p-1}) = \frac{2\pi^{p/2}}{\Gamma(p/2)}$.

Theorem 1.3.49 (Kolmogorov Extension Theorem [31]). *Let T denote some interval, and let $n \in \mathbb{N}$. For each $k \in \mathbb{N}$ and finite sequence of distinct times $t_1, \dots, t_k \in T$,*

let $\nu_{t_1 \dots t_k}$ be a probability measure on $(\mathbb{R}^n)^k$. Suppose that these measures satisfy two consistency conditions:

1. for all permutations π of $\{1, 2, \dots, k\}$ and measurable sets $F_i \subseteq \mathbb{R}^n$,

$$\nu_{t_{\pi(1)} \dots t_{\pi(k)}}(F_{\pi(1)} \times \dots \times F_{\pi(k)}) = \nu_{t_1 \dots t_k}(F_1 \times \dots \times F_k);$$

2. for all measurable sets $F_i \subseteq \mathbb{R}^n$, $m \in \mathbb{N}$

$$\nu_{t_1 \dots t_k}(F_1 \times \dots \times F_k) = \nu_{t_1 \dots t_k, t_{k+1}, \dots, t_{k+m}} \left(F_1 \times \dots \times F_k \times \underbrace{\mathbb{R}^n \times \dots \times \mathbb{R}^n}_m \right).$$

Then there exists a probability space $(\Omega, \mathcal{F}, \mathbb{P})$ and a stochastic process $X : T \times \Omega \rightarrow \mathbb{R}^n$ such that

$$\nu_{t_1 \dots t_k}(F_1 \times \dots \times F_k) = \mathbb{P}(X_{t_1} \in F_1, \dots, X_{t_k} \in F_k)$$

for all $t_i \in T$, $k \in \mathbb{N}$ and measurable sets $F_i \subseteq \mathbb{R}^n$, i.e. X has $\nu_{t_1 \dots t_k}$ as its finite-dimensional distributions relative to times $t_1 \dots t_k$.

Remark 1.3.50. An extension of any consistent family of probability measures $\{(\mathbb{R}^n, \mathcal{B}_n, \mathbb{P}_m)\}_{m=1}^\infty$ to a probability $\mathbb{P}(\cdot)$ on $(\mathbb{R}^\infty, \mathcal{B}_\infty)$ necessarily exists, and it is unique.

While the geometry of high-dimensional spheres provides a static picture of eigenvector distribution, certain asymptotic phenomena, particularly those involving sequences of empirical processes, call for a functional viewpoint. In such cases, the limiting behavior of high-dimensional objects is best described using tools from stochastic process theory. A key player in this context is the Wiener measure, which governs the law of Brownian motion and serves as a canonical example of a Gaussian process. In Chapter 2, we show that under suitable high-dimensional asymptotics, the empirical distribution of eigenvector coordinates can be viewed as converging (in distribution) to a Wiener process. To

prepare for that, we review the construction of the Wiener measure and the foundational Donsker's theorem that guarantees this convergence under proper scaling.

Definition 1.3.51 ([32]). Denote $C_0[0, 1]$ be the collection of all continuous function f defined on $[0, 1]$ such that $f(0) = 0$. The **Wiener measure** μ_W on $C_0[0, 1]$ is defined by

$$\mu_W(\{W : W(t) - W(s) \in D\}) = \frac{1}{(2\pi(t-s))^{1/2}} \int_D \exp\left(\frac{-\omega^2}{2(t-s)}\right) d\omega,$$

for $s < t$, W the Wiener process and Borel set $D \subset \mathbb{R}$.

Theorem 1.3.52 (Donsker's Theorem [33]). Let $X_1, X_2, \dots$ be a sequence of i.i.d. random variables with mean 0 and variance 1. Let $S_n = \sum_{i=1}^n X_i$, and $W_n(t) = \frac{S_{\lfloor nt \rfloor}}{\sqrt{n}}$, $t \in [0, 1]$. Then $W_n = (W_n(t))_{t \in [0, 1]} \Rightarrow W$, the Wiener process as $n \rightarrow \infty$.

1.4 Thesis Overview

This thesis investigates two interrelated problems in high-dimensional statistics under the asymptotic regime where the dimension $p \rightarrow \infty$ while the sample size n remains fixed. The first is the limiting distribution of sample eigenvectors under the null model, and the second is the behavior of shrinkage-based covariance estimators when the assumed model rank is misspecified. Both problems are examined through the lens of geometric intuition and asymptotic analysis.

Chapter 2 focuses on the asymptotic behavior of eigenvectors in the classical null Wishart setting. We begin by constructing the Wiener measure on the space $C_0[0, 1]$, following the formulation of [32], and invoke a result from [34] to establish a rigorous connection between the uniform measure on the infinite-dimensional sphere $\mathbb{S}^\infty$ and the Wiener measure. Using this framework, we prove that the eigenvectors of the sample covariance matrix $S = \frac{1}{n}YY^\top$, where the columns of Y are i.i.d. from $\mathcal{N}(0, \sigma^2 I_p)$,

converge in distribution to the Haar measure on $\mathbb{S}^\infty$ as $p \rightarrow \infty$, provided the eigenvectors lie outside the null space of S . This result formalizes the intuition that random high-dimensional eigenvectors become uniformly distributed over the unit sphere in the infinite-dimensional limit.

Chapter 3 turns to the analysis of shrinkage estimators under model misspecification. Specifically, we examine James–Stein-type shrinkage estimators of covariance matrices in the spiked model setting, where the true rank q is unknown and may be incorrectly specified by the user through an assumed rank K . Our main findings can be summarized as follows:

1. When the assumed rank K exceeds the true rank q , the optimization bias induced by the shrinkage estimator vanishes as $p \rightarrow \infty$. This suggests that moderate over-specification does not harm asymptotic performance in the high-dimensional setting.
2. When K underestimates the true rank q , the optimization bias persists and does not vanish in the limit, indicating a fundamental sensitivity to under-specification that must be carefully addressed in practice.
3. In the important special case of single-factor shrinkage ($K = 1$), we derive explicit asymptotic expressions for the signal-to-noise ratio ψ^2 , the squared projection error between empirical and population eigenvectors, and closed-form upper bounds on the resulting optimization bias.

Together, these two chapters reveal deep connections between high-dimensional geometry and the behavior of statistical procedures under structural misspecification. The results contribute both new theoretical insights and practical guidance for the use of shrinkage methods in high-dimensional inference.

Chapter 2

The Quadratic Optimization Bias of Misspecified PCA and Associated James-Stein Estimators.

Given a random vector $y \in \mathbb{R}^p$, a factor model is given by:

$$y = Bx + \varepsilon,$$

Here, $B \in \mathbb{R}^{p \times q}$ is the factor loading matrix, $x \in \mathbb{R}^q$ is a vector of latent variables, and $\varepsilon \in \mathbb{R}^p$ denotes idiosyncratic noise. This modeling framework originated from [35], and [36, 37] developed this technique to multiple factors. Now it has been extensively generalized in subsequent research, becoming a cornerstone for dimension reduction and covariance estimation in high-dimensional statistics, biological sciences and other areas (see, e.g. [38][39][40]). Under this model, the population covariance matrix $\Sigma = \text{Cov}(y)$

exhibits a distinctive low-rank plus sparse structure:

$$\Sigma = \underbrace{BB^\top}_{\text{Low-rank signal}} + \underbrace{\Gamma}_{\text{Isotropic noise}}.$$

In practice, we observe n independent samples $y_1, y_2, \dots, y_n$ from this model, which assemble into a $p \times n$ data matrix Y . A central empirical objective is to construct an effective estimator $\hat{\Sigma}$ from Y that accurately recovers Σ .

The theoretical justification for using principal component analysis (PCA) in factor models was fundamentally established by [41]. They demonstrated that, for the population covariance matrix Σ , PCA can recover the factor structure $BB^\top$ as $p \rightarrow \infty$, provided the first q eigenvalues of Σ diverge with p while the rest remain bounded. This provides the fundamental insight that the dominant eigenvectors of Σ span the true factor space $BB^\top$. However in practice, we must work with the sample covariance matrix $S = \frac{1}{n}YY^\top$. This leads to the standard PCA approach of using the eigenvectors H and the corresponding estimate is $\hat{\Sigma} = HH^\top + G$, with G being some residual corresponding to Γ . This framework encompasses various regularized estimators, including the POET estimator proposed by [42], which imposes a sparse structure on G .

Note that this classic consistency result from [41] relies on $p \rightarrow \infty$ with n fixed or growing slowly, and assumes the population covariance is known. In modern high-dimensional settings where $p/n \rightarrow \gamma > 0$, the sample eigenvectors exhibit a phase transition. [11] showed that the corresponding sample eigenvectors may fail to converge to their population counterparts if the signal strength is insufficient relative to the noise and dimensionality, and the column space of sample eigenvectors are not consistent estimators to $BB^\top$. The phenomenon of PCA inconsistency in high dimensions was first established by [43], who showed that under a bounded eigenvalue condition, the sample principal component directions become inconsistent estimators for their population

counterparts when $p/n \rightarrow \gamma \in (0, \infty)$. Later on, [12] extended the result and showed that the inconsistency holds when $p/(n\lambda) \rightarrow c \in (0, \infty)$, where λ is the eigenvalue. This result is also further discussed by [44]. Conversely, when λ is bounded, or $p/(n\lambda) \rightarrow 0$, then PCA consistently estimates $BB^\top$. See [45][46][38] as example. While the above results concern proportional growth $p/n \rightarrow \gamma \in (0, \infty)$, many modern applications operate in the HDLSS regime, with n fixed and $p \rightarrow \infty$. Under Spiked covariance setting (see [2])[47][48][49] showed the consistency of leading eigenvectors to population counterpart. However, spiked covariance model implies the limit of $B^\top B/p$ being not exist. When the limit exists, [50] showed maximal sample eigenvector is strongly inconsistent.

Despite these fundamental statistical limitations, eigenvector-based optimization remains at the heart of many critical decision problems, such as portfolio optimization [51], spatial filtering in signal processing [52], optimal fingerprinting in climate science [53], spectral graph partitioning [54]. Those problems can all be formulated as a (constrained) quadratic programming problem

$$\max_x \frac{1}{2} \langle x, \Sigma x \rangle \quad \text{subject to } A^\top x \leq c,$$

where A is the constraint matrix. Since the true population covariance Σ is unknown, it is typically replaced by a sample estimate $\hat{\Sigma}$. However, as [6] pointed out, “optimized” results based on the sample covariance are often unreliable and reflect estimation error more than genuine performance. This issue was formulated by [16] within terms of optimization bias $\mathcal{E}_p(\cdot)$, which is a vector-valued function in $\mathbb{R}^q$ of the eigenvector matrix H . In the special case where the constraint is a vector ($A = z$), they proposed a new estimator $H_\sharp$ that asymptotically eliminates $\mathcal{E}_p(H_\sharp)$ as $p \rightarrow \infty$. $H_\sharp$ takes form of the following:

$$H_\sharp = H\Psi^2 + \frac{(I_p - HH^\dagger)z}{1 - z^\top HH^\dagger z} z^\top H(I_q - \Psi^2)$$

where Ψ^2 is a diagonal matrix computed from eigenvalues of sample covariance matrix and $H^\dagger = (H^\top H)^{-1} H^\top$ is the pseudoinverse of H . They proved that although H yields a persistent bias $\mathcal{E}_p(H) \not\rightarrow 0$, $\mathcal{E}_p(H_\sharp) \rightarrow 0$ almost surely as $p \rightarrow \infty$ with n fixed.

This finding resonates with a fundamental insight from high-dimensional estimation known as Stein’s paradox [55]: in high dimensions, the usual unbiased estimator (i.e. sample mean) can be suboptimal under squared-error loss. Later, [56] generalized the decision-theoretic understanding of admissibility. He introduced an associated diffusion process for each estimator and showed that its admissibility is equivalent to the recurrence of the diffusion. One solution proposed by [14] is the James–Stein estimator, which introduces some degree of shrinkage toward a fixed target and achieves a uniformly lower mean-squared error. [57] provided an empirical Bayes interpretation of the James–Stein estimator, framing Stein’s result as part of a broader empirical Bayes principle for high-dimensional problems. [58] provided an intuitive explanation for the efficiency gain of shrinkage without appealing to Bayesian arguments. [59] further quantified the improvement of this new estimator by analyzing the limiting risk of James–Stein estimators, showing that the advantage of shrinkage becomes increasingly pronounced with higher dimension. [60] provided a historical interpretation that reframed the James–Stein estimator as a natural consequence of regression to the mean and multi-parameter estimation. [61] developed a geometrical explanation for this inadmissibility. Fourdrinier, [62] provided a comprehensive and modern synthesis of Stein-type shrinkage estimation from both frequentist and Bayesian perspectives.

Although our setting is different, the idea of Stein’s paradox and James–Stein estimator is directly relevant to our findings: allowing a small deviation from the sample covariance eigenvectors H can reduce the overall error in estimating their population counterparts. As a beginning of this storyline, [63] proposed a James–Stein estimator specifically for the leading eigenvector and clarified the connection between James–Stein

estimator and James-Stein for eigenvectors. [64] further introduced and analyzed the JSE estimator, and [65] proved its consistency for the leading eigenvector when additional information is available. Later, [66] listed the detailed recipe for the JSE estimator and showed it improves the volatility ratio of large optimized portfolios in terms of $x^\top \hat{\Sigma} x / x^\top \Sigma x$ being closer to 1.

Based on this line of work, we can extend the shrinkage idea to more general eigenvector settings by pulling H toward a structured subspace defined by problem-specific constraints. Conceptually, this corresponds to estimating the projection of the population eigenvectors B onto the joint space spanned by the sample directions H and the constraint matrix A . To apply the shrinkage idea to H , for any constraint matrix A , define $M = AA^\dagger H$. M is the projection of H onto $\text{col}(A)$. Following the James-Stein framework, M here is the shrinkage target, and the estimator is

$$H_{JS} = HC + M(I_q - C),$$

where C is the shrinkage coefficient matrix computed from the sample eigenvalues. In Section 2.2 and 2.4.4 we stated two results: (1) $H_\sharp$ is equivalent to H_{JS} when letting $A = z$; (2) H_{JS} is also a asymptotic root of the optimization bias equation, i.e. $\mathcal{E}_p(H_{JS}) \rightarrow 0$ almost surely as $p \rightarrow \infty$ and fixed n .

However, this approach rests on a strong and impractical assumption: the true number of factors q , i.e. $\text{rank}(B)$ is known. In reality, q is unknown and must be estimated by some K . We therefore analyze the behavior of this estimator under $K \neq q$. In Section 2.4.1 and 2.4.2, we show that (1) when $K < q$, $\mathcal{E}_p(H_\sharp) \not\rightarrow 0$ as $p \rightarrow \infty$. It implies $H_\sharp$ is necessarily inconsistent and fails to control the optimization bias. (2) when $K > q$, we have $\mathcal{E}_p(H_\sharp) \rightarrow 0$ almost surely as $p \rightarrow \infty$ and n fixed. It's showing $H_\sharp$ is consistent and optimization bias goes to 0. Therefore in practice, we want to choose K as large as

possible.

2.1 Literature Review

Research on factor models and principal component analysis (PCA) has produced a vast literature on eigenstructure estimation in high dimensions. Currently, most existing results rely on asymptotic regimes where both $p, n \rightarrow \infty$ with $p/n \rightarrow c$. In contrast, many modern biomedical and imaging datasets operate in the high-dimension low-sample-size (HDLSS) setting, where n is small or fixed while p is extremely large. In this regime, sample covariance matrices are ill-conditioned, and sample eigenvectors behave unpredictably. [49] and [67] show that in HDLSS data with mixing-type dependency the observations concentrate on a low-dimensional surface of the unit sphere as the dimension grows, and they propose a noise-reduction methodology to consistently estimate eigenvalues and PC directions in this regime. These findings motivated the development of shrinkage-based and regularized PCA methods designed to stabilize eigenvector estimation in HDLSS data.

A major contribution in this direction comes from Ledoit and Wolf [3, 15]. In their paper, they proposed linear and nonlinear covariance-shrinkage estimators that combine the sample covariance with structured targets. [68] extended these ideas to elliptical distributions, providing robust shrinkage estimators applicable beyond Gaussian assumptions. [69] introduced the POET(Principal Orthogonal Complement Thresholding) framework, which decomposes the covariance matrix into a low-rank factor component and a sparse idiosyncratic component via adaptive thresholding. Later papers by Fan and collaborators [70, 71] synthesized robust, factor-adjusted shrinkage methodologies and connected them to modern applications in finance and network analysis. Collectively, these works establish that appropriate shrinkage can mitigate HDLSS instability, but they typically

assume the correct factor rank is known, while it's an assumption rarely met in practice.

Determining the true number of latent factors q has therefore become a major research topic. The seminal information-criterion approach of [38] uses penalized loss functions to estimate q consistently as $p, n \rightarrow \infty$. [72] refined this criterion for robustness, and [73] proposed a residual-sum-of-squares rule minimizing prediction loss. More recently, [74] generalized information-criterion selection to mixed or discrete outcomes within high-dimensional latent-factor frameworks, establishing consistency even with missing data. Other approaches employ formal hypothesis tests: for example, [75] developed a likelihood-based test on the empirical spectral distribution, while [76] extended resampling-based significance tests to high-dimensional factor models. Spectral-gap and eigenvalue-ratio statistics [77, 78, 79] offer computationally simple alternatives by identifying sharp breaks between signal and noise eigenvalues. Related analyses by [80] examined the asymptotic effects of over-estimating the factor number. These procedures consistently estimate q under various asymptotics, but little is known about how errors in rank selection ($K \neq q$) affect subsequent eigenvector or shrinkage estimation, especially in HDLSS contexts.

Beyond determining the number of factors, a parallel line of work investigates the optimal estimation of high-dimensional eigenspaces themselves, particularly under sparsity or privacy constraints. [81] studied minimax rates for estimating both spiked covariance matrices and principal subspaces, establishing optimal convergence bounds under spectral norms. [82] extended these results to sparse subspace estimation without relying on the spiked covariance assumption, and provided nearly optimal nonasymptotic minimax rates. More recently, [83] investigated differentially private PCA in the spiked covariance model, characterizing the sensitivity of eigenvalues and eigenvectors and deriving matching minimax bounds under privacy constraints.

However, only a limited number of studies have explored the impact of rank misspec-

ification on PCA. [2] and [11] analyzed eigenvalue and eigenvector bias under spiked covariance models, showing that sample eigenvectors become inconsistent when the signal-to-noise ratio is weak or when p/n is large. [84] quantified finite-sample bias [80] studied over-estimated factor numbers in high-dimensional settings, showing increased variance and reduced identifiability. Yet these analyses remain confined to unregularized PCA and do not address how shrinkage estimators behave when the assumed rank is incorrect. In the HDLSS regime, where the sample covariance matrix is nearly singular, such misspecification can severely distort optimization geometry and bias the estimated eigendirections. This chapter builds on this gap by theoretically characterizing the optimization bias and discrepancy of James–Stein shrinkage eigenvector estimators under rank misspecification, providing new insight into the robustness of shrinkage methods when both dimensionality and model uncertainty are high.

2.2 Motivation

2.2.1 Problem Formulation and Theoretical Limits

We begin with the $p \times p$ population covariance matrix Σ with decomposition

$$\Sigma = BB^\top + \Gamma, \quad (2.1)$$

where Γ is a $p \times p$ invertible and symmetric matrix. We further denote $\mathcal{B}_{p \times q}$ as the k leading principal components of Σ , i.e. $\Sigma = \mathcal{B}\Lambda_p^2\mathcal{B}^\top + \Gamma$, where Λ_p^2 is the $q \times q$ diagonal matrix with diagonal elements being the q leading eigenvalues of Σ ranked of the largest to the smallest. The following approximate factor model presented below has become a standard in the literature.

Assumption 2.2.1. *The matrices $B = B_{p \times q}$ and $\Gamma = \Gamma_{p \times p}$ satisfy the following:*

$$(a) \ 0 < \liminf_{p \uparrow \infty} \inf_{|v|=1} \langle v, \Gamma v \rangle < \overline{\lim}_{p \uparrow \infty} \sup_{|v|=1} \langle v, \Gamma v \rangle < \infty.$$

$$(b) \ \lim_{p \uparrow \infty} B^\top B/p \text{ exists as an invertible } q \times q \text{ matrix with fixed } q \geq 1.$$

Starting from that, we can consider the general quadratic programming problem with constraint:

$$\begin{aligned} \min_{x \in \mathbb{R}^p} \frac{1}{2} \langle x, \Sigma x \rangle \\ A^\top x \leq c. \end{aligned} \tag{2.2}$$

where $A \in \mathbb{R}^{p \times r}$. A further satisfies the following assumption:

Assumption 2.2.2. *The column spaces of $A \in \mathbb{R}^{p \times r}$ and $B \in \mathbb{R}^{p \times q}$ satisfy:*

$$(a) \ \liminf_{p \rightarrow \infty} \|AA^\dagger - BB^\dagger\|_F^2 > 0.$$

$$(b) \ \text{There exists } \xi > 0 \text{ such that for all large } p \text{ and any } \ell \in \mathbb{R}^r,$$

$$\|(I - BB^\dagger)A\ell\|^2 \geq \xi p \|\ell\|^2.$$

The solution of Eq 2.2 is the maximizer of the Lagrangian

$$Q(x) = c_0 + c_1 \langle x, \zeta \rangle - \frac{1}{2} \langle x, \Sigma x \rangle, \tag{2.3}$$

where

$$\zeta = \sum_{i=1}^r \ell_i a_i \in \mathbb{R}^p, \tag{2.4}$$

with a_i the columns of A and ℓ_i the right multipliers. The maximizer $\hat{x} = c_1 \Sigma^{-1} \zeta$ is obtained by solving

$$c_1 \zeta - \Sigma x = 0,$$

and

$$\max_{x \in \mathbb{R}^p} Q(x) = c_0 + \frac{c_1^2}{2} \mu_p^2, \quad \mu_p^2 = \langle \zeta, \Sigma^{-1} \zeta \rangle.$$

Theorem 2.2.3. *Assume Assumption 2.2.1 holds, and there exists $\xi > 0$ such that for all large p and any $\ell \in \mathbb{R}^r$,*

$$\|(I - BB^\dagger)A\ell\|^2 \geq \xi p \|\ell\|^2.$$

then $\max_x Q(x) \rightarrow +\infty$

Remark 2.2.4. *The condition $\|(I - BB^\dagger)A\ell\|^2 \geq \xi p \|\ell\|^2$ ensures a uniform lower bound on the projection of A outside $\text{col}(B)$, i.e. $\text{col}(A)$ and $\text{col}(B)$ are not asymptotically collapsing.*

2.2.2 Discrepancy, Optimization Bias and Its Geometric Source

In practice, we don't know Σ , so we estimate $\hat{\Sigma}$ and replace Σ by $\hat{\Sigma}$ in Eq 2.2, with the Lagrangian

$$\hat{Q}(x) = c_0 + c_1 \langle x, \zeta \rangle - \frac{1}{2} \langle x, \hat{\Sigma} x \rangle. \quad (2.5)$$

The sample covariance matrix $\hat{\Sigma}$ can also be written in spectral decomposition, i.e.,

$$\hat{\Sigma} = \sum_{(s^2, h)} s^2 h h^\top = H H^\top + \gamma^2 I_p. \quad (2.6)$$

We further assume H has properties consistent B :

Assumption 2.2.5. *For $\zeta_H = H H^\dagger \zeta$, suppose $H = H_{p \times q}$ and $\zeta = \zeta_{p \times 1}$ satisfy $\overline{\lim}_{p \uparrow \infty} |\zeta_H|/|\zeta| < 1$ and $\underline{\lim}_{p \uparrow \infty} |\zeta| \neq 0$. Also $\lim_{p \uparrow \infty} H^\top H/p$ exists as a $q \times q$ invertible matrix and for $v = A\alpha$, $(I - H H^\dagger)v \neq 0$.*

Eq 2.5 attains $\max_{x \in \mathbb{R}^p} \hat{Q}(x) = \hat{Q}(\hat{x}) = c_0 + \frac{c_1^2 \hat{\mu}_p^2}{2}$ at $\hat{x}$ and with $\hat{\mu}_p^2 = \langle \zeta, \hat{\Sigma}^{-1} \zeta \rangle$. Plug $\hat{x}$ into Eq 2.3, we have

$$Q(\hat{x}) = c_0 + \frac{c_1^2 \hat{\mu}_p^2}{2} \hat{D}_p, \quad (2.7)$$

where

$$\hat{D}_p = 2 - \frac{\langle \hat{x}, \Sigma \hat{x} \rangle}{c_1^2 \hat{\mu}_p^2} = 2 - \frac{\zeta^\top \hat{\Sigma}^{-1} \Sigma \hat{\Sigma}^{-1} \zeta}{\zeta^\top \hat{\Sigma}^{-1} \zeta} = 2 - \frac{\alpha^\top (A^\top \hat{\Sigma}^{-1} \Sigma \hat{\Sigma}^{-1} A) \alpha}{\alpha^\top (A^\top \hat{\Sigma}^{-1} A) \alpha}, \quad \alpha = \frac{\ell}{|\ell|}. \quad (2.8)$$

$\hat{D}_p$ is called discrepancy, and it captures the out-of-sample deviation between the empirical and population objectives. In general, we have

Theorem 2.2.6. *If Assumption 2.2.1 and 2.2.5 holds, and there exists $\delta > 0$ such that*

$$\lambda_{\min}(B^\top H H^\top B/p) \geq \delta \quad (\text{T})$$

as $p \uparrow \infty$. Then $\hat{D}_p \rightarrow -\infty$ and $Q(\hat{x}) \rightarrow -\infty$.

Remark 2.2.7. *Note that the realized maximum $Q(\hat{x})$ tends to $-\infty$ but the true maximum $\max_x Q(x)$ tends to $+\infty$.*

One quick observation is that when we take $\hat{\Sigma} = \Sigma$,

$$\hat{D}_p = 2 - \frac{\alpha^\top (A^\top \Sigma^{-1} \Sigma \Sigma^{-1} A) \alpha}{\alpha^\top (A^\top \Sigma^{-1} A) \alpha} = 1.$$

To further understand the structural source of this deviation, we relate it to the geometric misalignment between the empirical eigenspace of $\hat{\Sigma}$ and the population eigenspace of Σ .

Consider $B, H \in \mathbb{R}^{p \times q}$ of the spectral decomposition of Σ and $\hat{\Sigma}$:

$$\Sigma = B B^\top + \Gamma; \quad \hat{\Sigma} = H H^\top + \hat{\gamma}^2 I.$$

the optimization bias along direction $A\alpha$ is defined as:

$$\mathcal{E}_p(H, A, a) = \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha) = \frac{\mathcal{B}^\top \mathcal{A} - (\mathcal{B}^\top \mathcal{H})(\mathcal{H}^\top \mathcal{A})}{I - \mathcal{A}^\top \mathcal{H} \mathcal{H}^\top \mathcal{A}} \alpha, \quad (2.9)$$

where $\mathcal{B}$ is the orthogonalization of B with $\Sigma = \mathcal{B} \Lambda_p^2 \mathcal{B} + \Gamma$, and $\mathcal{H}$ and $\mathcal{A}$ are orthogonalization of H and A respectively. Each constraint column a_i of A defines a bias term $\mathcal{E}_p(H, A, a_i)$ that measures the misalignment between the empirical eigenspace spanned by H and the population eigenspace spanned by B when projected along a_i . Since the feasible direction $A\alpha = \sum_i \alpha_i a_i$ is a linear combination of these columns, the total bias contribution combines their effects quadratically.

Remark 2.2.8. *Alternatively, we can write optimization bias in forms of $Q(\cdot)$ from Equation 2.3, where $\zeta = \sum_{i=1}^r a_i \ell_i$. This leads to the equivalent discussion on $\mathcal{E}_p$, and we will come back to this in Section 2.3.*

2.2.3 Decomposition and Shrinkage Estimator Framework

Noting that

$$\begin{aligned} \|\Lambda_p \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha)\|^2 &= \left\langle \Lambda_p \sum_i \alpha_i \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_i), \Lambda_p \sum_j \alpha_j \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_j) \right\rangle \\ &= \sum_{i,j} \alpha_i \alpha_j \langle \Lambda_p \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_i), \Lambda_p \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_j) \rangle. \end{aligned}$$

then following [16], $\hat{D}_p$ can decompose as a sum over all constraint directions:

Theorem 2.2.9. *Suppose regularity conditions Assumption 2.2.1 and 2.2.5 holds. For*

$$u_H = \frac{A\alpha - HH^\top A\alpha}{|A\alpha - HH^\top A\alpha|},$$

$$\hat{D}_p = -\frac{\alpha^\top Q \alpha}{\hat{\gamma}^2} + \left(2 - \frac{\langle u_H, \Gamma u_H \rangle}{\hat{\gamma}^2} \right) + O(\|\mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha)\| + \|\mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha)\|^2 + 1/p), \quad (2.10)$$

where

$$Q = (\langle \Lambda_p \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_i), \Lambda_p \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha_j) \rangle)_{i,j=1}^r$$

Remark 2.2.10. *The first term in Eq 2.10 represents a weighted sum of optimization biases across all constraint directions.*

Lastly, define the orthogonal projector be

$$P = AA^\dagger, \quad P_\perp = I_p - P.$$

For any matrix $H \in \mathbb{R}^{p \times q}$, we decompose

$$H = PH + P_\perp H =: M + R_H,$$

where $M = AA^\dagger H$ denotes the projection of H onto $\text{col}(A)$, and $R_H = P_\perp H$ is the orthogonal complement. Because $P_\perp M = (I_p - AA^\dagger)AA^\dagger H = 0$, all columns of M lie strictly within the constraint subspace.

For the projected estimator $M = \mathcal{A}\mathcal{A}^\top \mathcal{H}$, although its columns lie in $\text{col}(A)$, using M alone as an estimator may lose information outside $\text{col}(A)$ that is useful for overall estimation. Instead, we consider a linear combination of M and H :

$$H_{JS} = HC + M(I_q - C),$$

where C is a coefficient matrix. It is called James-Stein estimator and will make $\hat{D}_p$ bounded and asymptotically stable. In the next section, we will construct a new shrinkage estimator and study its behavior under rank misspecification in the simpler case by choosing $A = z \in \mathbb{R}^{p \times 1}$.

2.2.4 Application

To illustrate some examples in practice, we can consider the following minimum variance problem:

$$\begin{aligned} \min_{w \in \mathbb{R}^p} \quad & \langle w, \hat{\Sigma} w \rangle \\ \text{s.t.} \quad & \langle w, \zeta \rangle = 1 \\ \text{where} \quad & \zeta = \mathbf{1}_p \end{aligned} \tag{2.11}$$

The minimizer $\hat{w}$ corresponds to the weights of a minimum variance portfolio of financial assets, and $\zeta = \mathbf{1}_p$ represents the full invest constraint. The minimum of Equation 2.11 represents the variance of the estimated portfolio $\hat{w}$, and the out-of-sample variance (using population Σ instead of sample $\hat{\Sigma}$) is:

$$V_p^2 = \langle \hat{w}, \Sigma \hat{w} \rangle.$$

[16] proved that $\hat{D}_p = 2 - \hat{\mu}_p^2 V_p^2$, and further,

$$V_p^2 = \frac{|\Lambda_p + \mathcal{E}_p(H; \zeta/|\zeta|)|^2}{p|z - z_H|^2} + O(1/p).$$

We can also observe that $|\mathcal{E}_p(H; \zeta/|\zeta|)| \rightarrow 0$ leads to asymptotically riskless portfolio.

2.3 Model Setup and Problem Formulation

We now develop the main theoretical framework for the proposed shrinkage estimator $H_{\sharp}$ and establish its connection to the classical James–Stein (JS) estimator of eigenvectors. Without loss of generality, we focus on the case $A = z \in \mathbb{R}^p$ in all of the theorems. The reason to allow this simplification in the proofs is the optimizer $\hat{x} = c_1 \hat{\Sigma}^{-1} \zeta$ from Eq. 2.3 depends on A only through the composite direction $\zeta = A\ell$. Thus, any rank-

r constraint system is directionally equivalent to a one-dimensional effective constraint along ζ . For convenience, we will use $\mathcal{E}_p(H)$ denote the optimization bias of H for (fixed) single direction.

Under this simplified setting, we formalize the data-generating model and asymptotic assumptions, construct the shrinkage estimator $H_{\sharp}$, and analyze its behavior under rank misspecification. Finally, we show that $H_{\sharp}$ is a special case of James-Stein estimator.

2.3.1 Model Setup

Let Y denote a $p \times n$ data matrix of p variables and n observations. For a random $q \times n$ matrix X and random $p \times n$ matrix ε , Y follows the linear model

$$Y = BX + \varepsilon. \quad (2.12)$$

The $p \times q$ matrix B is the unknown factor, X is the matrix of latent variables, ε represents the noise, and only Y is observed.

We start with the following assumptions.

Assumption 2.3.1.

- (a) *Only Y is observed.*
- (b) *The true number of factors q is unknown with $n > q$.*
- (c) *$X^\top X$ is $q \times q$ invertible almost surely.*
- (d) *$\lim_{p \uparrow \infty} (B^\top B)/p = \Sigma_B$ exists as an invertible $q \times q$ matrix with fixed $q \geq 1$.*
- (e) *$\lim_{p \uparrow \infty} \varepsilon^\top \varepsilon / p = \gamma^2 I_n$ almost surely for some constant $\gamma > 0$.*
- (f) *$\overline{\lim}_{p \uparrow \infty} \|\varepsilon^\top B\| / p = 0$ almost surely for some matrix norm $\|\cdot\|$ in $\mathbb{R}^{n \times q}$.*

$$(g) \lim_{p \uparrow \infty} \|\epsilon^\top z\| / \sqrt{p} = 0$$

(h) $B = (\beta_1, \dots, \beta_q)$ with $\beta_i^\top \beta_j = 0$ for $i \neq j$, i.e. columns of B are orthogonal.

Moreover, we can denote $\lim_{p \rightarrow \infty} \frac{|\beta_i|^2}{p} = \eta_i$.

Remark 2.3.2. Without loss of generality, we can assume B to have orthogonal columns as we can write $B = \mathcal{B}\Lambda^2 U$ in eigen decomposition, and choose $B = \mathcal{B}\Lambda$ and new $X = \Lambda U X$. Here $\Lambda^2 = (\lambda_B(1)^2, \dots, \lambda_B(q)^2)$, the top q eigenvalues of B .

For the PCA estimate H of B

$$\hat{\Sigma} = \sum_{(s^2, h)} s^2 h h^\top = H H^\top + \hat{\gamma}^2 I_p, \quad (2.13)$$

the j -th column of the $p \times q$ matrix H is taken as $\eta = sh$ where s^2 is the j -th largest eigenvalue of $\hat{\Sigma}$ and h is the corresponding eigenvector. Ordering the eigenvalues of $\hat{\Sigma}$ as $s_{1,p}^2 \geq s_{2,p}^2 \geq \dots \geq s_{p,p}^2$, we have

$$\mathcal{H} = H S_p^{-1}; \quad H^\top H = S_p^2, \quad (2.14)$$

where S_p^2 is a $q \times q$ diagonal matrix with entries $(S_p^2)_{jj} = s_{j,p}^2$ and the columns of $\mathcal{H}$ are the associated sample eigenvectors h in Eq 2.13 with $\mathcal{H}^\top \mathcal{H} = I_q$. Denote the columns of $\mathcal{H}$ be $h_1, h_2, \dots, h_q$, we now introduce the relationship between h_i and $z = \frac{\zeta}{|\zeta|} \in \mathbb{R}^p$:

Lemma 2.3.3. For $i > q$, $\langle h_i, z \rangle \rightarrow 0$ as $p \uparrow \infty$.

Remark 2.3.4. This relationship implies the bulk eigenvectors are asymptotically orthogonal to any unit vector on $\mathbb{S}^{p-1}$ as p goes to infinity, which means that the bulk eigenvectors behave like random directions in high-dimensional space and are almost uncorrelated with any fixed vector.

2.3.2 Problem Formulation

One major issue in estimating Σ is that in practice, we don't know the number of latent factors q , so we need to propose an estimate of it, denote by K . Therefore, we decompose the sample covariance matrix S as

$$S = HH^\top + G, \quad (2.15)$$

where $H \in \mathbb{R}^{p \times K}$. For clarity, we use the subscript q to indicate quantities associated with the true underlying factors.

The following truncated consistency theorem is needed to ensure the validity of theorems and proofs in this paper.

Theorem 2.3.5. *Assume $A \in \mathbb{C}^{p \times p}$ is a Hermitian matrix with eigenvalues*

$$\lambda_1 \geq \lambda_2 \geq \cdots \geq \lambda_p,$$

with corresponding unit-norm eigenvectors $u_1, u_2, \dots, u_p \in \mathbb{C}^p$. Then for any integers $q < K \leq p$, let $u_1^{(K)}, \dots, u_K^{(K)}$ and $\lambda_1^{(K)}, \dots, \lambda_K^{(K)}$ be the first K eigenvectors and eigenvalues of A , and let $u_1^{(q)}, \dots, u_q^{(q)}$ and $\lambda_1^{(q)}, \dots, \lambda_q^{(q)}$ be the first q eigenvectors and eigenvalues computed directly from A .

Then:

1. *For all $i \leq q$, the eigenvalues are identical:*

$$\lambda_i^{(K)} = \lambda_i^{(q)}.$$

2. For all $i \leq q$, the eigenvectors differ only by a complex phase factor:

$$u_i^{(K)} = e^{i\theta_i} u_i^{(q)} \quad \text{for some } \theta_i \in \mathbb{R}.$$

That is, the first q eigenpairs obtained from the full top- K spectrum are identical (up to phase) to those computed directly from A . In the real symmetric case, the phase factors reduce to signs ± 1 .

Reference to the above theorem can be found in [18].

We also establish the relationship between eigenpairs of $\hat{\Sigma}$ and the dual covariance matrix $L = Y^\top Y/p$:

Lemma 2.3.6. $Lu = l^2 u$ for $l^2 = s^2 n/p > 0$ and $u = Y^\top v/(\sqrt{n}s)$ if and only if $\hat{\Sigma}v = s^2 v$ where $v = Yu/(\sqrt{p}l)$ and $s^2 = l^2 p/n > 0$.

Lastly, note that $HH^\dagger = \mathcal{H}\mathcal{H}^\top$, $\langle z, \mathcal{H}\mathcal{H}^\top z \rangle = \|\mathcal{H}\mathcal{H}^\top z\|^2$, and feasible set $Ax \leq a$ can be written as $\mathcal{A}x = \alpha$, we have

$$\mathcal{E}_p(H, A, a) = \mathcal{E}_p(\mathcal{H}, \mathcal{A}, \alpha),$$

i.e., we can construct the orthogonal estimator (namely $\mathcal{H}_\sharp$) and then analysis the optimization bias based on it.

2.3.3 $\mathcal{H}_\sharp$ Construction

Recall the definitions we have: for $K \times K$ diagonal matrix $S_p^2 = \text{diag}(s_{1,p}^2, s_{2,p}^2, \dots, s_{K,p}^2)$ with the elements being the top K ranked eigenvalues of the sample covariance matrix S , the sample covariance matrix can be decomposed as:

$$S = \mathcal{H} S_p^2 \mathcal{H}^\top + G, \tag{2.16}$$

with the columns of $\mathcal{H}$ are the associated sample eigenvectors with $\mathcal{H}^\top \mathcal{H} = I_K$.

The signal-to-noise ratio (SNR) matrix Ψ^2 quantifies the relative strength of the latent signal components with respect to the background noise. It is defined as

$$\Psi^2 = I_K - l_p^2 S_p^{-2}, \quad l_p^2 = \frac{\sum_{j>K} s_{j,p}^2}{n - K}, \quad (2.17)$$

where $S_p = \text{diag}(s_{1,p}, \dots, s_{K,p})$ contains the leading empirical eigenvalues, and l_p^2 represents the average energy of the bulk (noise) eigenvalues.

Intuitively, Ψ^2 serves as a diagonal matrix of normalized signal-to-noise ratios: when $s_{k,p}^2 \gg l_p^2$, the corresponding entry of Ψ^2 is close to 1, indicating a strong signal component; whereas when $s_{k,p}^2 \approx l_p^2$, the entry approaches 0, implying that the estimated component is dominated by noise.

This matrix plays a central role in the subsequent derivations, as it adjusts the contribution of each principal direction according to its estimated SNR and provides a natural scaling for shrinkage or bias correction of the sample eigenvectors.

Furthermore, l_p^2 satisfies the following:

Theorem 2.3.7.

$$\lim_{p \uparrow \infty} \frac{l_p^2}{p} = \begin{cases} \frac{\gamma^2}{n}, & K \geq q \\ \frac{(n-q)\gamma^2}{n(n-K)} + \frac{\sum_{K < j \leq q} \lambda_j^2}{n(n-K)}, & K < q \end{cases} \quad (2.18)$$

for constant γ in Assumption 2.3.1 and λ_i^2 being the i th largest eigenvalue of the dual covariance matrix $L = \frac{1}{p} Y^\top Y$.

For proof of Theorem 2.3.7, please see Appendix 2.6.3.

The $p \times K$ James-Stein estimator $\mathcal{H}_\#$ is then defined as

$$\mathcal{H}_\# = \mathcal{H} \Psi^2 + z_{\perp \mathcal{H}} \frac{z^\top \mathcal{H} (I_K - \Psi^2)}{|z - z_{\mathcal{H}}|}, \quad (2.19)$$

where

$$z_{\mathcal{H}} = \mathcal{H}\mathcal{H}^\top z \text{ and } z_{\perp\mathcal{H}} = \frac{z - z_{\mathcal{H}}}{|z - z_{\mathcal{H}}|}$$

, the orthogonal projection of z onto the orthogonal complement of $\text{col}(\mathcal{H})$.

Remark 2.3.8. *The James-Stein estimators defined in this paper will need to be orthonormalized, but it won't change $|\mathcal{E}_p(\mathcal{H}_\#)|$.*

2.4 Theoretical Results on Rank Misspecification Bias

2.4.1 James-Stein Behavior Under Overspecified rank $K > q$

We first start with $K > q$ case. From Equation 2.9 and noting that $HH^\dagger = \mathcal{H}\mathcal{H}^\top$, the $q \times 1$ optimization bias $\mathcal{E}_p(\mathcal{H}_\#)$ is defined as

$$\mathcal{E}_p(\mathcal{H}_\#) = \frac{\mathcal{B}^\top z - (\mathcal{H}^\top \mathcal{B})(\mathcal{H}^\top z)}{|z - z_{\perp\mathcal{H}}|}. \quad (2.20)$$

We now introduce the $p \times q$ James-Stein estimator $\mathcal{H}_{\#,q}$ in the same form of Eq 2.19:

$$\mathcal{H}_{\#,q} = \mathcal{H}_q \Psi_q^2 + z_{\perp\mathcal{H}_q} \frac{z^\top \mathcal{H}_q (I_q - \Psi_q^2)}{|z - z_{\mathcal{H}_q}|}, \quad (2.21)$$

where $\mathcal{H}_q$ is the $p \times q$ matrix of eigenvectors, and Ψ_q is the $q \times q$ signal-to-noise ratio $\Psi_q^2 = I_q - l_{p,q}^2 S_{p,q}^{-2}$. Here, $l_{p,q}^2$ and $S_{p,q}^2$ corresponds to the same definition as stated in Eq 2.17 with dimension being q instead of K . The connection between Ψ^2 and Ψ_q^2 is described as following:

Lemma 2.4.1. *Denote the first $q \times q$ block of Ψ^2 as $\tilde{\Psi}^2$, then as $p \uparrow \infty$,*

$$\frac{1}{p} \left\| \tilde{\Psi}^2 - \Psi_q^2 \right\| \rightarrow 0. \quad (2.22)$$

It was shown that $\mathcal{E}_p(\mathcal{H}_{\#,q}) \rightarrow 0$ almost surely [16] and we can link the $q \times 1$ vector $\mathcal{E}_p(\mathcal{H}_{\#})$ to $\mathcal{E}_p(\mathcal{H}_{\#,q})$ by the following theorem:

Theorem 2.4.2. *When $K \geq q$, $\mathcal{E}_p(\mathcal{H}_{\#}) \sim \mathcal{E}_p(\mathcal{H}_{\#,q}) \rightarrow 0$.*

The proof of Theorem 2.4.2 is provided in Appendix 2.6.4.

Recall Theorem 2.2.9 and similar to previous discussions, the discrepancy using $\mathcal{H}_{\#}$ approximately equals to

$$\hat{D}_p = -\frac{|\Lambda_p \mathcal{E}_p(\mathcal{H}_{\#})|^2}{\hat{\gamma}^2} + \left(2 - \frac{\langle u_{\mathcal{H}_{\#}}, \Gamma u_{\mathcal{H}_{\#}} \rangle}{\hat{\gamma}^2} \right) + O \left(|\mathcal{E}_p(\mathcal{H}_{\#})| + |\mathcal{E}_p(\mathcal{H}_{\#})|^2 + \frac{1}{p} \right), \quad (2.23)$$

Note that $\Lambda_p = O(\sqrt{p})$, $\hat{\gamma}^2$ and $\langle u_{\mathcal{H}_{\#}}, \Gamma u_{\mathcal{H}_{\#}} \rangle$ are finite, we can conclude $\hat{D}_p \rightarrow 0$ as $|\mathcal{E}_p(\mathcal{H}_{\#})| \rightarrow 0$ for true number of factors $q < K$.

2.4.2 James-Stein Behavior Under Underspecified rank $K < q$

Now we look at the $K < q$ case. Note here the definition of $\mathcal{H}_{\#}$ and $\mathcal{E}_p(\mathcal{H}_{\#})$ remains the same as in $K > q$ case, and we restate them below:

$$\mathcal{H}_{\#} = \mathcal{H} \Psi^2 + z_{\perp \mathcal{H}} \frac{z^{\top} \mathcal{H} (I_K - \Psi^2)}{|z - z_{\mathcal{H}}|}, \quad (2.24)$$

and

$$\mathcal{E}_p(\mathcal{H}_{\#}) = \frac{\mathcal{B}^{\top} z - (\mathcal{B}^{\top} \mathcal{H}_{\#})(\mathcal{H}_{\#}^{\top} z)}{\sqrt{1 - |\mathcal{H}_{\#}^{\top} z|^2}}, \quad (2.25)$$

the result is different from $K > q$ case.

Theorem 2.4.3. *Define the following notations for simplicity:*

1. $v = \mathcal{H}^\top z \in \mathbb{R}^K$
2. $\alpha = \mathcal{B}^\top \mathcal{H} \Psi^2 v \in \mathbb{R}^q$
3. $u = \mathcal{B}^\top (I_p - \mathcal{H} \mathcal{H}^\top) z \in \mathbb{R}^q$
4. $r = z^\top \mathcal{H} (I_K - \Psi^2) v \in \mathbb{R}$,

we have

$$\|\mathcal{E}_p(\mathcal{H}_\#)\|^2 = \frac{\|\mathcal{B}^\top z - \alpha\|^2}{1 - \|v\|^2} + \frac{r^2 \|\mathcal{E}_p(\mathcal{H})\|^2 - 2r \langle \mathcal{B}^\top z - \alpha, u \rangle}{(1 - \|v\|^2)^2}. \quad (2.26)$$

Remark 2.4.4. *Note that $\mathcal{B}^\top z - \alpha = \mathcal{B}^\top (I_p - \mathcal{H} \Psi^2 \mathcal{H}^\top) z$, so Eq 2.26 is decomposing the total squared optimization bias of the shrinkage estimator $\mathcal{H}_\#$ into three parts:*

- *the squared discrepancy between the projection of z and the shrinkage-retained component in $\text{Col}(\mathcal{B})$*
- *the squared optimization bias of PCA estimator $\mathcal{H}$*
- *the alignment between the residual component of z after shrinkage and the pure projection residual of z in $\text{Col}(\mathcal{B})$*

Corollary 2.4.5. *When $K < q$, $\mathcal{E}_p(\mathcal{H}_\#) \not\rightarrow 0$ as $p \uparrow \infty$.*

The proof of Theorem 2.4.3 and Corollary 2.4.5 are provided together in Appendix 2.6.5. Recall Theorem 2.2.9 and similar to previous discussions, the discrepancy using $\mathcal{H}_\#$ approximately equals to

$$\hat{D}_p = -\frac{|\Lambda_p \mathcal{E}_p(\mathcal{H}_\#)|^2}{\hat{\gamma}^2} + \left(2 - \frac{\langle u_{\mathcal{H}_\#}, \Gamma u_{\mathcal{H}_\#} \rangle}{\hat{\gamma}^2}\right) + O\left(|\mathcal{E}_p(\mathcal{H}_\#)| + |\mathcal{E}_p(\mathcal{H}_\#)|^2 + \frac{1}{p}\right), \quad (2.27)$$

Note that $\Lambda_p = O(\sqrt{p})$, $\hat{\gamma}^2$ and $\langle u_{\mathcal{H}_\#}, \Gamma u_{\mathcal{H}_\#} \rangle$ are finite, we can conclude $\hat{D}_p \rightarrow -\infty$ as $|\mathcal{E}_p(\mathcal{H}_\#)| \not\rightarrow 0$ for true number of factors $q > K$.

2.4.3 A Closer Look: Choose number of factors $K = 1$

We further study the trivial case when $K = 1$, and obtain the exact expression of the first element of the optimization bias. Follows the same decomposition, we write the sample covariance matrix as

$$S = s_{1,p}^2 h h^\top + G, \quad (2.28)$$

where $(s_{1,p}^2, h)$ is the leading eigenpair of $\hat{\Sigma}$. The signal-to-noise ratio $\psi^2 = \frac{s_p^2 - l_p^2}{s_p^2}$ is now a number with $l_p^2 = \frac{\sum_{i>1} s_{i,p}^2}{n-K}$ being the average non-zero bulk eigenvalues. Define $N = \frac{\langle h, z \rangle - \psi^2 \langle h, z \rangle}{\sqrt{1 - \langle h, z \rangle^2}}$ and $D = \sqrt{\psi^4 + N^2}$ as the normalizing constants. The updated vector $h_\#$ is a combination of h and $z_{\perp h}$, the unit vector in the direction of the component of z that is orthogonal to h :

$$h_\# = \frac{1}{D}(\psi^2 h + N z_{\perp h}), \quad (2.29)$$

with $z_h = \langle h, z \rangle h$ and $z_{\perp h} = \frac{z - z_h}{|z - z_h|}$. The optimization bias of $h_\#$ is:

$$\mathcal{E}_p(h_\#) = \frac{\mathcal{B}^\top z - (\mathcal{B}^\top h_\#)(h_\#^\top z)}{\sqrt{1 - \langle h_\#, z \rangle^2}}. \quad (2.30)$$

Denote the eigenvalues of the dual covariance matrix be $\lambda_1^2 \geq \lambda_2^2 \geq \dots \geq \lambda_n^2$, and ω be the eigenvector corresponds to λ_1^2 , we have the main theorem to characterize the asymptotic behavior for this case:

Theorem 2.4.6.

(a) The $q \times 1$ vector $\mathcal{E}_p(h_\#) \neq \mathcal{E}_p(\mathcal{H}_\#)$ with the first element $\mathcal{E}_{p,1}(h_\#) = \frac{\psi^2 \langle b, z \rangle - \langle h, b \rangle \langle h, z \rangle}{\sqrt{\psi^4 + (1 - 2\psi^2) \langle h, z \rangle^2}}$.

Furthermore, $|\mathcal{E}_p(h_\#)|$ is bounded away from 0 eventually.

(b) $\psi^2 \sim 1 - \left(\frac{\gamma^2}{\lambda_1^2} + \frac{\sum_{1 \leq j \leq q} (\lambda_j^2 - \gamma^2)}{\lambda_1^2(n-1)} \right)$ and is decreasing in q as $p \rightarrow \infty$.

(c) $\langle h, b \rangle^2 \sim \psi^2 - \frac{1}{\lambda_1^2} \sum_{i=2}^q \eta_i (X_i \omega)^2$ and is decreasing in q as $p \rightarrow \infty$ for η_i defined in Assumption 2.3.1 and X_i being the i -th row of X .

(d) $\langle h, z \rangle^2 \sim \frac{1}{\lambda_1^2} \sum_{i=1}^q \eta_i (X\omega)_i^2 \langle b_i, z \rangle^2$ and is decreasing in q as $p \rightarrow \infty$, where $(X\omega)_i$ is the i -th element of the $q \times 1$ vector $X\omega$.

Remark 2.4.7. Theorem 2.4.6 provides an asymptotic characterization of the estimated leading direction when the model rank is misspecified ($K = 1 < q$). It shows that even in this simplest case, the sample-based estimator $h_\#$ exhibits a non-vanishing bias relative to the true population direction b . The term ψ^2 captures the shrinkage strength toward the isotropic component, which decreases as the true factor dimension q increases, indicating stronger regularization in more complex latent structures. Meanwhile, the inner products $\langle h, b \rangle$ and $\langle h, z \rangle$ respectively quantify the alignment with the true signal and the bias toward the shrinkage target. Together, these results reveal how rank misspecification induces systematic deviation of the estimated eigenvector from its population counterpart, even in high-dimensional asymptotics.

Recall Theorem 2.2.9 and similar to previous discussions, the discrepancy using $h_\#$ approximately equals to

$$\hat{D}_p = -\frac{|\Lambda_p \mathcal{E}_p(h_\#)|^2}{\hat{\gamma}^2} + \left(2 - \frac{\langle u_{h_\#}, \Gamma u_{h_\#} \rangle}{\hat{\gamma}^2} \right) + O\left(|\mathcal{E}_p(h_\#)| + |\mathcal{E}_p(h_\#)|^2 + \frac{1}{p} \right), \quad (2.31)$$

and $\hat{D}_p \rightarrow -\infty$ as $|\mathcal{E}_p(h_\#)| \not\rightarrow 0$ for true number of factors $q > 1$.

2.4.4 Equivalence with a James-Stein type estimator

Let $M = AA^\dagger H$ for any $A \in \mathbb{R}^{p \times r}$, $C = I_K - l^2 N^{-1}$ and $N = (H - M)^\top (H - M)$.

We have the James-Stein estimator for eigenvectors from [14]:

$$H_{JS} = HC + M(I_q - C). \quad (2.32)$$

Recall the estimator we propose in Sec 2.3.3:

$$H_\sharp = H\Psi^2 + \frac{(I_p - HH^\dagger)z}{1 - z^\top HH^\dagger z} z^\top H(I_q - \Psi^2), \quad (2.33)$$

we have the following lemma:

Lemma 2.4.8. *Take $A = z$ in Eq 2.32, then $H_{JS} = H_\sharp$.*

Proof. Let $A = z$. Note that $zz^\dagger = zz^\top$, we have

$$N = H^\top (I_p - zz^\top) H \quad (2.34)$$

and

$$N^{-1} = (H^\top H)^{-1} + \frac{H^\dagger zz^\top (H^\dagger)^\top}{1 - z^\top HH^\dagger z}, \quad (2.35)$$

By noting $H^\top H = S^2$, $H^\dagger = S^{-2}H^\top$ and recall $\Psi^2 = I_q - l^2 S^{-2}$,

$$\begin{aligned}
H_{JS} &= H - H(H^\top H)^{-1}l^2 - \frac{HH^\dagger zz^\top (H^\dagger)^\top}{1 - z^\top HH^\dagger z}l^2 + zz^\top H(H^\top H)^{-1}l^2 + \frac{zz^\top HH^\dagger zz^\top (H^\dagger)^\top}{1 - z^\top HH^\dagger z}l^2 \\
&= H\Psi^2 + zz^\top HS^{-2}l^2 + \frac{zz^\top HH^\dagger zz^\top (H^\dagger)^\top - HH^\dagger zz^\top (H^\dagger)^\top}{1 - z^\top HH^\dagger z}l^2 \\
&= H\Psi^2 + \left(zz^\top H - \frac{(I_p - zz^\top)HH^\dagger zz^\top H}{1 - z^\top HH^\dagger z} \right) (I_q - \Psi^2) \\
&= H\Psi^2 + \frac{(I_p - HH^\dagger)zz^\top H + zz^\top HH^\dagger zz^\top H - (z^\top HH^\dagger z)zz^\top H}{1 - z^\top HH^\dagger z} (I_q - \Psi^2) \\
&= H\Psi^2 + \frac{(I_p - HH^\dagger)z}{1 - z^\top HH^\dagger z} z^\top H (I_q - \Psi^2) \\
&= H_\sharp.
\end{aligned} \tag{2.36}$$

The second from the last equation comes from the fact that

$$z(z^\top HH^\dagger z) = (z^\top HH^\dagger z)z. \tag{2.37}$$

□

Remark 2.4.9. *The Lemma above indicates $H_\sharp$ is a special case of H_{JS} . Although it's a specialization, it captures the subspace-adjusted shrinkage along the projection residual $(I_p - HH^\dagger)z$, which corresponds to the residual direction along which optimization bias arises under rank misspecification. Focusing on this case lets us obtain closed-form expressions and sharp asymptotics, while still reflecting the behavior of the regular James-Stein estimator.*

2.5 Portfolio Optimization with Misspecified Factor Models: A Simulation Study

2.5.1 Simulation Framework and Methodology

We illustrate the results on a numerical example to provide a verification of Theorems 2.4.3, 2.4.3 and 2.4.6. We consider i.i.d observations of $y \in \mathbb{R}^{p_{max}}$ with

$$y = \alpha + Bx + \epsilon, \quad (2.38)$$

where $\alpha \in \mathbb{R}^{p_{max}}$ represents the vector of asset-specific alphas, $B \in \mathbb{R}^{p_{max} \times q}$ represents the factor loadings, $x \in \mathbb{R}^q$ and $\epsilon \in \mathbb{R}^{p_{max}}$. x and ϵ are uncorrelated with $Var(x) = I_q$ and $Var(\epsilon) = \Gamma$. Thus, the variance of y is given by

$$Var(y) = \Sigma = BB^\top + \Gamma. \quad (2.39)$$

In our simulation, we first fix $n = 120$, $q = 7$ and $p_{max} = 128,000$, and generate $Y \in \mathbb{R}^{p_{max} \times n}$ with y as its columns.

The covariance matrix calibration follows the same as specified in [16]. The random

vector of factor returns $f \in \mathbb{R}^7$ and the $p_{max} \times 7$ exposure matrix Ξ satisfy

$$\text{Var}(f) = \begin{pmatrix} 250 & 0 & 0 & 55 & 44 & 68 & -22 \\ 0 & 64 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 16 & 0 & 0 & 0 & 0 \\ 55 & 0 & 0 & 481 & 192 & -108 & 0 \\ 44 & 0 & 0 & 192 & 260 & -8 & 22 \\ 68 & 0 & 0 & -108 & -8 & 160 & -44 \\ -22 & 0 & 0 & 0 & 22 & -44 & 121 \end{pmatrix}, \quad Bx = \Xi f. \quad (2.40)$$

The factor returns f are Gaussian with mean 0 and variance specified in Eq 2.40. The columns of Ξ are exposures to fundamental risk factors.

- The first column of the exposure matrix Ξ , representing market risk exposures, is generated with entries independently drawn from a normal distribution with mean 1.0 and standard deviation 0.25.
- The second and third columns of Ξ , corresponding to style risk factors, are filled with i.i.d. entries from normal distributions with mean zero and standard deviations of 0.5 and 1.0, respectively.
- The remaining four columns of Ξ , representing industry exposures, are initialized to zero. For each row i , we independently sample two distinct industry indices I_1 and I_2 uniformly without replacement from the set $\{1, 2, 3, 4\}$. We then draw two independent random variables $U_1, U_2 \sim \text{Unif}(0, 1)$, and assign

$$\Xi_{iI_1} = U_1, \quad \Xi_{iI_2} = \Xi_{iI_1} + U_2.$$

The asset-specific return vector $\epsilon \in \mathbb{R}^{p_{max}}$ is modeled as a mean-zero Gaussian random

vector with diagonal covariance matrix $\text{var}(\epsilon) = \Gamma$. Each diagonal entry Γ_{ii} is defined as γ_i^2 , where γ_i denotes the idiosyncratic volatility of asset i . These volatilities are sampled independently according to

$$\gamma_i = 25 + 75 \times Z, \quad \text{where } Z \sim \text{Beta}(4, 16).$$

The resulting values γ_i are interpreted as annualized volatilities in percent.

For the simulated matrix $Y_{p_{\max} \times n}$, we select subsets $Y_{p \times n}$ with $p = 500, 2000, 8000, 32000, 128000$ to study the asymptotics of PCA estimator $\hat{\Sigma}$ and the James-Stein estimator $\hat{\Sigma}_{\sharp}$. They are all based on the sample covariance matrix $S = YY^{\top}/n$. For $q = 7$ and $K = 1, q - 1, q, q + 1$, and $n - 1$, we can compute the $K \times K$ diagonal matrix S_p^2 with the K largest eigenvalues of S and Ψ^2 as stated in Eq 2.17. The final estimators of the population covariance matrix has the form:

$$\hat{\Sigma}_{\natural} = \mathcal{H}_{\natural}(S_p \Psi)^2 \mathcal{H}_{\natural}^{\top} + \hat{\gamma}^2 I, \mathcal{H}_{\natural} \in \{\mathcal{H}, \mathcal{H}_{\sharp}\}, \quad (2.41)$$

where $\hat{\gamma}^2 = nl_p^2/p$ for l_p^2 in Eq 2.17, and $\mathcal{H}_{\natural}$ represents the $p \times K$ matrix of eigenvectors.

To evaluate the performance of the covariance estimators defined in Eq 2.41, we consider two key metrics: the subspace overlap $|\mathcal{E}_p(\mathcal{H}_{\natural})|$, and the associated discrepancy measure

$$\hat{D}_p = 2 - \hat{\mu}_p^2 \mathcal{V}_p^2,$$

where the scalar quantity $\hat{\mu}_p^2$ is given by

$$\hat{\mu}_p^2 = \langle \mathbf{1}_p, \hat{\Sigma}_{\natural}^{-1} \mathbf{1}_p \rangle,$$

and the variance component $\mathcal{V}_p^2$ is defined as

$$\mathcal{V}_p^2 = \langle \hat{w}_{\natural}, \Sigma \hat{w}_{\natural} \rangle.$$

Here, $\hat{w}_{\natural}$ denotes the portfolio weights obtained by solving the minimum variance portfolio problem under the estimated covariance matrix $\hat{\Sigma}_{\natural}$.

Furthermore, we also tested $|\mathcal{E}_p(\mathcal{H}_{\natural})|$ and $\hat{D}_p$ for various value of qs .

2.5.2 Empirical Results and Analysis

Table 2.1 summarizes the optimization bias $\mathcal{E}_p(\mathcal{H}_{\natural})$ for all three cases: $K < q$, $K = q$ and $K > q$. The first two columns verify our result in Theorem 2.4.6, 2.4.3 and 2.4.2. It can be seen that

- The James-Stein eigenvector estimator $\mathcal{H}_{\natural}$ outperforms the PCA estimator $\mathcal{H}$ on reducing the optimization bias $\mathcal{E}_p(\cdot)$,
- When $K < q$, it is shown that the optimization bias $\mathcal{E}_p(\mathcal{H})$ and $\mathcal{E}_p(\mathcal{H}_{\natural})$ both don't converge to 0 as $p \uparrow \infty$. Correspondingly, both $\sqrt{p}\mathcal{E}_p(\mathcal{H})$ and $\mathcal{E}_p(\mathcal{H}_{\natural})$ diverges to ∞ . Combining with Theorem 2.2.9, $\hat{D}_p$ for $\hat{\Sigma}$ and $\hat{\Sigma}_{\natural}$ diverges to $-\infty$, as verified in Figure 2.1.
- When $K \geq q$, we have $\mathcal{E}_p(\mathcal{H}_{\natural})$ converges to 0 as $p \uparrow \infty$. Furthermore, $\sqrt{p}\mathcal{E}_p(\mathcal{H}_{\natural})$ is bounded which results in $\hat{D}_p \rightarrow 0$ as $p \uparrow \infty$ as provided in Theorem 2.2.9 and verified in Figure 2.1.

2.5.3 Interpretation and Practical Recommendations

Our theoretical analysis and numerical simulations demonstrate that the James-Stein estimator maintains stable eigenvector estimation even when employing an intentionally

p	K	$\mathbb{E}(\mathcal{E}_p(\mathcal{H}))$	$\mathbb{E}(\mathcal{E}_p(\mathcal{H}_\dagger))$	$\sqrt{p} \mathbb{E}(\mathcal{E}_p(\mathcal{H}))$	$\sqrt{p} \mathbb{E}(\mathcal{E}_p(\mathcal{H}_\dagger))$
500	1	0.5130	0.5043	11.4706	2.7986
	6	0.2164	0.1270	4.8397	2.7954
	7	0.2149	0.1251	4.8055	2.7983
	8	0.2138	0.1233	4.7812	2.7983
	127	0.1654	0.0823	3.6993	1.8395
2000	1	0.5259	0.5177	23.5229	23.1528
	6	0.1993	0.0769	8.9137	3.4411
	7	0.1969	0.0662	8.8039	2.9609
	8	0.1964	0.0645	8.7815	2.8835
	127	0.1818	0.0505	8.1311	2.2590
8000	1	0.5559	0.5482	49.7244	49.0342
	6	0.2040	0.0811	18.2434	7.2605
	7	0.1965	0.0298	17.5772	2.6688
	8	0.1965	0.0297	17.5795	2.6546
	127	0.1921	0.0301	17.1858	2.6909
32000	1	0.5621	0.5546	100.5569	99.2142
	6	0.2013	0.0720	36.0226	12.8722
	7	0.1942	0.0128	34.7336	2.2831
	8	0.1942	0.0127	34.7448	2.2775
	127	0.1929	0.0162	34.5070	2.8900
128000	1	0.5576	0.5499	199.4869	196.7399
	6	0.1996	0.0706	71.4061	25.2424
	7	0.1923	0.0065	68.8033	2.3344
	8	0.1924	0.0065	68.8480	2.3180
	127	0.1918	0.0085	68.6124	3.0459

Table 2.1: Optimization bias metrics computed with $\mathcal{H}_\dagger$ for $q = 7$. The metric under 'correct' estimate ($K = q$) is highlighted in bold.

inflated factor number (e.g., $q = n - 1$). This finding suggests a practical hybrid approach combining conventional factor selection methods with our shrinkage estimation: First, practitioners may determine a preliminary factor dimension q_0 by applying techniques such as information criterion [38] and hypothesis testing framework [76]. Then, they can deliberately expand the factor set to $q_0 + k$ ($k \geq 1$) before applying our shrinkage estimation. This integrated protocol leverages the statistical efficiency of traditional methods

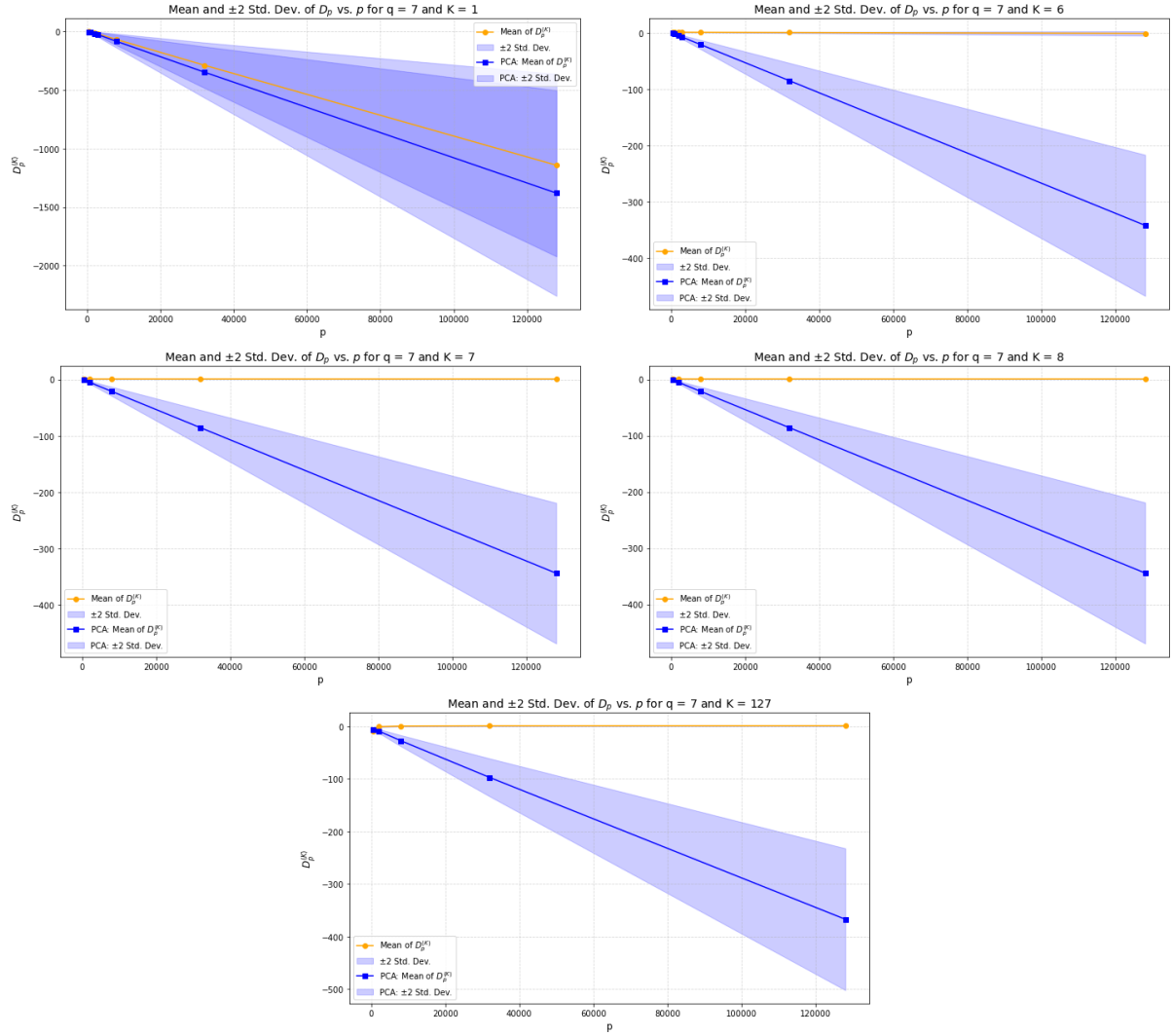


Figure 2.1: Discrepancy $\hat{D}_p$ v.s. p with the sample means and two standard deviation confidence intervals for $q = 7$.

while mitigating under-specification risks through our stabilized estimator, achieving enhanced robustness in real-world applications.

The proposed hybrid strategy offers three distinct advantages: First, it preserves the theoretical rigor of conventional factor selection procedures. Second, the shrinkage estimation effectively compensates for potential specification errors inherent in discrete factor-number selection. Third, as demonstrated in our numerical experiments,

the method exhibits superior robustness across varying signal-to-noise ratios and sample sizes. The expansion parameter k can be further optimized via cross-validation, balancing estimation accuracy with computational efficiency. This framework not only extends the applicability of existing methods but also provides a principled solution to model selection challenges in high-dimensional factor analysis, particularly when the true factor dimension is ambiguous.

2.6 Proofs for Chapter 2

2.6.1 Proofs for Section 2.2

Proof of Theorem 2.2.3. For the multipliers $\ell = (\ell_1, \dots, \ell_r)$, write $\alpha = \ell/|\ell|$, we then have $c_0 = -\ell^\top c$ and $c_1 = \|\ell\|$. So

$$\max_x Q(x) = c_0 + \frac{c_1^2}{2} \mu_p^2 = -\ell^\top c + \frac{c_1^2}{2} (A\ell)^\top \Sigma^{-1} (A\ell).$$

Let $P_B = BB^\dagger$, the orthogonal projector onto $\text{col}(B)$, and decompose ζ into $\zeta = \zeta_B + \zeta_\perp = P_B \zeta + (I - P_B)\zeta$. Under Assumption 2.2.1 and let $c_+ = \overline{\lim}_{p \uparrow \infty} \sup_{|v|=1} \langle v, \Gamma v \rangle$, we have on $\text{col}(B)^\perp$,

$$\Sigma^{-1} \succeq \Gamma^{-1} \succeq \frac{1}{c_+} I,$$

so

$$\zeta^\top \Sigma^{-1} \zeta \geq \zeta_\perp^\top \Sigma^{-1} \zeta_\perp \geq \frac{1}{c_+} \|\zeta_\perp\|^2 = \frac{1}{c_+} \|(I - P_B)A\ell\|^2 \geq \frac{\xi}{c_+} p \|\ell\|^2.$$

Hence,

$$\max_x Q(x) = -\ell^\top c + \frac{c_1^2}{2} \zeta^\top \Sigma^{-1} \zeta \geq -\ell^\top c + \frac{c_1^2}{2} \frac{\xi}{c_+} p \|\ell\|^2 \rightarrow \infty.$$

□

□

Proof of Theorem 2.2.6. Denote

$$R_1(\alpha) = \alpha^\top (A^\top \hat{\Sigma}^{-1} \Sigma \hat{\Sigma}^{-1} A) \alpha = \langle \hat{\Sigma}^{-1} v, \Sigma \hat{\Sigma}^{-1} v \rangle$$

and

$$R_2(\alpha) = \alpha^\top (A^\top \hat{\Sigma}^{-1} A) \alpha = \langle v, \hat{\Sigma}^{-1} v \rangle.$$

From Woodbury formula,

$$\hat{\Sigma}^{-1} = \hat{\gamma}^{-2} I - \hat{\gamma}^{-4} H (I + \hat{\gamma}^{-2} H^\top H)^{-1} H^\top.$$

Apply SVD on H : $H = USV^\top$ where $U, V \in \mathbb{R}^{p \times k}$ and $S = \text{diag}(s_1, \dots, s_k) \succ 0$, then $HH^\top = US^2U^\top$ and the projection operator

$$P_H = HH^\dagger = UU^\top.$$

Write $v = P_H v + (I - P_H)v = v_H + v_\perp$, we have

$$\hat{\Sigma}^{-1} v = U \text{diag} \left(\frac{1}{s_i^2 + \hat{\gamma}^2} \right) U^\top v + \hat{\gamma}^{-2} v_\perp. \quad (2.42)$$

Therefore,

$$R_2(\alpha) = v_H^\top U \text{diag} \left(\frac{1}{s_i^2 + \hat{\gamma}^2} \right) U^\top v_H + \hat{\gamma}^{-2} \|v_\perp\|^2 \leq \hat{\gamma}^{-2} (\|v_H\|^2 + \|v_\perp\|^2) = \hat{\gamma}^{-2} \|v\| < \infty.$$

Next, note that

$$R_1(\alpha) = \langle \hat{\Sigma}^{-1} v, \Sigma \hat{\Sigma}^{-1} v \rangle \geq \left\| B^\top \hat{\Sigma}^{-1} v \right\|^2.$$

From triangle inequality and Eq 2.42,

$$\left\| B^\top \hat{\Sigma}^{-1} v \right\| \geq \hat{\gamma}^{-2} \left\| B^\top v_\perp \right\| - \left\| B^\top U \text{diag} \left(\frac{1}{s_i^2 + \hat{\gamma}^2} \right) U^\top v \right\|,$$

and

$$\left\| B^\top U \text{diag} \left(\frac{1}{s_i^2 + \hat{\gamma}^2} \right) U^\top v \right\| \leq \frac{1}{\hat{\gamma}^2} \left\| B^\top v_H \right\|,$$

so

$$\left\| B^\top \hat{\Sigma}^{-1} v \right\| \geq \frac{1}{\hat{\gamma}^2} (\left\| B^\top v_\perp \right\| - \left\| B^\top v_H \right\|) \quad (2.43)$$

By the transversality condition (T), there exists $\delta > 0$ such that

$$\frac{1}{p} B^\top P_{H^\perp} B \succeq \delta I \implies \left\| B^\top v_\perp \right\|^2 = v_\perp^\top B B^\top v_\perp \geq v_\perp^\top B P_{H^\perp} B^\top v_\perp \geq \delta p \left\| v_\perp \right\|^2. \quad (\text{T}')$$

Meanwhile, since $B^\top B/p \rightarrow \Sigma_B \succ 0$, there is $C < \infty$ with

$$\left\| B^\top v_H \right\|^2 \leq \lambda_{\max}(B^\top B) \left\| v_H \right\|^2 \leq C p \left\| v_H \right\|^2. \quad (\text{BH})$$

Combining (T') and (BH) with Eq 2.43 and using $\|a - b\|^2 \geq (1 - \varepsilon)\|a\|^2 - \varepsilon^{-1}\|b\|^2$ for any $\varepsilon \in (0, 1)$ with $a = B^\top v_\perp$, $b = B^\top v_H$, we obtain

$$\left\| B^\top \hat{\Sigma}^{-1} v \right\|^2 \geq \frac{1}{\hat{\gamma}^4} \left((1 - \varepsilon) \left\| B^\top v_\perp \right\|^2 - \varepsilon^{-1} \left\| B^\top v_H \right\|^2 \right) \geq \frac{p}{\hat{\gamma}^4} \left((1 - \varepsilon) \delta \left\| v_\perp \right\|^2 - \varepsilon^{-1} C \left\| v_H \right\|^2 \right).$$

Choose a fixed $\varepsilon \in (0, 1)$ small enough so that the bracket is positive (this is ensured by the transversality (T) and Assumption 2.2.5). Hence there exists $c > 0$ independent of p such that

$$\left\| B^\top \hat{\Sigma}^{-1} v \right\|^2 \geq c \hat{\gamma}^{-4} p.$$

Therefore,

$$R_1(\alpha) = \langle \hat{\Sigma}^{-1}v, \Sigma \hat{\Sigma}^{-1}v \rangle \geq \|B^\top \hat{\Sigma}^{-1}v\|^2 \geq c \hat{\gamma}^{-4} p \xrightarrow[p \rightarrow \infty]{} \infty.$$

and $\hat{D}_p \rightarrow -\infty$. Lastly, since $c_1^2 \hat{\mu}_p^2$ is bounded, we have $Q(\hat{x}) \rightarrow \infty$. $\square$ $\square$

2.6.2 Proofs for Section 2.3.1

Proof of Lemma 2.3.3. Write $\hat{\Sigma} = S + E$, where

$$S = \frac{1}{n} B X X^\top B^\top, \quad E = \frac{1}{n} (\varepsilon X^\top B^\top + B X \varepsilon^\top + \varepsilon \varepsilon^\top) \quad (2.44)$$

We have $\frac{1}{n} \varepsilon \varepsilon^\top \rightarrow I$ by law of large numbers. The cross term in E satisfies

$$\left\| \frac{1}{n} \varepsilon X^\top B^\top \right\| \leq \frac{1}{n} \|\varepsilon\| \|X\| \|B\| \rightarrow 0 \quad (2.45)$$

based on Assumption 2.2.1, and similarly $\left\| \frac{1}{n} B X \varepsilon^\top \right\| \rightarrow 0$. Therefore $E \xrightarrow{P} \gamma^2 I_p$.

For the eigenpair $(s_{i,p}^2, h_i)$ of $\hat{\Sigma}$, we have $\hat{\Sigma} h_i = s_{i,p}^2 h_i$. Multiply $z^\top$ on the left hand side, we have

$$z^\top \hat{\Sigma} h_i = s_{i,p}^2 \langle z, h_i \rangle = z^\top (S + E) h_i. \quad (2.46)$$

Note $\text{rank}(S) = q$, and $h_i \perp S$ as $p \uparrow \infty$ for $i > q$, we have $z^\top S h_i \rightarrow 0$. Combined with asymptotics of E , we have

$$\gamma^2 z^\top h_i \sim s_{i,p}^2 \langle z, h_i \rangle. \quad (2.47)$$

Therefore,

$$\begin{aligned} \langle z, h_i \rangle &\sim \frac{z^\top E h_i}{s_{i,p}^2} \\ &= \frac{1}{n s_{i,p}^2} \sum_{j=1}^n (z^\top \varepsilon_{:,j}) (\varepsilon_{:,j}^\top h_i). \end{aligned} \quad (2.48)$$

Note $\mathbb{E}(z^\top \varepsilon_{:,k}) = \mathbb{E}(\varepsilon_{:,k}^\top) = 0$, $\text{Var}(z^\top \varepsilon_{:,k}) = \gamma^2 \|z\|^2 = \gamma^2$ and $\text{Var}(\varepsilon_{:,k}^\top h_i) = \gamma^2 \|h_i\|^2 = \gamma^2$, we have $\langle z, h_i \rangle \rightarrow 0$. $\square$

2.6.3 Proofs for Section 2.3.3

For γ^2 in Assumption 2.3.1, define the $n \times K$ matrix M as:

$$L = MM^\top + \gamma^2 I_n. \quad (2.49)$$

Let $\lambda_{j,n}^2(M)$ denote the j -th largest eigenvalue of $MM^\top$ associated with the j -th column of eigenvectors of $MM^\top$. By Assumption 2.3.1, we have $\lambda_{j,n}^2(M) > 0$ for $j \leq q$ and $\lambda_{j,n}^2(M) = 0$ otherwise.

Lemma 2.6.1. *Under Assumption 2.2.1 and 2.3.1, almost surely, $\lim_{p \uparrow \infty} Y^\top Y/p = L$ and*

$$\lim_{p \uparrow \infty} \frac{s_{j,p}^2}{p} = \begin{cases} (\lambda_{j,n}^2(M) + \gamma^2)/n & 1 \leq j < n \\ \gamma^2/n & j = n \\ 0 & \text{otherwise} \end{cases} \quad (2.50)$$

The proof can be found in [16].

Proof of Theorem 2.3.7. Case 1: $q \leq K \leq n$: we have $s_{j,p}^2/p \rightarrow \gamma^2/n$ for $q+1 \leq j \leq n$ via Lemma 2.6.1. So $l_p^2/p \rightarrow \gamma^2/n$ almost surely.

Case 2: $K < q$: we have

$$\begin{aligned}
\lim_{p \uparrow \infty} \frac{l_p^2}{p} &= \lim_{p \uparrow \infty} \frac{1}{n - K} \sum_{j > K} \frac{s_{j,p}^2}{p} \\
&= \frac{1}{n - K} \sum_{j > K} \frac{\lambda_{j,n}^2(M) + \gamma^2}{n} \\
&= \frac{\gamma^2}{n} + \frac{\sum_{j > K} \lambda_{j,n}^2(M)}{n(n - K)} \\
&= \frac{\gamma^2}{n} + \frac{\sum_{j > K} (\lambda_j^2 - \gamma^2)}{n(n - K)} \\
&\sim \frac{\gamma^2}{n} + \frac{\sum_{K < j \leq q} (\lambda_j^2 - \gamma^2)}{n(n - K)}
\end{aligned} \tag{2.51}$$

The equations above used the assumption that Y is of full rank. $\square$

2.6.4 Proofs for Section 2.4.1

Proof of Lemma 2.4.1. Since

$$\Psi^2 = \text{diag} \left(1 - \frac{l_p^2}{s_{1,p}^2}, 1 - \frac{l_p^2}{s_{2,p}^2}, \dots, 1 - \frac{l_p^2}{s_{K,p}^2} \right), \tag{2.52}$$

and

$$\Psi_q^2 = \text{diag} \left(1 - \frac{l_{p,q}^2}{s_{1,p}^2}, 1 - \frac{l_{p,q}^2}{s_{2,p}^2}, \dots, 1 - \frac{l_{p,q}^2}{s_{q,p}^2} \right) \tag{2.53}$$

by Theorem 2.3.5, we have

$$\tilde{\Psi}^2 = \Psi_q^2 + \text{diag} \left(\frac{l_{p,q}^2 - l_p^2}{s_{1,p}^2}, \frac{l_{p,q}^2 - l_p^2}{s_{2,p}^2}, \dots, \frac{l_{p,q}^2 - l_p^2}{s_{q,p}^2} \right). \tag{2.54}$$

Combine with $K > q$ case of Theorem 2.3.7, we have $\frac{1}{p} \left\| \tilde{\Psi}^2 - |s_i^2| \right\| \rightarrow 0$. $\square$

Proof of Theorem 2.4.2. For any $1 < q < K$, based on Lemma 2.3.3, we have

$$\left\| \mathcal{H}^\top z - \begin{pmatrix} \mathcal{H}_q^\top z \\ 0 \\ \vdots \\ 0 \end{pmatrix} \right\| \rightarrow 0. \quad (2.55)$$

Also, from Lemma 2.4.1, we have

$$\left\| \Psi^2 - \begin{pmatrix} \Psi_q^2 & 0 \\ 0 & 0 \end{pmatrix} \right\| \rightarrow 0. \quad (2.56)$$

Therefore, we have

$$\|\mathcal{H}\Psi^2 - (\mathcal{H}_q\Psi_q^2, 0, \dots, 0)\| \rightarrow 0, \quad (2.57)$$

and

$$\left\| \frac{z_{\perp \mathcal{H}}^\top \mathcal{H} (I - \Psi^2)}{|z - z_{\mathcal{H}}|} - \left(\frac{z_{\perp \mathcal{H}}^\top \mathcal{H}_q (I_q - \Psi_q^2)}{|z - z_{\mathcal{H}}|}, 0, \dots, 0 \right) \right\| \rightarrow 0. \quad (2.58)$$

Write $\mathcal{H}_\#$ in terms of $\mathcal{H}_q$ and Ψ_q^2 , we have

$$\left\| \mathcal{H}_\# - \left(\mathcal{H}_q\Psi_q^2 + \frac{z_{\perp \mathcal{H}}^\top \mathcal{H}_q (I_q - \Psi_q^2)}{|z - z_{\mathcal{H}}|}, 0, \dots, 0 \right)_{p \times K} \right\| \rightarrow 0. \quad (2.59)$$

Lemma 2.6.2. When $q < K$, $\|z_{\mathcal{H}} - z_{\mathcal{H}_q}\| \rightarrow 0$ and $\|z_{\perp \mathcal{H}} - z_{\perp \mathcal{H}_q}\| \rightarrow 0$.

Proof. We first note $\mathcal{H}^\dagger = (\mathcal{H}^\top \mathcal{H})^{-1} \mathcal{H}^\top = \mathcal{H}^\top$, so $z_{\mathcal{H}} = \mathcal{H} \mathcal{H}^\top z$. Rewrite $z_{\perp \mathcal{H}_q}$ and $z_{\perp \mathcal{H}}$ as

$$z_{\perp \mathcal{H}_q} = \frac{(I_p - \mathcal{H}_q \mathcal{H}_q^\top)z}{|(I_p - \mathcal{H}_q \mathcal{H}_q^\top)z|}, \quad z_{\perp \mathcal{H}} = \frac{(I_p - \mathcal{H} \mathcal{H}^\top)z}{|(I_p - \mathcal{H} \mathcal{H}^\top)z|}. \quad (2.60)$$

Further note

$$\mathcal{H}\mathcal{H}^\top = \sum_{i=1}^K h_i h_i^\top, \quad (2.61)$$

where $h_1, h_2, \dots, h_K$ are other orthogonal columns of $\mathcal{H}$. Therefore,

$$(I_p - \mathcal{H}\mathcal{H}^\top)z = (I_p - \mathcal{H}_q\mathcal{H}_q^\top)z - \sum_{i=q+1}^K h_i h_i^\top z \quad (2.62)$$

Recall Lemma 2.3.3, we have $\langle h_i, z \rangle \rightarrow 0$ for $i = q+1, q+2, \dots, K$, So we finish this proof. $\square$

By Lemma 2.6.2 and the continuity of limits, we have

$$\left\| \mathcal{H}_{\#,q} - \left(\mathcal{H}_q \Psi_q^2 + \frac{z_{\perp \mathcal{H}} z^\top \mathcal{H}_q (I_q - \Psi_q^2)}{|z - z_{\mathcal{H}}|} \right) \right\| \rightarrow 0, \quad (2.63)$$

and thus

$$\|\mathcal{H}_{\#} - (\mathcal{H}_{\#,q}, 0, \dots, 0)\| \rightarrow 0. \quad (2.64)$$

Now,

$$\begin{aligned} \|\mathcal{E}_p(\mathcal{H}_{\#}) - \mathcal{E}_p(\mathcal{H}_{\#,q})\| &= \left\| \frac{\mathcal{B}^\top z - (\mathcal{B}^\top \mathcal{H}_{\#})(\mathcal{H}_{\#}^\top z)}{\sqrt{1 - |\mathcal{H}_{\#}^\top z|^2}} - \frac{\mathcal{B}^\top z - (\mathcal{B}^\top \mathcal{H}_{\#,q})(\mathcal{H}_{\#,q}^\top z)}{\sqrt{1 - |\mathcal{H}_{\#,q}^\top z|^2}} \right\| \\ &\rightarrow 0, \end{aligned} \quad (2.65)$$

by Lemma 2.3.3 and 2.6.2, which concludes our proof by combining with $\mathcal{E}_p(\mathcal{H}_{\#,q}) \rightarrow 0$ [16]. $\square$

2.6.5 Proofs for Section 2.4.2

Proof of Theorem 2.4.3 and Corollary 2.4.5. The proofs of the theorem and the corollary are presented together, as Corollary 2.4.5 serves as an intermediate step in the proof of Theorem 2.4.3.

We first compute the $K \times 1$ vector $\mathcal{H}_\#^\top z$ and $q \times K$ matrix $\mathcal{B}^\top \mathcal{H}_\#$. Since

$$\begin{aligned}
 z_{\perp \mathcal{H}}^\top z &= \frac{z^\top z - (\mathcal{H} \mathcal{H}^\top z)^\top z}{|z - z_{\mathcal{H}}|} \\
 &= \frac{z^\top (I_p - \mathcal{H} \mathcal{H}^\top) z}{\sqrt{(z - \mathcal{H} \mathcal{H}^\top z)^\top (z - \mathcal{H} \mathcal{H}^\top z)}} \\
 &= \frac{z^\top (I_p - \mathcal{H} \mathcal{H}^\top) z}{\sqrt{z^\top (I_p - \mathcal{H} \mathcal{H}^\top)^\top (I_p - \mathcal{H} \mathcal{H}^\top) z}} \\
 &= \sqrt{z^\top (I_p - \mathcal{H} \mathcal{H}^\top) z} = \sqrt{1 - \|\mathcal{H}^\top z\|^2},
 \end{aligned} \tag{2.66}$$

we have

$$\begin{aligned}
 \mathcal{H}_\#^\top z &= \left(\Psi^2 \mathcal{H}^\top + \frac{(I_K - \Psi^2) \mathcal{H}^\top z z_{\perp \mathcal{H}}^\top}{|z - z_{\mathcal{H}}|} \right) z \\
 &= \Psi^2 \mathcal{H}^\top z + (I_K - \Psi^2) \mathcal{H}^\top z \\
 &= \mathcal{H}^\top z,
 \end{aligned} \tag{2.67}$$

and therefore,

$$\begin{aligned}
 (\mathcal{B}^\top \mathcal{H}_\#)(\mathcal{H}_\#^\top z) &= \mathcal{B}^\top \mathcal{H} \Psi^2 \mathcal{H}^\top z + \mathcal{B}^\top z_{\perp \mathcal{H}} \frac{z^\top \mathcal{H} (I_K - \Psi^2)}{|z - z_{\mathcal{H}}|} \mathcal{H}^\top z \\
 &= \mathcal{B}^\top \left(\mathcal{H} \Psi^2 \mathcal{H}^\top + z_{\perp \mathcal{H}} \frac{z^\top \mathcal{H} (I_K - \Psi^2)}{|z - z_{\mathcal{H}}|} \right) z \\
 &= \mathcal{B}^\top \left[(I_p - C) \mathcal{H} \Psi^2 + C \mathcal{H} \right] \mathcal{H}^\top z,
 \end{aligned} \tag{2.68}$$

by letting the $p \times p$ matrix $C = \frac{z_{\perp \mathcal{H}} z_{\perp \mathcal{H}}^\top}{\sqrt{1 - \|\mathcal{H}^\top z\|^2}}$.

Claim 2.6.3. *The following statement holds:*

$$\|(I_p - C)\mathcal{H}\Psi^2 + C\mathcal{H}\mathcal{H}^\top - I_p\| \not\rightarrow 0. \quad (2.69)$$

Proof. We prove this claim by contradiction. Assume $\|(I_p - C)\mathcal{H}\Psi^2 + C\mathcal{H}\mathcal{H}^\top - I_p\| \rightarrow 0$, then take the part within $\|\cdot\|$ for left side of Eq 2.69 and multiply by $\mathcal{H}^\top$ to the left. Then

$$\begin{aligned} & \mathcal{H}^\top(I_p - C)\mathcal{H}\Psi^2\mathcal{H}^\top + \mathcal{H}^\top C\mathcal{H}\mathcal{H}^\top - \mathcal{H}^\top \\ &= \Psi^2\mathcal{H}^\top - \mathcal{H}^\top C\mathcal{H}\Psi^2\mathcal{H}^\top + \mathcal{H}^\top C\mathcal{H}\mathcal{H}^\top - \mathcal{H}^\top \\ &= \Psi^2\mathcal{H}^\top + \mathcal{H}^\top C\mathcal{H}(I_K - \Psi^2)\mathcal{H}^\top - \mathcal{H}^\top \\ &= (\Psi^2 + \mathcal{H}^\top C\mathcal{H}(I_K - \Psi^2) - I_K)\mathcal{H}^\top, \end{aligned} \quad (2.70)$$

which implies

$$\|\Psi^2 + \mathcal{H}^\top C\mathcal{H}(I_K - \Psi^2) - I_K\| \rightarrow 0. \quad (2.71)$$

Denote $D = I_K - \Psi^2$, D is a positive definite diagonal matrix. Continuing on Eq 2.71, we have

$$\|(\mathcal{H}^\top C\mathcal{H})D - D\| = \|(\mathcal{H}^\top C\mathcal{H} - I_K)D\| \rightarrow 0. \quad (2.72)$$

The smallest singular value of D satisfies $\sigma_{\min}(D) \geq \delta > 0$ for some positive number δ , and therefore

$$\|\mathcal{H}^\top C\mathcal{H} - I_K\| \leq \frac{\|(\mathcal{H}^\top C\mathcal{H} - I_K)D\|}{\|D\|} \leq \frac{\|(\mathcal{H}^\top C\mathcal{H} - I_K)D\|}{\delta} \rightarrow 0, \quad (2.73)$$

and hence leads to the contradiction as $\|C - I_p\| \not\rightarrow 0$ as $p \uparrow \infty$.

□

Claim 2.6.4. *Let $R \in \mathbb{R}^{p \times p}$ be a linear operator such that $\|R\| \not\rightarrow 0$. Let $z \in \mathbb{R}^p$ be a fixed unit vector, and let $B \in \mathbb{R}^{p \times q}$ be a full-rank matrix (i.e., $\text{rank}(B) = q$).*

If $Rz \not\rightarrow 0$ and $Rz \notin \ker(B^\top)$, then:

$$\|B^\top Rz\| \not\rightarrow 0.$$

Proof. Since $z \in \mathbb{R}^p$ is a unit vector and $\|R\| \not\rightarrow 0$, it follows that $Rz \not\rightarrow 0$; in other words, the vector Rz remains bounded away from zero in norm:

$$\|Rz\| \geq \delta > 0.$$

Next, because $B \in \mathbb{R}^{p \times q}$ is full-rank, its transpose $B^\top$ has full row rank. Therefore, the singular values of $B^\top$ are strictly positive, and in particular, the minimum singular value satisfies:

$$\sigma_{\min}(B^\top) > 0.$$

Now, if $Rz \notin \ker(B^\top)$, then $B^\top Rz \neq 0$. Furthermore, we apply the inequality for operator norms:

$$\|B^\top Rz\| \geq \sigma_{\min}(B^\top) \cdot \|Rz\| \geq \sigma_{\min}(B^\top) \cdot \delta > 0.$$

Hence, $\|B^\top Rz\| \not\rightarrow 0$, as desired. $\square$

Now, denote $R = I_p - (I_p - C)\mathcal{H}\Psi^2 + C\mathcal{H}\mathcal{H}^\top$, we have

$$\|\mathcal{B}^\top z - (\mathcal{B}^\top \mathcal{H}_\#)(\mathcal{H}_\#^\top z)\| = \|\mathcal{B}^\top Rz\| \not\rightarrow 0, \quad (2.74)$$

and thus when $K < q$,

$$\mathcal{E}_p(\mathcal{H}_\#) \not\rightarrow 0 \quad (2.75)$$

as $p \uparrow \infty$.

Next, we consider $\|\mathcal{E}_p(\mathcal{H}_\#)\|^2$. Recall in the theorem, we denote $v = \mathcal{H}^\top z$, $\alpha = \mathcal{B}^\top \mathcal{H} \Psi^2 v$, $u = \mathcal{B}^\top (I_p - \mathcal{H} \mathcal{H}^\top) z$ and $r = z^\top \mathcal{H} (I_K - \Psi^2) v$, and denote $M = I_p - R$, we can write the numerator of $\mathcal{E}_p(\mathcal{H}_\#)$ as

$$\mathcal{B}^\top R z = \mathcal{B}^\top (I_p - M) z = \mathcal{B}^\top z - \left(\alpha + \frac{r}{1 - \|v\|^2} u \right),$$

with

$$\begin{aligned} \|\mathcal{B}^\top M z\|^2 &= \left\| \alpha + \frac{r}{1 - \|v\|^2} u \right\|^2 \\ &= \|\alpha\|^2 + \left(\frac{r}{1 - \|v\|^2} \right)^2 \|u\|^2 + \frac{2r}{1 - \|v\|^2} \langle \alpha, u \rangle. \end{aligned} \quad (2.76)$$

Note that

$$\|u\| = \|\mathcal{B}^\top z - (\mathcal{B}^\top \mathcal{H})(\mathcal{H}^\top z)\| = \sqrt{1 - \|v\|^2} \|\mathcal{E}_p(\mathcal{H})\|, \quad (2.77)$$

we then have

$$\|\mathcal{B}^\top M z\|^2 = \|\alpha\|^2 + \frac{r^2}{1 - \|v\|^2} \|\mathcal{E}_p(\mathcal{H})\|^2 + \frac{2r}{1 - \|v\|^2} \langle \alpha, u \rangle, \quad (2.78)$$

and thus

$$\begin{aligned} \|\mathcal{B}^\top R z\|^2 &= \|\mathcal{B}^\top z\|^2 + \|\mathcal{B}^\top M z\|^2 - 2\langle \mathcal{B}^\top z, \mathcal{B}^\top M z \rangle \\ &= \|\mathcal{B}^\top z\|^2 + \|\alpha\|^2 + \frac{r^2}{1 - \|v\|^2} \|\mathcal{E}_p(\mathcal{H})\|^2 + \frac{2r}{1 - \|v\|^2} \langle \alpha, u \rangle \\ &\quad - 2 \left(z^\top \mathcal{B} \alpha + \frac{r}{1 - \|v\|^2} z^\top \mathcal{B} u \right) \\ &= \|\mathcal{B}^\top z\|^2 + \|\alpha\|^2 + \frac{r^2}{1 - \|v\|^2} \|\mathcal{E}_p(\mathcal{H})\|^2 + \frac{2r}{1 - \|v\|^2} \langle \alpha - \mathcal{B}^\top z, u \rangle \\ &\quad - 2\langle \mathcal{B}^\top z, \alpha \rangle \\ &= \|\mathcal{B}^\top z - \alpha\|^2 + \frac{r^2}{1 - \|v\|^2} \|\mathcal{E}_p(\mathcal{H})\|^2 - \frac{2r}{1 - \|v\|^2} \langle \mathcal{B}^\top z - \alpha, u \rangle, \end{aligned} \quad (2.79)$$

and

$$\|\mathcal{E}_p(\mathcal{H}_\#)\|^2 = \frac{1}{1 - \|v\|^2} \|\mathcal{B}^\top R z\|^2 = \frac{\|\mathcal{B}^\top z - \alpha\|^2}{1 - \|v\|^2} + \frac{r^2 \|\mathcal{E}_p(\mathcal{H})\|^2 - 2r \langle \mathcal{B}^\top z - \alpha, u \rangle}{(1 - \|v\|^2)^2} \quad (2.80)$$

□

2.6.6 Proofs for Section 2.4.3

Proof of Thm 2.4.6. We start with part (a).

1. Recall we can write $\mathcal{B} = (b, b_2, \dots, b_q)$, so the optimization bias $\mathcal{E}_p(h_\#)$ is

$$\mathcal{E}_p(h_\#) = \frac{\begin{pmatrix} \langle b, z \rangle \\ \vdots \\ \langle b_q, z \rangle \end{pmatrix} - \begin{pmatrix} \langle b, h_\# \rangle \langle h_\#, z \rangle \\ \vdots \\ \langle b_q, h_\# \rangle \langle h_\#, z \rangle \end{pmatrix}}{\sqrt{1 - \langle h_\#, z \rangle^2}}. \quad (2.81)$$

$\mathcal{E}_p(h_\#) \neq \mathcal{E}_p(\mathcal{H}_\#)$ is a direct result from the above equation by noting the 2nd to the q th elements in $\mathcal{E}_p(h_\#)$ and $\mathcal{E}_p(\mathcal{H}_\#)$ are different. To prove $\mathcal{E}_p(h_\#) \not\rightarrow 0$, we only need to show $\mathcal{E}_{p,1}(h_\#)$ is bounded away from zero.

2. We have

$$\begin{aligned} \langle h_\#, z \rangle &= \frac{\psi^2 \langle h, z \rangle}{D} + \frac{N \langle z_{\perp h}, z \rangle}{D} \\ &= \frac{\psi^2 \langle h, z \rangle}{D} + \frac{N}{D} \frac{1 - \langle z_h, z \rangle}{|z - z_h|} \\ &= \frac{\psi^2 \langle h, z \rangle}{D} + \frac{N}{D} \sqrt{1 - \langle h, z \rangle^2} \\ &= \frac{\psi^2 \langle h, z \rangle}{D} + \frac{\langle h, z \rangle - \psi^2 \langle h, z \rangle}{D} \\ &= \frac{\langle h, z \rangle}{D} \end{aligned} \quad (2.82)$$

and

$$\begin{aligned}
\langle h_{\sharp}, b \rangle &= \frac{\psi^2}{D} \langle h, b \rangle + \frac{N}{D} \frac{\langle z, b \rangle - \langle h, z \rangle \langle h, b \rangle}{\sqrt{1 - \langle h, z \rangle^2}} \\
&= \frac{\psi^2}{D} \langle h, b \rangle + \frac{(\langle h, z \rangle - \psi^2 \langle h, z \rangle)(\langle z, b \rangle - \langle h, z \rangle \langle h, b \rangle)}{D(1 - \langle h, z \rangle^2)} \\
&= \frac{\psi^2 \langle h, b \rangle + (1 - \psi^2) \langle h, z \rangle \langle z, b \rangle - \langle h, b \rangle \langle h, z \rangle^2}{D(1 - \langle h, z \rangle^2)}.
\end{aligned} \tag{2.83}$$

Therefore,

$$\begin{aligned}
\mathcal{E}_{p,1}^*(h_{\sharp}) &= \frac{D \langle b, z \rangle}{\sqrt{D^2 - \langle h, z \rangle^2}} - \frac{\psi^2 \langle h, b \rangle \langle h, z \rangle + (1 - \psi^2) \langle h, z \rangle^2 \langle b, z \rangle - \langle h, z \rangle^3 \langle h, b \rangle}{D(1 - \langle h, z \rangle^2) \sqrt{D^2 - \langle h, z \rangle^2}} \\
&= \frac{(\psi^2 - \langle h, z \rangle^2)(\psi^2 \langle b, z \rangle - \langle h, b \rangle \langle h, z \rangle)}{D(1 - \langle h, z \rangle^2) \sqrt{D^2 - \langle h, z \rangle^2}} \\
&= \frac{\psi^2 \langle b, z \rangle - \langle h, b \rangle \langle h, z \rangle}{\sqrt{\psi^4 + (1 - 2\psi^2) \langle h, z \rangle^2}}
\end{aligned} \tag{2.84}$$

3. Recalling from [16], we can define $h_z = (h \ z_{\perp h})$, $t_* = \frac{(\langle h, b \rangle, \mathcal{E}_p(h))}{\sqrt{\langle h, b \rangle^2 + \mathcal{E}_p^2(h)}}$ and $t_{\sharp} = t_* \frac{\langle h, b \rangle}{|\langle h, b \rangle|}$.

Note $\mathcal{E}_p(h_z t_*) = 0$ and $h_{\sharp} = h_z t_{\sharp}$. Define $\tilde{\mathcal{E}}_p(\cdot) : t \mapsto \mathcal{E}_p(h_z t)$ and note that $\langle h_z t, z \rangle < 1$ and $\langle b, z \rangle - \langle b, h_z t \rangle \langle h_z t, z \rangle = \langle b, z \rangle - t^2 \langle b, h_z \rangle \langle h_z, z \rangle$ is continuous w.r.t t , we have $t \mapsto \tilde{\mathcal{E}}_p(t)$ is continuous in $\mathbb{R}$.

Therefore when $K = 1 = q$, we have (1) $\tilde{\mathcal{E}}_p(t)$ is continuous (2) $\tilde{\mathcal{E}}_p(t_{\sharp}) = \tilde{\mathcal{E}}_p(t_{\sharp} - t_* + t_*)$ (3) $\tilde{\mathcal{E}}_p(t_*) = \mathcal{E}_p(h_z t_*) = 0$ (4) $|t_{\sharp} - t_*| \rightarrow 0$ and thus $\mathcal{E}_p(h_{\sharp}) \rightarrow 0$. When $K = 1 < q$, the condition $|t_{\sharp} - t_*|$ breaks so $\mathcal{E}_{p,1}(h_{\sharp}) \not\rightarrow 0$.

Part (b) is the direct result from Lemma 2.6.1 and Theorem 2.3.7.

For part (c), we have

$$\frac{Y^{\top} b b^{\top} Y}{p} = \frac{1}{p} X^{\top} B^{\top} b b^{\top} B X + \frac{1}{p} \epsilon^{\top} b b^{\top} B X + \frac{1}{p} X^{\top} B^{\top} b b^{\top} \epsilon + \frac{1}{p} \epsilon^{\top} b b^{\top} \epsilon. \tag{2.85}$$

Comparing with

$$L = \frac{Y^\top Y}{p} = \frac{1}{p} X^\top B^\top B X + \frac{1}{p} \epsilon^\top B X + \frac{1}{p} X^\top B^\top \epsilon + \frac{1}{p} \epsilon^\top \epsilon, \quad (2.86)$$

We have

$$\frac{Y^\top b b^\top Y}{p} = L + M_p - \Gamma_p + \frac{X^\top D_1 X}{p} + \frac{\epsilon^\top D_2 X}{p} + \frac{X^\top D_3 \epsilon}{p}, \quad (2.87)$$

where $M_p = \frac{\epsilon^\top b b^\top \epsilon}{p}$, $\Gamma_p = \frac{\epsilon^\top \epsilon}{p}$, $D_1 = B^\top b b^\top B - B^\top B$, $D_2 = b b^\top B - B$ and $D_3 = B^\top b b^\top - B^\top$. For simplicity, we also denote $N_p = \frac{X^\top D_1 X}{p} + \frac{\epsilon^\top D_2 X}{p} + \frac{X^\top D_3 \epsilon}{p}$.

Recall (λ_1^2, ω) is the leading eigenpair of L and denote $\mathcal{W} = \omega \sqrt{p/(n s_p^2)}$, then we have (because $\langle \omega, L \omega \rangle = \lambda_1^2$),

$$\begin{aligned} \langle h, b \rangle^2 &= \omega^\top (L + M_p - \Gamma_p + N_p) \omega / \lambda_1^2 \\ &= 1 + \omega^\top (M_p - \Gamma_p + N_p) \omega / \lambda_1^2. \end{aligned} \quad (2.88)$$

Combined with $\psi^2 = 1 - l_p^2/s_{1,p}^2$, we can define $\phi^2 = \psi^2 + \mathcal{W}^\top (N_p + \frac{n l_p^2}{p} I_n - \Gamma_p) \mathcal{W}$ and $|\langle h, b \rangle^2 - \phi^2| \rightarrow 0$.

To analysis the relationship between $\langle h, b \rangle^2$ and q , we start with following lemmas.

Lemma 2.6.5.

1. $n l_p^2 / p \mathcal{W}^\top \mathcal{W} = \frac{\sum_{j \geq 1} \lambda_j^2}{\lambda_1^2(n-1)}$ is bounded
2. $\mathcal{W}^\top \Gamma_p \mathcal{W} \sim \frac{p \gamma^2}{n s_p^2} \rightarrow \frac{\gamma^2}{\lambda_1^2}$ as $p \rightarrow \infty$.
3. $\mathcal{W}^\top N_p \mathcal{W}$ is decreasing in q .

Proof.

1. Note that $|\omega| = 1$, so $\mathcal{W}^\top \mathcal{W} = \frac{p}{ns_p^2}$. So We have

$$\frac{nl_p^2}{p} \mathcal{W}^\top \mathcal{W} = \frac{\sum_{j>1} \lambda_j^2}{n-1} \quad (2.89)$$

So $nl_p^2/p \mathcal{W}^\top \mathcal{W} = \frac{\sum_{j>1} \lambda_j^2}{\lambda_1^2(n-1)}$ is bounded.

2. It directly comes from Assumption 2.3.1, $\Gamma_p \rightarrow \gamma^2 I$ and $\mathcal{W}^\top \mathcal{W} = p/(ns_p^2)$
3. We discuss on the D_1, D_2, D_3 terms separately.

(a) $D_1 = B^\top b b^\top B - B^\top B$ with $B^\top B = \text{diag}(|\beta_1|^2, \dots, |\beta_q|^2)$ and $B^\top b b^\top B = |\beta_1|^2 E_{11}$ where E_{11} is a $q \times q$ matrix with the 1, 1th element equals to 1 and others equal to 0. So $D_1 = \text{diag}(0, -|\beta_2|^2, |\beta_3|^2, \dots, -|\beta_q|^2)$. Denote the i th row of X be X_i , then

$$X^\top D_1 X = \sum_{i=1}^q D_{1,ii} (X_i)^\top X_i = - \sum_{i=2}^q |\beta_i|^2 (X_i^\top X_i). \quad (2.90)$$

and

$$\frac{1}{p} \mathcal{W}^\top X^\top D_1 X \mathcal{W} = - \frac{1}{ns_p^2} \sum_{i=2}^q |\beta_i|^2 (X_i \omega)^2, \quad (2.91)$$

for X_i being the i -th row of X . Each term in the sum is non-negative, and therefore $\frac{1}{p} \mathcal{W}^\top X^\top D_1 X \mathcal{W}$ is decreasing in q .

- (b) $D_2 = B^\top b b^\top - B^\top$. $B^\top b b^\top$ is a $q \times p$ matrix with first row equals to $|\beta_1| b^\top$ and the others equal to 0. Therefore, D_2 is a $q \times p$ matrix with first row equals to 0 and the i th row equals to $-\beta_i^\top$ for $i = 2, 3, \dots, q$. Then

$$\epsilon^\top D_2^\top X = - \sum_{i=2}^q (\epsilon^\top \beta_i) X_i \quad (2.92)$$

with $\epsilon^\top D_2^\top X/p \rightarrow 0$. We have

$$\frac{1}{p} \mathcal{W}^\top \epsilon^\top D_2^\top X \mathcal{W} = -\frac{1}{ns^2} \sum_{i=2}^q (X_i \omega) (\epsilon \omega)^\top \beta_i. \quad (2.93)$$

for $X_{.i}$ being the i -th column of X and it goes to 0.

(c) The discussion on $D_3 = bb^\top B - B$ and $X^\top D_3^\top \epsilon$ follows the similar discussion as above, and $X^\top D_3^\top \epsilon/p \rightarrow 0$. Also

$$\frac{1}{p} \mathcal{W}^\top X^\top D_3^\top \epsilon \mathcal{W} = -\frac{1}{ns^2} \sum_{i=2}^q (\beta_i^\top \epsilon \omega) (\omega^\top X_{.i}). \quad (2.94)$$

for $X_{.i}$ being the i -th column of X and it goes to 0.

□

Therefore we finish the proof of part (c).

Lastly for part (d), recall $\omega = u \sqrt{\frac{p}{ns_{1,p}^2}}$, where u is the $n \times 1$ right singular vector of Y . We have

$$h = \frac{Yv}{\sqrt{p\lambda_1^2}} = \frac{1}{\sqrt{p\lambda_1^2}} (BX\omega + \epsilon v), \quad (2.95)$$

and

$$\langle h, z \rangle = h^\top z = \frac{1}{\sqrt{p\lambda_1^2}} (v^\top X^\top B^\top z + v^\top \epsilon^\top z). \quad (2.96)$$

Recall we write $B = (\beta_1, \beta_2, \dots, \beta_q)$ with each β_i orthogonal to each other, we have

$$\langle h, z \rangle = \frac{1}{\sqrt{p\lambda_1^2}} \left(v^\top \epsilon^\top z + \sum_{i=1}^q (X\omega)_i (\beta_i^\top z) \right), \quad (2.97)$$

with $(X\omega)_i$ being the i -th element of the $q \times 1$ vector $X\omega$. Therefore

$$\begin{aligned} \langle h, z \rangle^2 &= \frac{1}{p\lambda_1^2} \left(\left(\sum_{i=1}^q (X\omega)_i \langle \beta_i, z \rangle \right)^2 + 2 \left(\sum_{i=1}^q (X\omega)_i \langle \beta_i, z \rangle \right) v^\top \epsilon^\top z + (v^\top \epsilon^\top z)^2 \right) \\ &= \frac{1}{p\lambda_1^2} \left(\sum_{i=1}^q (X\omega)_i^2 \langle \beta_i, z \rangle^2 + 2 \left(\sum_{i=1}^q (X\omega)_i \langle \beta_i, z \rangle \right) v^\top \epsilon^\top z + (v^\top \epsilon^\top z)^2 \right) \end{aligned} \quad (2.98)$$

by using $\beta_i^\top \beta_j = 0$ for $i \neq j$. The first term in above equation is increasing in q . The middle term and last term goes to 0 from Assumption 2.3.1. Write $\beta_i = b_i |\beta_i|$ for $i = 1, 2, \dots, q$, we have $\langle \beta_i, z \rangle^2 = |\beta_i|^2 \langle b_i, z \rangle^2$, and recall $\frac{|\beta_i|^2}{p} \rightarrow \eta_i$. Thus

$$\begin{aligned} \frac{1}{p\lambda_1^2} \sum_{i=1}^q (X\omega)_i^2 \langle \beta_i, z \rangle^2 &= \frac{1}{p} \sum_{i=1}^q (X\omega)_i^2 |\beta_i|^2 \langle b_i, z \rangle^2 \\ &= \frac{1}{\lambda_1^2} \sum_{i=1}^q \eta_i (X\omega)_i^2 \langle b_i, z \rangle^2. \end{aligned} \quad (2.99)$$

And we can now write

$$\left\| \langle h, z \rangle^2 - \frac{1}{\lambda_1^2} \sum_{i=1}^q \eta_i (X\omega)_i^2 \langle b_i, z \rangle^2 \right\| \rightarrow 0, \quad (2.100)$$

and is increasing in q . □

Remark 2.6.6. When $q = 1$, note that $v \rightarrow \frac{X}{|X|}$, so we have $\langle h, z \rangle^2 \sim \frac{|\beta|^2}{p} v^\top X^\top X \omega \langle b, z \rangle^2 = \langle h, b \rangle^2 \langle b, z \rangle^2$.

Chapter 3

A Stochastic Block Model With Preferences and James-Stein Estimation for Spiked Wigner Matrices

In this chapter, we extend the James–Stein shrinkage framework for eigenvectors beyond the high-dimension, low-sample-size (HDLSS) regime by allowing the sample size to grow along with the dimension.

In Chapter 2 our analysis focused on settings where the number of observations n remains fixed while the dimension p increases. There we showed that, under spiked covariance model with true rank q and fixed n , a James–Stein type correction of the sample eigenvectors can asymptotically eliminate optimization bias when the rank is correctly specified, and we characterized how rank misspecification ($K \neq q$) affects this bias. However in practice, many modern applications happen in regimes where both the data dimension and the sample size are large, and the spectral behavior is well stud-

ied in random matrix theory (see [85] for key results in RMT). Classical results such as the Marchenko–Pastur law [25], Wigner’s semicircle law [86], and their modern refinements by Erdős, Yau, and collaborators [87, 88, 89, 90], describe the limiting eigenvalue distributions of large random matrices and the delocalization of their eigenvectors. In this asymptotic regime, the sample covariance exhibits stochastic geometry rather than deterministic bias.

[?] studied the model of this form:

$$Y_N = BB^\top + W_N,$$

where $BB^\top$, W_N are all $N \times N$ ¹ symmetric matrix, L_N has rank $r \leq n$ and W_N is orthogonally invariant. They uncovered the phase transition that the extreme eigenvalues of Y_N deviate from the bulk limit if and only if the eigenvalues of L_N exceed a critical threshold determined by the eigenvalues of W_N . It coincides with the BBP transition and is accompanied by an analogous shift in the eigenvectors.

In this chapter, we adopt spiked Wigner perspective as a special case of above, i.e. W_N is a Wigner matrix. Beyond the classical BBP transition, a growing body of work has extended the spiked Wigner framework to nonlinear and structured settings.

[91] established one of the first rigorous analyses of finite-rank deformations of large Wigner matrices. They showed that sufficiently strong spikes produce outlier eigenvalues beyond the semicircle bulk, and that while their locations are universal, their fluctuations depend on the distribution of the noise entries. [92] analyzed nonlinear Wigner spiked models and showed that the critical signal-to-noise threshold depends on the first nonzero generalized information coefficient of the nonlinearity. [93] further investigated the transformed spiked Wigner model, and proved the largest eigenvalue follows a Gaus-

¹To clarify we now use N as dimensions instead of n .

sian law above the critical signal-to-noise ratio and a Tracy–Widom law below it. In the context of hypothesis testing, [94] studied detection problems in spiked random matrix models, deriving sharp phase transition thresholds and demonstrating how entry-wise transformations can improve statistical power when the noise is non-Gaussian. Separately, [95] investigated block-structured spiked Wigner models, demonstrating that the optimal spectral method achieves the BBP threshold even in the presence of inhomogeneous or correlated noise.

In our work, we do the following: (1) find the limit of eigenvalues of Y_N in eigenvalues of L_N ; (2) show the eigenvectors of Y_N is not consistent estimator of those of L_N ; (3) provided a James-Stein estimator of spiked eigenvectors of L_N ; (4) introduce a new model called stochastic block model with preference, and illustrate James-Stein estimator we derived provides a better estimator comparing to sample eigenvectors.

3.1 Literature Review

The understandings of eigenvector behavior for Wigner matrices originally comes from a sequence of works by Erdos and his collaborators. The earliest work appeared in [88], where they showed eigenvectors of Wigner matrices cannot concentrate their mass on small subsets of coordinates with overwhelming probability. Shortly after, [87] provided local semicircle law estimates in the bulk, and showed bulk eigenvectors are of order $O(N^{-1/2})$, implying every eigenvector is completely delocalized.

These results were extended by [89], and they developed a general framework for universality of generalized Wigner matrices. Their analysis showed that the local semicircle law holds without assuming identical distributions across matrix entries and that eigenvectors are fully delocalized under broad conditions. Moreover, they gave estimates on the matrix elements of the Green function, which demonstrated bulk and edge universal-

ity under four- and two-moment matching. [96] further studied a dynamical perspective based on Dyson Brownian motion (DBM), which provided precise resolvent estimates and connected eigenvector delocalization to the relaxation dynamics of the matrix flow.

Rather than analysing the norms of eigenvectors, [97] characterized the distribution of eigenvectors. They proved that eigenvectors near the spectral edge exhibit universal Gaussian fluctuations under two-moment matching, while bulk universality requires four-moment matching. Soon after, [98] provided the isotropic delocalization bound for the eigenvectors of matrices of the form X^*X . [99] introduced the eigenvector moment flow. Building on the DBM framework, they proved bulk and edge quantum unique ergodicity and restated the isotropic local semicircle law and isotropic delocalization of eigenvectors. Besides, [100] studied Wigner matrices with uniformly subexponential entries, and showed delocalization with the optimal rate $\sqrt{\log(N)/N}$ holds with arbitrary small probability as long as optimal delocalization constant for the ℓ^∞ norm of eigenvectors of generalized Wigner matrices.

Beyond homogeneous Wigner matrices, delocalization phenomena have been studied in other settings. [101] considered random matrices of the form $H = W + \lambda V$, and showed top eigenvectors remain delocalized when λ is below a critical threshold but become partially localized when λ exceeds it. In parallel, [102] demonstrated delocalization for random matrices with independent entries under minimal assumptions by employing geometric approaches. Finally, the behavior of eigenvectors in sparse or structured settings has been clarified by results such as [103], where they showed a universal localization–delocalization transition for Gaussian matrices with deterministic sparsity patterns.

Recent advances have extended these ideas to general Wigner-type matrices and stochastic block models. [104] proved a local law and eigenvector delocalization for general Wigner-type matrices, specialized their results to sparse SBM with finite many blocks, and extended to infinite number of blocks under certain conditions. For sparse

SBMs specifically, [105] established a weak local semicircle law and obtained Tracy–Widom fluctuations at the spectral edge, and implied bulk eigenvectors of sparse SBMs are fully delocalized.

3.2 Motivation

Modern high-dimensional data analysis frequently relies on spectral methods, i.e. procedures that use the eigenstructure of large random matrices to uncover latent low-dimensional patterns. Principal component analysis, spectral clustering, and factor modeling all depend on the accurate estimation of eigenvectors that represent meaningful directions in the data. Yet in many contemporary applications, the data dimension and sample size are of comparable magnitude, so classical consistency results for eigenvectors no longer hold. Empirical directions fluctuate substantially across realizations, and their alignment with the underlying population structure becomes partial and noisy.

This instability motivates a careful study of regularized eigenvector estimators. Among various approaches to stabilization, James–Stein–type shrinkage occupies a special conceptual role. In Chapter 2, we examine the inconsistency principle in the context of rank misspecification, i.e. the number of estimated latent components K differs from the true dimension q in a low-rank factor model. There, the shrinkage adjustment was shown to mitigate the bias to some extent: when $K > q$, it asymptotically removed optimization bias, whereas for $K < q$, a nonvanishing bias persisted. That analysis revealed how James–Stein regularization operates in finite-sample high-dimensional settings.

The present chapter extends this perspective from model-based misspecification to asymptotic geometric randomness. Rather than assuming a fixed latent structure, we now consider eigenvectors that fluctuate intrinsically for underlying data matrix being Wigner ensemble. In this proportional regime, eigenvectors are delocalized due to the

stochastic geometry of the data itself, and thus applying James–Stein shrinkage in this setting is conceptually different but philosophically related. Our goal is still to stabilize an uncertain direction by borrowing strength from a reference, but we now isolate the effect of high-dimensional spectral randomness without introducing an additional issue of model misspecification.

Translating the James–Stein idea to eigenvector estimation under random matrix asymptotics remains nontrivial. Eigenvectors are nonlinear functions of the data, constrained to the unit sphere, and their fluctuations depend on global spectral interactions rather than independent coordinate noise. Nonetheless, the underlying intuition persists: if an empirical eigenvector fluctuates around a population direction with substantial random dispersion, one might systematically shrink it toward a stable reference direction to reduce mean-squared misalignment. In this chapter, we explore whether such shrinkage can yield measurable improvement in the proportional limit.

3.3 Problem formulation

Let W be a $N \times N$ Wigner matrix, i.e., $W = (w_{ij})_{1 \leq i, j \leq N}$ where the upper triangular coefficients w_{ij} for $i \leq j$ are jointly independent, mean zero and of unit variance. The lower triangular entries w_{ij} for $i < j$ satisfy $w_{ij} = w_{ji}$. For simplicity we assume real entries (but our theory applied directly to complex entries, i.e., Hermitian W). We will work with the following rescaled matrix.

$$W_N = W/\sqrt{N} \tag{3.1}$$

One example of W (and W_N) that arises in graph theory applications is that of the symmetric Bernoulli ensemble ([106]). Here, W equals the centered adjacency matrix of

an Erdos-Reyni graph as discussed in Section 1.1.2.

In the limit $N \rightarrow \infty$, the eigenvalues of W_N are well-known to converge (almost surely) in distribution to the so called semi-circular law μ_{sc} , i.e.,

$$\mu_{\text{sc}}(z) = \frac{1}{2\pi} \sqrt{(4 - x^2)_+} dx \quad (x \in [-2, 2]). \quad (3.2)$$

This probability law is deterministic and has support on $[-2, 2]$. To make the statement more precise, let μ_N denote the empirical measure of the eigenvalues of the matrix W_N , i.e., for $\lambda_i(W_N)$ the i th largest eigenvalue of W_N , define

$$\mu_N(\mathcal{A}) = \frac{1}{N} \sum_{i=1}^N 1_{\mathcal{A}}(\lambda_i(W_N)) \quad (\text{Borel } \mathcal{A} \subset \mathbb{R}), \quad (3.3)$$

where $1_{\mathcal{A}}$ is the indicator function of the set $\mathcal{A}$.

The random measure μ_N evaluates to the fraction $\mu_N(\mathcal{A})$ of the eigenvalues of the random matrix W_N that fall within the set $\mathcal{A}$. Regarding (3.2), we have

$$\mu_N \Rightarrow \mu_{\text{sc}} \quad (3.4)$$

as $N \rightarrow \infty$ almost surely (e.g., [107]).

The utility of the measure μ_N stems from its connection to the trace of a function f of the matrix W_N (i.e. $f(W_N) = V f(\Lambda) V^\top$ where $V \Lambda V^\top$ is the eigen-decomposition of W_N (see [108])). In particular,

$$\frac{1}{N} \text{tr} f(W_N) = \int_{\mathbb{R}} f(\lambda) d\mu_N(\lambda)$$

where the left side is the trace of the matrix $f(W_N)$. Letting $f(\lambda) = \frac{1}{\lambda - z}$ for z not in the

spectrum of W_N yields the so-called Stieltjes transform of μ_N , i.e.,

$$s_N(z) = \frac{1}{N} \text{tr}((W_N - zI)^{-1}) = \int_{\mathbb{R}} \frac{d\mu_N(\lambda)}{\lambda - z}.$$

By (3.4), we have $s_N(z) \rightarrow s_{\text{sc}}(z)$ almost surely as $N \rightarrow \infty$ where s_{sc} is the Stieltjes transform of μ_{sc} that satisfies (see Appendix 3.6.3 for the derivation)

$$s_{\text{sc}}(z) = \int_{\mathbb{R}} \frac{d\mu_{\text{sc}}(\lambda)}{\lambda - z} = \frac{-z + \sqrt{z^2 - 4}}{2} \quad (z \in \mathbb{C}_+). \quad (3.5)$$

The matrix W_N in (3.1) discussed so far represents noise and its spectrum does not carry any structure beyond that of the semicircular law (3.2). There are no outlier eigenvalues in the limit and all eigenvectors are delocalized.

To make use of this random starting point in applications, we add a “signal” to the random matrix W_N that takes the form of a few “spiked” eigenvalues.

3.3.1 The spiked Wigner matrix model

We suppose that we observe a symmetric $N \times N$ data matrix Y_N of the form

$$Y_N = BB^\top + W_N \quad (3.6)$$

for the scaled Wigner matrix W_N per (3.1) and where B is a full rank matrix of dimensions $N \times r$ for some fixed (rank) r . The low dimensional component

$$L_N = BB^\top \quad (3.7)$$

of the data Y_N is a symmetric, positive semidefinite matrix which has rank r and dimensions $N \times N$. In particular, we write its nonzero eigenvalues as follows.

$$\lambda_1(L_N) \geq \cdots \geq \lambda_r(L_N) > 0$$

These eigenvalues, when they are outside of the support $[-2, 2]$ of the spectrum of W_N , are often referred to as spikes. As L_N is latent, it is a natural question if they may be estimated from the spectrum of Y_N , which is observed.

The following theorem generalizes the corresponding results of [109] and [110] to our Wigner random matrix setting. For their approach to apply, one of the matrices W_N or L_N must be assumed (at least) to be orthogonally invariant.

Theorem 3.3.1. *Suppose that $\lambda_j(L_N) \rightarrow \ell_j > 1$ for $j = 1, \dots, r$. Then, for $\lambda_j(Y_N)$ denoting the j th (with $j \leq r$) largest eigenvalue of Y_N ,*

$$\lim_{N \rightarrow \infty} \lambda_j(Y_N) = \ell_j + \frac{1}{\ell_j} \quad (3.8)$$

almost surely. Moreover, if $\ell_j \leq 1$ then $\lambda_j(Y_N) \rightarrow 2$ almost surely.

Remark 3.3.2. *The condition $\ell_j > 1$ is analogous to the BBP threshold. We mention in passing that we allow $\ell_j = \infty$ in which case the behavior of $\lambda_j(Y_N)$ is consistent in the sense that $\lambda_j(Y_N)/\ell_j \rightarrow 1$ as $N \rightarrow \infty$ almost surely.*

In the applications we have in mind, Theorem 3.3.1 is applied to obtain estimates of the unknown spiked eigenvalues $\lambda_j(L_N)$. By (3.8), for large N ,

$$\lambda_j(Y_N) \approx \lambda_j(L_N) + \frac{1}{\lambda_j(L_N)}$$

which becomes an equality in the limit $N \rightarrow \infty$. Consequently,

$$\ell_j^2 - \lambda_j(Y_N)\ell_j + 1 \approx 0$$

with an equality in the limit so that

$$\hat{\ell}_{j,N} = \frac{\lambda_j(Y_N) + \sqrt{\lambda_j^2(Y_N) - 4}}{2} \quad (3.9)$$

forms a consistent estimator of $\ell_j > 1$.

Remark 3.3.3. *If we choose the smaller root $\frac{\lambda_j(Y_N) - \sqrt{\lambda_j^2(Y_N) - 4}}{2}$, then from Theorem 3.3.1, we have in the limit $\ell_j = \frac{(\ell_j + \ell_j^{-1}) - \sqrt{(\ell_j + \ell_j^{-1})^2 - 4}}{2} = \frac{(\ell_j + \ell_j^{-1}) - (\ell_j - \ell_j^{-1})}{2} = \ell_j^{-1}$, which contradicts to $\ell_j > 1$.*

3.3.2 Spiked eigenvector inconsistency

Having derived estimators (e.g., (3.9)) for the spiked eigenvalues, it is natural to ask if the same can be accomplished for eigenvectors. We shed insight into the bias of eigenvectors, and in Section 3.4 we derive superior estimators. Note,

$$\mathbb{E}(Y_N) = L_N = BB^\top \quad (3.10)$$

since $Y_N = L_N + W_N$ per (3.6) and the entries of the scaled Wigner matrix W_N have zero mean. As the estimation of the spiked eigenvalues and eigenvectors entails a low dimensional approximation of Y_N , we write

$$Y_N = HH^\top + G_N \quad (3.11)$$

where H is a $N \times r$ full rank matrix of eigenvectors of Y_N corresponding to its r largest eigenvalues (i.e., spikes). The residual G_N satisfies $H^\top G_N = 0_r$ and $H^\top H$ is a diagonal matrix of the spiked eigenvalues $\lambda_j(Y_N)$ for $j = 1, \dots, r$.

In view of (3.11), without loss of generality we assume the columns of B are orthogonal, so that any column β of $L_N = \mathbb{E}(Y_N)$ is an eigenvector, i.e.,

$$\mathbb{E}(Y_N)\beta = |\beta|^2 \beta \quad b = \frac{\beta}{|\beta|} \quad (3.12)$$

where the length squared, $|\beta|^2 = \langle \beta, \beta \rangle$, is the eigenvalue corresponding to column β and the unit length normalized version of the eigenvector is b . For the j th column we have $|\beta|^2 = \lambda_j(L_N)$. Similarly, for η a column of H in (3.11),

$$Y_N \eta = |\eta|^2 \eta \quad h = \frac{\eta}{|\eta|} \quad (3.13)$$

where $|\eta|^2$ is the corresponding eigenvalue and h is the unit-length normalized eigenvector. As before, for the j th column we have $|\eta|^2 = \lambda_j(Y_N)$.

In the notation we have set up above, Theorem 3.3.1 lead to the approximation

$$H^\top H \approx B^\top B + (B^\top B)^{-1}$$

which becomes an equality almost surely in the limit $N \rightarrow \infty$.

Comparing (3.12) and (3.13) justifies the statement that η is a biased estimator of β (and similarly for h and b). The inner product $\langle h, b \rangle$ is the cosine of the angle between h and b and is one way of measuring the inconsistency of principal component analysis. Our next result concerns this quantity.

From (3.5) and restricting s_{sc} to $z \in \mathbb{R}_+$, we compute the derivative,

$$s'_{\text{sc}}(z) = \frac{z}{\sqrt{z^2 - 4}} - \frac{1}{2}. \quad (3.14)$$

Theorem 3.3.4. *Let h and b in $\mathbb{R}^N$ denote the (unit-length) eigenvectors of Y_N and L_N corresponding to the j th largest eigenvalues $\lambda_j(Y_N)$ and $\lambda_j(L_N)$ respectively with $1 \leq j \leq r$. Recalling s'_{sc} in (3.14) and $\hat{\ell}_{j,N}$ in (3.9), we define*

$$\psi_N = \frac{1}{\sqrt{\hat{\ell}_{j,N} s'_{\text{sc}}(\lambda_j(Y_N))}}. \quad (3.15)$$

Then, if $\lambda_j(L_N) \rightarrow \ell_j > 1$ we have that $\psi_N^2 \rightarrow \frac{1}{\ell_j s'_{\text{sc}}(\ell_j + 1/\ell_j)} \in (0, 1]$ and

$$\lim_{N \rightarrow \infty} \frac{\langle h, b \rangle^2}{\psi_N^2} = 1 \quad (\text{almost surely}). \quad (3.16)$$

Moreover, $\psi_N^2 \rightarrow 1$ when $\ell_j = \infty$. Lastly, if $\ell_j \leq 1$ above or h and b correspond to eigenvalues $\lambda_j(Y_N)$ and $\lambda_k(L_N)$ for $j \neq k$ then $\langle h, b \rangle^2 \rightarrow 0$ almost surely.

That ψ_N is eventually in $(0, 1)$ unless $\ell_j = \infty$ tells us that the sample eigenvectors h are inconsistent estimates of their population counterparts b . The fact that ψ_N is eventually bounded away from zero highlights that for $\ell_j > 1$ the spiked eigenvalues inform us that the true eigenvectors may at least partially be reconstructed. We next seek to improve the naive estimator h .

3.4 James-Stein type estimator

Theorem 3.3.4 concerns the signal noise separation in (3.11) as would be performed in principal component analysis. Combining these results with Theorem 3.3.1 allows us

to develop limiting expressions for the matrix of inner products,

$$H^\top B \tag{3.17}$$

between the sample eigenvectors H and the true matrix B . Indeed, the diagonal entries of (3.17) are zero in the limit and j th diagonal entry is of the form

$$\langle \eta, \beta \rangle = |\beta| |\eta| \langle h, b \rangle = \langle h, b \rangle \sqrt{\lambda_j(L_N) \lambda_j(Y_N)} \tag{3.18}$$

where η and β correspond to the j th columns of H and B respectively.

When $\langle h, b \rangle$ does not tend to 1 as $N \rightarrow \infty$, the low rank recovery $HH^\top$ is a biased estimate of the true low-rank component $L_N = BB^\top$. Since the limit of (3.17) is eventually a diagonal matrix, unlike Chapter 2, the sample eigenvectors in H may be treated independently of one another to derive their corresponding James-Stein estimators. Thus, we consider the shrinkage formula,

$$\eta(c) = \eta c + g(1 - c) \tag{3.19}$$

for some shrinkage target $g \in \mathbb{R}^N$ and shrinkage intensity $c \in [0, 1]$. In this context, one can seek optimal values of c that optimizes the error metrics,

$$|\eta(c) - \beta|^2 \quad \text{and} \quad \frac{\langle \eta(c), \beta \rangle}{|\eta(c)| |\beta|}. \tag{3.20}$$

The mean-squared error $|\eta(c) - \beta|^2$ (or the Euclidian length squared) measures errors in the eigenvalue $|\eta(c)|$ (relative to $|\beta|^2$) as well as the direction vector $h(c) = \eta(c)/|\eta(c)|$, a unit length vector in $\mathbb{R}^N$ (relative to $b = \beta/|\beta|$). The second metric in (3.20) concerns only the angle between $h(c)$ and b .

The optimal values of c in (3.19) will depend on the unknown quantities $\langle h, b \rangle$ and $\langle b, z \rangle$ where $z = g/|g|$. Theorem 3.3.4 allows for consistent estimation of $\langle h, b \rangle^2$, and we will propose a solution to the fact that the sign of $\langle h, b \rangle$ is arbitrary (i.e., eigenvectors are defined upto sign only). Consistent estimates of the inner product $\langle b, z \rangle$ will be facilitated by the main result of Section 3.4.1.

3.4.1 A concentration of measure result

We derive a concentration of measure result that characterizes the impact that the noise matrix W_N has on the sample eigenvector h relative to b with respect to any (independent) direction vector $z \in \mathbb{R}^N$ of unit-length.

Theorem 3.4.1. *Let h and b in $\mathbb{R}^N$ denote the (unit-length) eigenvectors of Y_N and L_N corresponding to the j th largest eigenvalues $\lambda_j(Y_N)$ and $\lambda_j(L_N)$ respectively with $1 \leq j \leq r$. Then, for any sequence of vectors $g \in \mathbb{R}^N$ indexed by $N = 1, 2, \dots$, each independent of h , and with $z = g/|g|$, we have*

$$\lim_{N \rightarrow \infty} |\langle h, z \rangle - \langle h, b \rangle \langle b, z \rangle| = 0 \quad (\text{almost surely}). \quad (3.21)$$

In words, this says that the angle of h away from the point z is at least as large as the one between b and z . It is strictly larger provided that $\langle b, z \rangle$ is bounded away from zero and $|\langle h, b \rangle|$ is eventually in $(0, 1)$. The latter occurs whenever the limit ℓ_j of $\lambda_j(L_N)$ satisfies $\ell_j > 1$ as in Theorem 3.3.4.

Figure 3.1 illustrates the relationship in (3.21) geometrically.

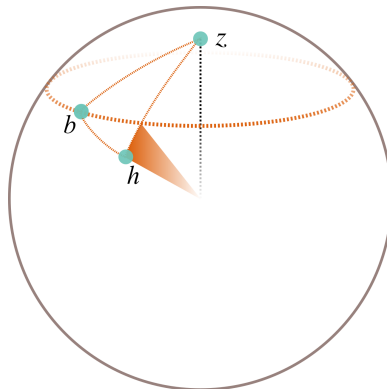


Figure 3.1: An illustration of the consequences of (3.21) with the angle between h and z larger than the one between b and z . Source: [64]

3.5 Numerical Simulation

3.5.1 Stochastic block model with preferences

Recall in Sec 1.1.2 we gave a brief introduction of Stochastic block model, and studied the eigen behavior of the expected adjacency matrix $\mathbb{E}(A)$. However, the classical SBM assumes that all vertices within the same community have identical expected degrees. This homogeneity assumption is often violated in real-world networks. For example, in a scientific collaboration network, researchers naturally form communities corresponding to research areas, but individual researchers differ substantially in their activity levels: senior investigators and large research groups form many more collaborations than early-career researchers. Such degree variability cannot be captured by the classical two-block SBM, even when the block structure is correctly specified.

To incorporate degree heterogeneity while preserving the two-community structure, we consider a preference-based SBM (heterogeneous SBM). Each vertex i is assigned a

positive preference weight $\beta_i > 0$ with normalization

$$\sum_{i=1}^N \beta_i^2 = N, \quad \sum_{i \in V_1} \beta_i^2 = \sum_{i \in V_2} \beta_i^2 \quad (3.22)$$

Edges are then generated according to

$$\mathbb{P}(A_{ij} = 1) = a_{ij} = \begin{cases} p_N \beta_i \beta_j, & i, j \in V_1 \text{ or } i, j \in V_2, \\ q_N \beta_i \beta_j, & i \in V_1, j \in V_2 \text{ or vice versa.} \end{cases}$$

for the scaling

$$p_N = p_0/N \text{ and } q_N = q_0/N.$$

We need N to be large enough to ensure $p_N, q_N \in [0, 1]$.

This formulation allows vertices within the same community to have different expected degrees while maintaining higher within-block connectivity (p_N) and lower cross-block connectivity (q_N). Setting $\beta_i \equiv 1$ recovers the classical two-block SBM, so the preference model is a strict generalization.

To illustrate the behavior of the preference-based two-block SBM, we generate a synthetic network using the parameters above. Figure 3.2 displays a realization with $N = 1000$ vertices, within-community probability $p_N = 0.05$, between-community probability $q_N = 0.005$, and preference weights drawn independently from $\text{Unif}(0.8, 1.2)$ and normalized via 3.22. Two dense communities are clearly visible, with certain vertices exhibiting higher connectivity due to the heterogeneous preference weights.

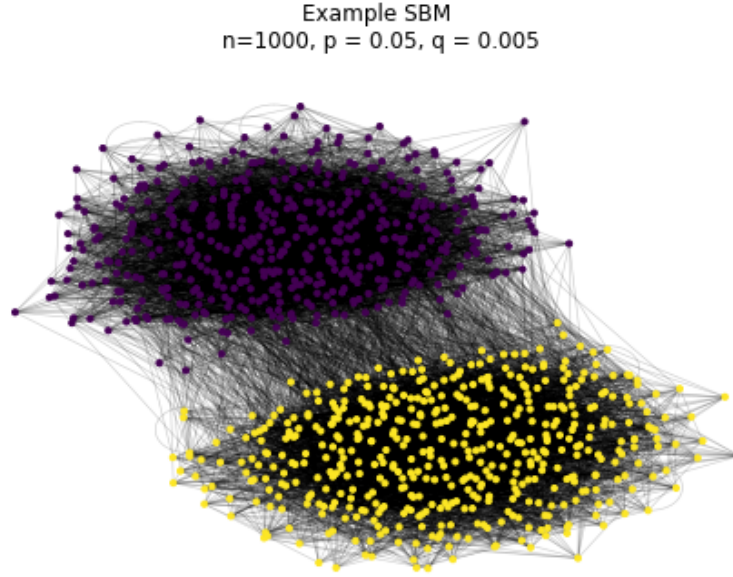


Figure 3.2: A random graph generated according to the stochastic block model $G(n, p, q)$ with $N = 1000$, $p = 0.05$, $q = 0.005$, $\beta \sim \text{Unif}(0.8, 1.2)$.

For later spectral analysis, it is convenient to introduce the mean and variance of each adjacency entry. Since $A_{ij} \sim \text{Bernoulli}(a_{ij})$ with parameter a_{ij} specified above, we have

$$a_{ij} = \mathbb{E}(A_{ij}), \quad \nu_{ij}^2 = \text{Var}(A_{ij}) = a_{ij}(1 - a_{ij}).$$

To quantify the overall fluctuation level of the random graph, we define the average variance

$$\ell_N^2 = \frac{1}{N^2} \sum_{i=1}^N \sum_{j=1}^N \nu_{ij}^2, \quad \ell_N = \sqrt{\ell_N^2}. \quad (3.23)$$

Based on these quantities, we introduce the rescaled adjacency matrix

$$A_N = \frac{1}{\ell_N \sqrt{N}} A. \quad (3.24)$$

The eigenvalue histogram of A_N for the network in Figure 3.2 is shown in Figure 3.3.

The spectrum of A_N exhibits two spike eigenvalues that lie outside the bulk ones. These two spikes arise from the rank-two structure of the mean matrix $\mathbb{E}(A_N)$ associated with the two blocks. The rest of the eigenvalues concentrate in a compact interval with a semicircle-like shape. This separation between the signal eigenvalues and the bulk were formalized in Section 3.3.1.

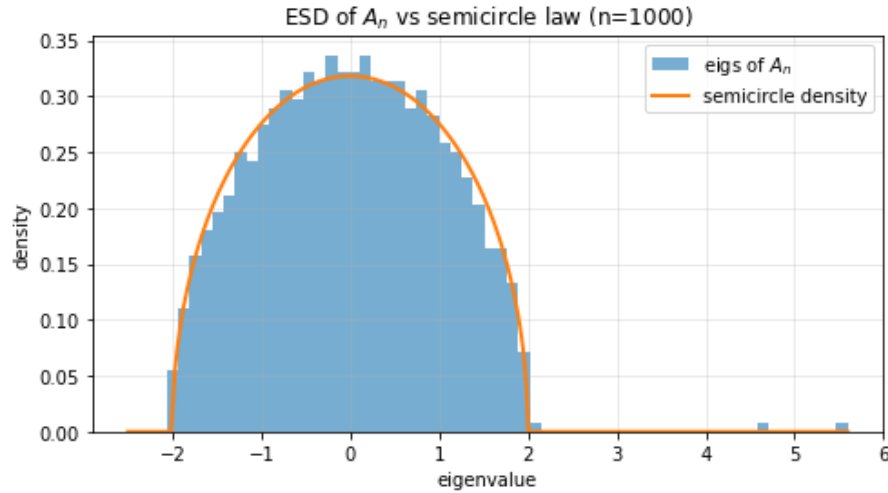


Figure 3.3: Eigenvalues of A_N scaled from adjacency matrix of Figure 3.2.

To better understand the role of probability scaling in p_N and q_N , Figure 3.4 contrasts two representative regimes.

The left panel corresponds to extremely sparse case with $p_N = p_0/N^2$ and $q_N = q_0/N^2$. In this case, the graph consists mostly of isolated vertices, and the two latent communities are no longer visually apparent. From a spectral perspective, the “spike” eigenvalues of population mean $\mathbb{E}(A_N)$ decays at rate $N^{-1/2}$ and will vanish asymptotically. In contrast, the spectral norm of the centered noise matrix $W_N = A_N - \mathbb{E}(A_N)$ collapse to a point mass at zero, with only two large, noise-induced outliers and no semicircular bulk.

The right panel corresponds to extremely dense case with $p_N = p_0/\sqrt{N}$ and $q_N = q_0/\sqrt{N}$. In this case, most edges occur with high probability, and the two latent commu-

nities remain visually distinguishable but are connected by many cross-community edges. Because the model allows self-loops and uses edge probabilities of the form $p_N\beta_i\beta_j$ and $q_N\beta_i\beta_j$, diagonal entries of A also appears to be non-zero with high probability. Therefore we observe certain number of self-loops. From a spectral perspective, the spike eigenvalues of the population mean $\mathbb{E}(A_N)$ grow at rate $N^{1/2}$ due to the normalization factor $\ell_N \asymp N^{-1/4}$, and therefore diverge as $N \rightarrow \infty$. In contrast, W_N continues to exhibit semicircular bulk with support on $[-2, 2]$.

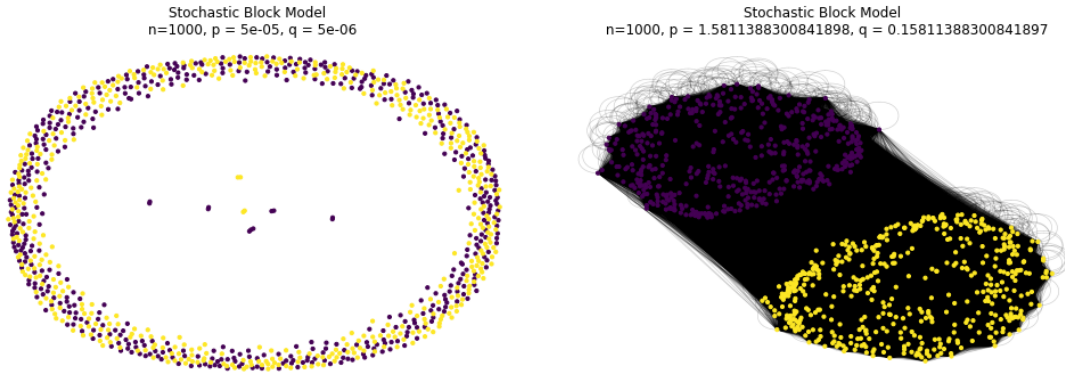


Figure 3.4: Effect of probability scaling on the observed two-block structure in preference-based SBMs. Left: $p_N, q_N = O(N^{-2})$. Right: $p_N, q_N = O(N^{-1/2})$.

3.5.2 A numerical study on preference recovery

We now conduct numerical experiments to validate the theoretical predictions regarding the spike locations, eigenvector alignment, and semicircle behavior of the bulk spectrum. We consider a two-community stochastic block model with

$$N \in \{1000, 2000, 4000, 6000\},$$

where the two communities both have size $N/2$. For each N , we generate five independent replications and report the averaged results.

The degree-heterogeneity vector β is sampled from $\text{Unif}(0.8, 1.2)$ and then rescaled to satisfy the normalization constraints. We fix $p_0 = 50$ and $q_0 = 5$, and generate edges according to the heterogeneous SBM model with $p_N = p_0/N$ and $q_N = q_0/N$. For numerical stability we truncate edge probabilities at 1 when necessary.

We then verify the semicircle behavior of the scaled noise matrix $W_N = A_N - L_N$. Figure 3.5 compares the empirical bulk spectrum with the semicircle density for each N .

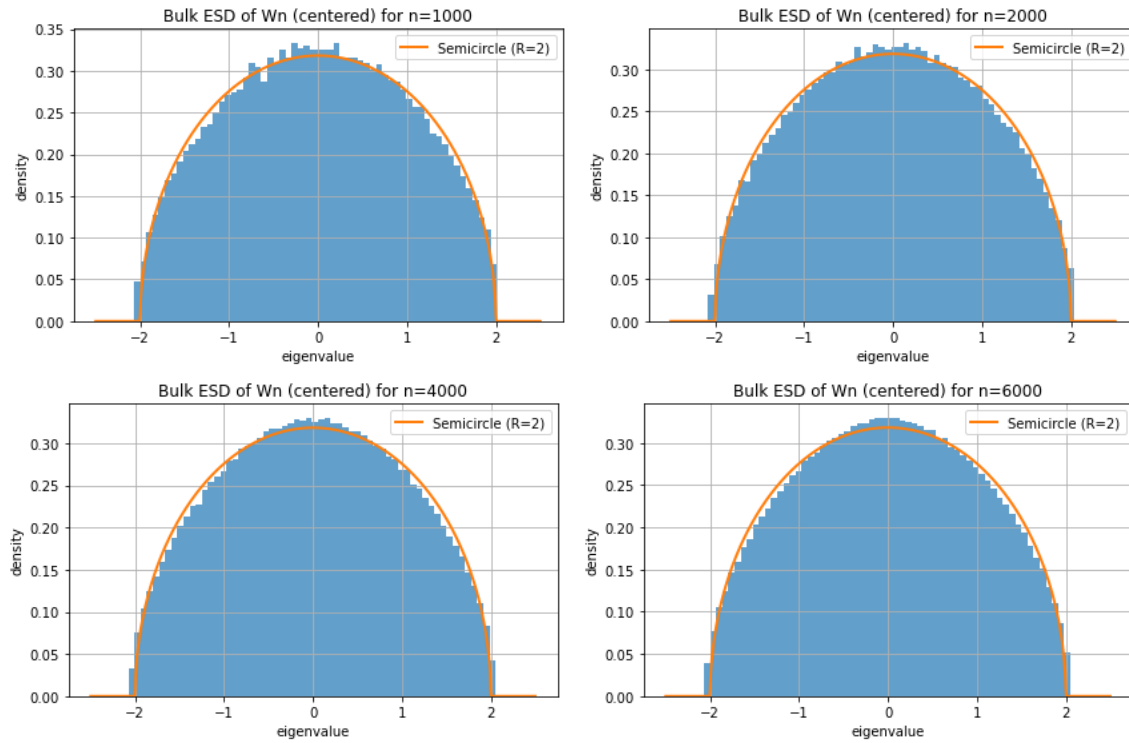


Figure 3.5: Histogram of bulk eigenvalues of W_n for $n = 1000, 2000, 4000, 6000$

Lastly, we apply James-Stein shrinkage to h , and compare the performance of h and h_{js} . Figure 3.6 displays the angle between h and b and also the angles between h_{js} and b . We can observe that h_{js} and b has smaller angle comparing to h and b , indicating h_{js} better estimates the performance direction. Furthermore, to evaluate how accurately the empirical directions recover the latent preference vector, we consider the

oracle reconstructions

$$\hat{\beta} = \langle h, \beta \rangle h, \quad \hat{\beta}_{\text{js}} = \langle h_{\text{js}}, \beta \rangle h_{\text{js}},$$

which quantify how well the PCA direction h and the James-Stein direction h_{js} align with the true preference vector β . The quality of preference recovery is measured using the relative ℓ_2 error

$$\frac{\|\hat{\beta} - \beta\|}{\|\beta\|}$$

and the numerical results are summarized in Table 3.1. We can see under different graph sizes, the James-Stein estimator reduces the relative ℓ_2 error by approximately 43%, and implies the James-Stein shrinkage enables more robust and accurate estimation of the preference vector β .

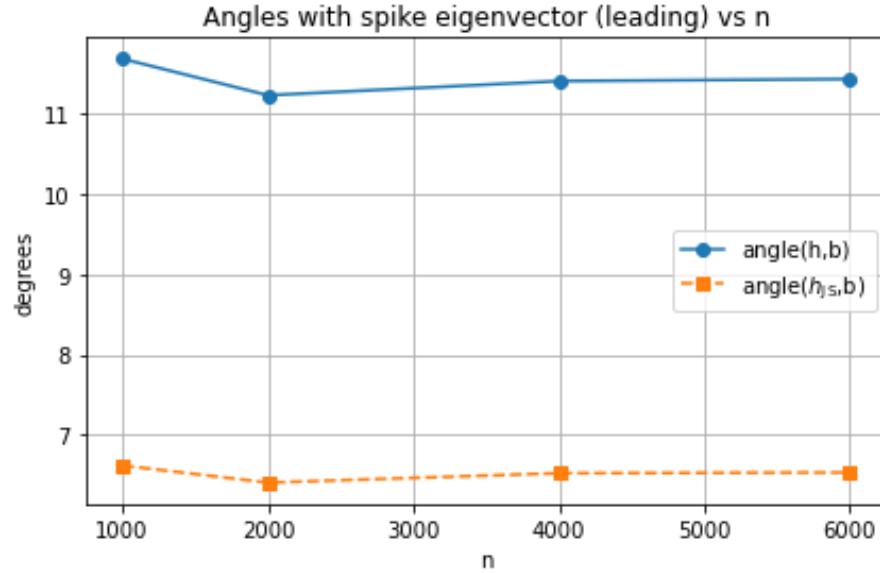


Figure 3.6: Line plot of $\arccos(h, b)$ and $\arccos(h_{\text{js}}, b)$ for $n = 1000, 2000, 4000, 6000$.

3.6 Proofs for Chapter 3

We prove Theorem 3.3.1 below and Theorems 3.3.4 and 3.4.1 together after.

Table 3.1: Relative ℓ_2 error $\|\hat{\beta} - \beta\|/\|\beta\|$ for PCA and JS estimators.

N	$\ \hat{\beta} - \beta\ /\ \beta\ $	$\ \hat{\beta}_{js} - \beta\ /\ \beta\ $	Improvement
1000	0.2025	0.1153	43.1%
2000	0.1947	0.1116	42.7%
4000	0.1977	0.1136	42.5%
6000	0.1982	0.1138	42.6%

3.6.1 Proof of Theorem 3.3.1

The former theorem concerning the eigenvalues is a corollary of Theorem 2.9 in [110]. It implies that for the spiked eigenvalue behavior is determined by the transform,

$$m(z) = \int_{\mathbb{R}} \frac{z}{z^2 - \lambda^2} \mu(dz)$$

where μ is the limit of the empirical measure μ_N of the singular values of the random matrix. In particular, the limit of the j eigenvalue tends to the inverse the function m^2 evaluated at $1/\ell_j$, the limit of the deterministic sequence $\lambda_j(L_N)$ per our assumptions.

In our case, this random matrix is W_N of Section 3.3. There $\mu_N \Rightarrow \mu_{sc}$, the semicircular law in (3.2). We prove in Section 3.6.3 below that, for the semicircular law,

$$m_{sc}(z) = -s_{sc}(z) = \frac{z - \sqrt{z^2 - 4}}{2} \quad (3.25)$$

We therefore have the limit relation,

$$m_{sc}^2(\lim_{N \rightarrow \infty} \lambda_j(Y_N)) = \frac{1}{\ell_j^2}$$

from which we obtain, letting $\lambda_j = \lim_{N \rightarrow \infty} \lambda_j(Y_N)$, that

$$\frac{\lambda_j - \sqrt{\lambda_j^2 - 4}}{2} = \frac{1}{\ell_j}$$

and since $(\ell_j + 1/\ell_j)^2 - 4 = \ell_j^2 - (1/\ell_j)2 = (\ell_j - 1/\ell_j)^2$, the limit is identified as $\lambda_j = \ell_j + 1/\ell_j$ as claimed. Since the largest value of the support for the semicircular law is 2, any j th largest eigenvalue $\lambda_j(Y_N)$ with $j > r$ must converge to 2, again by Theorem 2.9 of [110].

3.6.2 Proofs of Theorems 3.3.4 and 3.4.1

As we will be taking inner product of sample eigenvectors with $b = \beta/|\beta|$ for a column β of B in (3.6) and since the columns of B are orthogonal, it suffices to consider

$$Y_N = \beta\beta^\top + W_N$$

i.e., the simpler case of a single spiked model. Its eigen decomposition satisfies,

$$\begin{aligned} \mathcal{J}h &= Y_N h = \beta|\beta|\langle b, h \rangle + W_N h \\ &= b|\beta|^2\langle b, h \rangle + W_N h \end{aligned}$$

where h is the eigenvector of Y_N with largest eigenvalue $\mathcal{J}$ (i.e. $\lambda_1(Y_N)$). Then,

$$\begin{aligned} W_N W_N h &= \mathcal{J}W_N h - W_N b|\beta|^2\langle b, h \rangle \\ &= \mathcal{J}^2 h - \mathcal{J}W_N b|\beta|^2\langle h, b \rangle - W_N b|\beta|^2\langle b, h \rangle \\ &= \mathcal{J}^2 h - W_N b(1 + \mathcal{J})|\beta|^2\langle b, h \rangle \end{aligned}$$

which leads to,

$$h = (\mathcal{J}^2 I - W_N W_N)^{-1} W_N b (1 + \mathcal{J}) |\beta|^2 \langle h, b \rangle \quad (3.26)$$

Since $|h|^2 = 1$, we obtain the relation

$$1 = b^\top W_N (\mathcal{J}^2 I - W_N W_N)^{-1} W_N b (1 + \mathcal{J})^2 |\beta|^4 \langle h, b \rangle^2$$

By Theorem 3.3.1 we have $\lambda_1(Y_N) = \mathcal{J} \rightarrow \ell_1 + 1/\ell_1$ as $N \rightarrow \infty$ almost surely and $|\beta|^2 \rightarrow \ell_1$ under our hypotheses. Letting $\rho^2 = (\ell_1 + 1/\ell_1)^2$, the limit of $\mathcal{J}^2$ we have

$$\frac{\lim_{N \rightarrow \infty} \langle h, b \rangle^{-2}}{(1 + \ell_1 + 1/\ell_1) \ell_1^2} = \lim_{N \rightarrow \infty} b^\top W_N (\rho^2 I - W_N W_N)^{-1} W_N b$$

Since the eigenvectors of the matrices W_N , $W_N W_N$, and $\rho^2 I - W_N W_N$ all have the same eigenvectors. Thus,

$$\begin{aligned} \frac{\lim_{N \rightarrow \infty} \langle h, b \rangle^{-2}}{(1 + \ell_1 + 1/\ell_1) \ell_1^2} &= \lim_{N \rightarrow \infty} b^\top W_N (\rho^2 I - W_N W_N)^{-1} W_N b \\ &= \lim_{N \rightarrow \infty} \frac{1}{N} \text{tr}(b b^\top W_N (\rho^2 I - W_N W_N)^{-1} W_N) \end{aligned}$$

We deduce that,

$$\begin{aligned} \frac{\lim_{N \rightarrow \infty} \langle h, b \rangle^{-2}}{(1 + \ell_1 + 1/\ell_1) \ell_1^2} &= \lim_{N \rightarrow \infty} \frac{1}{N} \text{tr}(W_N (\rho^2 I - W_N W_N)^{-1} W_N) \\ &= \lim_{N \rightarrow \infty} \int_{\mathbb{R}} \frac{\lambda^2}{\rho^2 - \lambda^2} \mu_N(\lambda) \\ &= \lim_{N \rightarrow \infty} \int_{\mathbb{R}} \left(\frac{\rho^2}{\rho^2 - \lambda^2} - 1 \right) \mu_N(\lambda) \\ &= \frac{1}{\rho} m_{\text{sc}}(\rho) - 1 = -\frac{1}{\rho} s_{\text{sc}}(\rho) - 1 \end{aligned}$$

where the transform m_{sc} was defined in (3.25).

Collecting the terms above with $\rho^2 = \ell_1 + 1/\ell_1$ and noting the expression for s'_{sc} in (3.14), we obtain after collecting the terms that,

$$\lim_{N \rightarrow \infty} \langle h, b \rangle^2 = \frac{1}{\ell_j s'_{sc}(\rho)}. \quad (3.27)$$

Recalling that $\hat{\ell}_{1,N} \rightarrow \ell_1$ and $\lambda_1(Y_N) = \mathfrak{s} \rightarrow \rho$ concludes the proof of (3.16).

Taking the dot product with b on both sides of (3.26), yields

$$1 = \langle b, (\mathfrak{s}^2 I - W_N W_N)^{-1} W_N b \rangle (1 + \mathfrak{s}) |\beta|^2 \quad (3.28)$$

While taking the dot product with the vector z in the statement of Theorem 3.4.1, and following steps similar to the ones above,

$$\begin{aligned} \lim_{N \rightarrow \infty} \langle h, z \rangle &= \lim_{N \rightarrow \infty} z^\top ((\mathfrak{s}^2 I - W_N W_N)^{-1} W_N b (1 + \mathfrak{s}) |\beta|^2 \langle h, b \rangle) \\ &= \lim_{N \rightarrow \infty} \frac{1}{N} \text{tr} (z b^\top (\mathfrak{s}^2 I - W_N W_N)^{-1} W_N) (1 + \mathfrak{s}) |\beta|^2 \langle h, b \rangle \\ &= \lim_{N \rightarrow \infty} \langle b, z \rangle \frac{1}{N} \text{tr} (b b^\top (\mathfrak{s}^2 I - W_N W_N)^{-1} W_N) (1 + \mathfrak{s}) |\beta|^2 \langle h, b \rangle \\ &= \lim_{N \rightarrow \infty} \langle b, z \rangle \langle h, b \rangle \frac{1}{N} \text{tr} (b b^\top (\mathfrak{s}^2 I - W_N W_N)^{-1} W_N) (1 + \mathfrak{s}) |\beta|^2 \\ &= \lim_{N \rightarrow \infty} \langle b, z \rangle \langle h, b \rangle \langle b, (\mathfrak{s}^2 I - W_N W_N)^{-1} W_N b \rangle (1 + \mathfrak{s}) |\beta|^2 \end{aligned}$$

Applying (3.28) concludes the proof of Theorem 3.4.1.

3.6.3 Proofs of Section 3.3

We compute the transform,

$$\begin{aligned}
s_{\text{sc}}(z) &= \int_{\mathbb{R}} \frac{1}{x-z} d\mu(x) \\
&= \int_{-2}^2 \frac{1}{x-z} \frac{1}{2\pi} \sqrt{4-x^2} dx \\
&= -\frac{2}{\pi} \int_0^\pi \frac{\sin^2 \theta}{z-2\cos \theta} d\theta \\
&= \frac{1}{\pi} \int_0^\pi \frac{\cos(2\theta)-1}{z-2\cos \theta} d\theta \\
&= \frac{1}{\pi} (I_2(z) - I_0(z))
\end{aligned}$$

where $I_n(z) = \int_0^\pi \frac{\cos(n\theta)}{z-2\cos \theta} d\theta$ and $I_0(z) = \frac{\pi}{\sqrt{z^2-4}}$. Note that

$$\cos n\theta \cos \theta = \frac{1}{2} (\cos((n+1)\theta) + \cos((n-1)\theta)),$$

we divide both sides by $z-2\cos \theta$ and integrate on $[0, \pi]$, we have

$$\int_0^\pi \frac{2\cos \theta \cos n\theta}{z-2\cos \theta} d\theta = \int_0^\pi \frac{\cos((n+1)\theta)}{z-2\cos \theta} d\theta + \int_0^\pi \frac{\cos((n-1)\theta)}{z-2\cos \theta} d\theta$$

The LHS equals to

$$\int_0^\pi \left(\frac{z}{z-2\cos \theta} - 1 \right) \cos n\theta d\theta = zI_n(z) - \int_0^\pi \cos n\theta d\theta,$$

and the latter term equals to 0. Therefore

$$zI_n(z) = I_{n+1}(z) + I_{n-1}(z).$$

Next, write $\cos \theta = \frac{1}{2}(z - (z - 2 \cos \theta))$, we have

$$\begin{aligned} I_1(z) &= \int_0^\pi \frac{\cos \theta}{z - 2 \cos \theta} d\theta \\ &= \frac{z}{2} I_0(z) - \frac{\pi}{2} \\ &= \frac{\pi}{2} \left(\frac{z}{\sqrt{z^2 - 4}} - 1 \right). \end{aligned} \tag{3.29}$$

Therefore,

$$\begin{aligned} I_2(z) &= z I_1(z) - I_0(z) \\ &= \frac{z\pi}{2} \left(\frac{z}{\sqrt{z^2 - 4}} - 1 \right) - \frac{\pi}{\sqrt{z^2 - 4}} \\ &= \pi \left(\frac{z^2 - 2}{2\sqrt{z^2 - 4}} - \frac{z}{2} \right) \\ &= \frac{\pi}{\sqrt{z^2 - 4}} \left(\frac{z - \sqrt{z^2 - 4}}{2} \right)^2 \end{aligned} \tag{3.30}$$

and

$$s_{\text{sc}}(z) = \frac{1}{\pi} (I_2(z) - I_0(z)) = \frac{\sqrt{z^2 - 4} - z}{2}$$

Now we move on to $m(z)$. We have

$$m(z) = \int_{-2}^2 \frac{z}{z^2 - \lambda^2} \frac{\sqrt{4 - \lambda^2}}{2\pi} d\lambda$$

Let $\lambda = 2 \cos \theta$, $d\lambda = -2 \sin \theta d\theta$, so

$$m(z) = \frac{2z}{\pi} \int_0^\pi \frac{\sin^2 \theta}{z^2 - 4 \cos^2 \theta} d\theta = \frac{2z}{\pi} I.$$

Note that

$$I = \int_0^\pi \frac{1 - \cos^2 \theta}{z^2 - 4 \cos^2 \theta} d\theta = I_1 - I_2,$$

and

$$\begin{aligned} I_2 &= \int_0^\pi \frac{z^2 - (z^2 - 4 \cos^2 \theta)}{z^2 - 4 \cos \theta} d\theta \\ &= \frac{z^2}{4} I_1 - \frac{\pi}{4}, \end{aligned}$$

we are left to calculate I_1 . Now note that

$$z^2 - 4 \cos^2 \theta = z^2 - 2(1 + \cos 2\theta) = (z^2 - 2) - 2 \cos 2\theta,$$

so

$$\begin{aligned} I_1 &= \int_0^\pi \frac{1}{(z^2 - 2) - 2 \cos 2\theta} d\theta \\ &= \frac{1}{2} \int_0^{2\pi} \frac{1}{(z^2 - 2) - 2 \cos u} du \\ &= \frac{1}{2} \frac{2\pi}{\sqrt{(z^2 - 2)^2 - 4}} = \frac{\pi}{z\sqrt{z^2 - 4}}. \end{aligned}$$

Therefore,

$$\begin{aligned} m(z) &= \frac{2z}{\pi} I = \frac{2z}{\pi} \left(\frac{4 - z^2}{4} \right) I_1 + \frac{z}{2} \\ &= \frac{z(4 - z^2)}{2\pi} \frac{\pi}{z\sqrt{z^2 - 4}} + \frac{z}{2} \\ &= \frac{z - \sqrt{z^2 - 4}}{2} \end{aligned}$$

Bibliography

- [1] T. Hastie and R. Tibshirani, *Efficient quadratic regularization for expression arrays*, *Biostatistics* **5** (2004), no. 3 329–340.
- [2] I. M. Johnstone, *On the distribution of the largest eigenvalue in principal components analysis*, *The Annals of statistics* **29** (2001), no. 2 295–327.
- [3] O. Ledoit and M. Wolf, *A well-conditioned estimator for large-dimensional covariance matrices*, *Journal of multivariate analysis* **88** (2004), no. 2 365–411.
- [4] N. El Karoui, *Spectrum estimation for large dimensional covariance matrices using random matrix theory*, . arXiv preprint arXiv:0804.2019.
- [5] P. J. Bickel and E. Levina, *Covariance regularization by thresholding*, *The Annals of Statistics* **36** (2008), no. 6 2577–2604.
- [6] R. O. Michaud, *The markowitz optimization enigma: Is ‘optimized’ optimal?*, *Financial analysts journal* **45** (1989), no. 1 31–42.
- [7] N. E. Karoui, *On the realized risk of high-dimensional markowitz portfolios*, *SIAM Journal on Financial Mathematics* **4** (2013), no. 1 737–783.
- [8] M. Newman, *Networks*. Oxford university press, 2018.
- [9] R. Vershynin, *High-dimensional probability: An introduction with applications in data science*, vol. 47. Cambridge university press, 2018.
- [10] T. W. Anderson, *Asymptotic theory for principal component analysis*, *The Annals of Mathematical Statistics* **34** (1963), no. 1 122–148.
- [11] D. Paul, *Asymptotics of sample eigenstructure for a large dimensional spiked covariance model*, *Statistica Sinica* **17** (2007), no. 4 1617–1642.
- [12] D. Shen, H. Shen, H. Zhu, and J. Marron, *Surprising asymptotic conical structure in critical sample eigen-directions*, . arXiv preprint arXiv:1303.6171.
- [13] S. Jung, A. Sen, and J. Marron, *Boundary behavior in high dimension, low sample size asymptotics of pca*, *Journal of Multivariate Analysis* **109** (2012) 190–203.

- [14] W. James and C. Stein, *Estimation with quadratic loss*, *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability* **1** (1961) 361–379.
- [15] O. Ledoit and M. Wolf, *Analytical nonlinear shrinkage of large-dimensional covariance matrices*, *Annals of Statistics* **48** (2020), no. 5 3043–3065.
- [16] H. Gurdogan and A. Shkolnik, *The quadratic optimization bias of large covariance matrices*, . arXiv preprint arXiv:2410.03053.
- [17] W. Qiu, *Advanced Algebra*. Higher Education Press, Beijing, 2015. Title in original Chinese: .
- [18] R. A. Horn and C. R. Johnson, *Matrix analysis*. Cambridge university press, 2012.
- [19] G. Zhang, Y. Lin, and M. Guo, *Lectures on Functional Analysis*. Peking University Press, Beijing, 1990. Title in original Chinese: .
- [20] M. Zhou, *Real Analysis*. Peking University Press, Beijing, 2016. Title in original Chinese: .
- [21] A. Ben-Israel and T. N. Greville, *Generalized inverses: theory and applications*, .
- [22] Y. Yu, T. Wang, and R. J. Samworth, *A useful variant of the davis–kahan theorem for statisticians*, *Biometrika* **102** (2015), no. 2 315–323.
- [23] E. P. Wigner, *Characteristic vectors of bordered matrices with infinite dimensions i* , in *The Collected Works of Eugene Paul Wigner: Part A: The Scientific Papers*, pp. 524–540. Springer, 1993.
- [24] A. R. Feier, *Methods of proof in random matrix theory*, vol. 9. Harvard University, 2012.
- [25] V. Marchenko and L. A. Pastur, *Distribution of eigenvalues for some sets of random matrices*, *Mat. Sb.(NS)* **72** (1967), no. 114 507–536.
- [26] J. Baik, G. Ben Arous, and S. Péché, *Phase transition of the largest eigenvalue for nonnull complex sample covariance matrices*, . arXiv preprint math/0506683.
- [27] A. Bloemendal, A. Knowles, H.-T. Yau, and J. Yin, *On the principal components of sample covariance matrices*, *Probability theory and related fields* **164** (2016), no. 1 459–552.
- [28] I. Todhunter, *Spherical trigonometry, for the use of colleges and schools: with numerous examples*. Macmillan, 1863.
- [29] G. B. Folland, *A course in abstract harmonic analysis*. CRC press, 2016.

- [30] G. B. Folland, *Real analysis: modern techniques and their applications*, vol. 40. John Wiley & Sons, 1999.
- [31] B. Oksendal, *Stochastic differential equations: an introduction with applications*. Springer Science & Business Media, 2013.
- [32] I. L. Dryden, *Statistical analysis on high-dimensional spheres and shape spaces*, *The Annals of Statistics* **33** (2005), no. 4 1643–1665.
- [33] M. D. Donsker, *Justification and extension of doob’s heuristic approach to the kolmogorov-smirnov theorems*, *The Annals of mathematical statistics* (1952) 277–281.
- [34] N. Cutland and S.-A. Ng, *The wiener sphere and wiener measure*, *The Annals of Probability* (1993) 1–13.
- [35] C. Spearman, ” *general intelligence” objectively determined and measured.*, *American Journal of Psychology* **15** (1904), no. 2 201–292.
- [36] L. L. Thurstone, *Multiple factor analysis.*, *Psychological review* **38** (1931), no. 5 406–427.
- [37] L. L. Thurstone, *The vectors of mind.*, *Psychological review* **41** (1934), no. 1 1–32.
- [38] J. Bai and S. Ng, *Determining the number of factors in approximate factor models*, *Econometrica* **70** (2002), no. 1 191–221.
- [39] C. Subbarao, N. Subbarao, and S. Chandu, *Characterization of groundwater contamination using factor analysis*, *Environmental Geology* **28** (1996), no. 4 175–180.
- [40] S. Hochreiter, D.-A. Clevert, and K. Obermayer, *A new summarization method for affymetrix probe level data*, *Bioinformatics* **22** (2006), no. 8 943–949.
- [41] G. Chamberlain and M. Rothschild, *Arbitrage, factor structure, and mean-variance analysis on large asset markets*, 1982.
- [42] J. Fan, Y. Liao, and M. Mincheva, *Large covariance estimation by thresholding principal orthogonal complements*, *Journal of the Royal Statistical Society Series B: Statistical Methodology* **75** (2013), no. 4 603–680.
- [43] I. M. Johnstone and A. Y. Lu, *Sparse principal components analysis*, . arXiv preprint arXiv:0901.4392.
- [44] W. Wang and J. Fan, *Asymptotics of empirical eigenstructure for high dimensional spiked covariance*, *Annals of statistics* **45** (2017), no. 3 1342–1374.

- [45] J. H. Stock and M. W. Watson, *Forecasting using principal components from a large number of predictors*, *Journal of the American statistical association* **97** (2002), no. 460 1167–1179.
- [46] J. Bai, *Inferential theory for factor models of large dimensions*, *Econometrica* **71** (2003), no. 1 135–171.
- [47] P. Hall, J. S. Marron, and A. Neeman, *Geometric representation of high dimension, low sample size data*, *Journal of the Royal Statistical Society Series B: Statistical Methodology* **67** (2005), no. 3 427–444.
- [48] J. Ahn, J. Marron, K. M. Muller, and Y.-Y. Chi, *The high-dimension, low-sample-size geometric representation holds under mild conditions*, *Biometrika* **94** (2007), no. 3 760–766.
- [49] S. Jung and J. S. Marron, *Pca consistency in high dimension, low sample size context*, . arXiv preprint arXiv:0910.2930.
- [50] D. Shen, H. Shen, and J. S. Marron, *A general framework for consistency of principal component analysis*, *Journal of Machine Learning Research* **17** (2016), no. 150 1–34.
- [51] H. M. Markowitz, *Portfolio selection, the journal of finance. 7 (1)*, *the Journal of Finance* **7** (1952), no. 1 71–91.
- [52] J. Capon, *High-resolution frequency-wavenumber spectrum analysis*, *Proceedings of the IEEE* **57** (2005), no. 8 1408–1418.
- [53] G. C. Hegerl, H. von Storch, K. Hasselmann, B. D. Santer, U. Cubasch, and P. D. Jones, *Detecting greenhouse-gas-induced climate change with an optimal fingerprint method*, *Journal of Climate* **9** (1996), no. 10 2281–2306.
- [54] M. Fiedler, *Algebraic connectivity of graphs*, *Czechoslovak mathematical journal* **23** (1973), no. 2 298–305.
- [55] C. Stein, *Inadmissibility of the usual estimator for the mean of a multivariate normal distribution*, in *Proceedings of the third Berkeley symposium on mathematical statistics and probability, volume 1: Contributions to the theory of statistics*, vol. 3, pp. 197–207, University of California Press, 1956.
- [56] L. D. Brown, *Admissible estimators, recurrent diffusions, and insoluble boundary value problems*, *The Annals of Mathematical Statistics* **42** (1971), no. 3 855–903.
- [57] B. Efron and C. Morris, *Stein’s paradox in statistics*, *Scientific American* **236** (1977), no. 5 119–127.

- [58] A. K. Gupta and E. A. Peña, *A simple motivation for james-stein estimators*, *Statistics & probability letters* **12** (1991), no. 4 337–340.
- [59] G. Casella and J. T. Hwang, *Limit expressions for the risk of james-stein estimators*, *Canadian Journal of Statistics* **10** (1982), no. 4 305–309.
- [60] S. M. Stigler, *The 1988 neyman memorial lecture: a galtonian perspective on shrinkage estimators*, *Statistical Science* **5** (1990), no. 2 147–155.
- [61] L. D. Brown and L. H. Zhao, *A geometrical explanation of stein shrinkage*, . arXiv preprint arXiv:1206.3499.
- [62] D. Fourdrinier, W. E. Strawderman, and M. T. Wells, *Shrinkage estimation*. Springer, 2018.
- [63] A. Shkolnik, *James–stein estimation of the first principal component*, *Stat* **11** (2022), no. 1 e419.
- [64] L. R. Goldberg, A. Papanicolaou, and A. Shkolnik, *The dispersion bias*, *SIAM Journal on Financial Mathematics* **13** (2022), no. 2 521–550.
- [65] L. R. Goldberg and A. N. Kercheval, *James–stein for the leading eigenvector*, *Proceedings of the National Academy of Sciences* **120** (2023), no. 2 e2207046120.
- [66] A. Shkolnik, A. Kercheval, H. Gurdogan, L. R. Goldberg, and H. Bar, *Portfolio selection revisited*, *Annals of Operations Research* **346** (2025), no. 1 137–155.
- [67] K. Yata and M. Aoshima, *Effective pca for high-dimension, low-sample-size data with noise reduction via geometric representations*, *Journal of multivariate analysis* **105** (2012), no. 1 193–215.
- [68] Y. Chen, A. Wiesel, and A. O. Hero, *Robust shrinkage estimation of high-dimensional covariance matrices*, *IEEE Transactions on Signal Processing* **59** (2011), no. 9 4097–4107.
- [69] J. Fan, Y. Liao, and M. Mincheva, *High dimensional covariance matrix estimation in approximate factor models*, *Annals of statistics* **39** (2011), no. 6 3320–3356.
- [70] J. Fan, W. Wang, and Z. Zhu, *A shrinkage principle for heavy-tailed data: High-dimensional robust low-rank matrix recovery*, *Annals of statistics* **49** (2021), no. 3 1239–1266.
- [71] Y. Chen, J. Fan, C. Ma, and Y. Yan, *Bridging convex and nonconvex optimization in robust pca: Noise, outliers, and missing data*, *Annals of statistics* **49** (2021), no. 5 2948–2971.

- [72] L. Alessi, M. Barigozzi, and M. Capasso, *A robust criterion for determining the number of static factors in approximate factor models*, .
- [73] G. Kapetanios, *A testing procedure for determining the number of factors in approximate factor models with large datasets*, *Journal of Business & Economic Statistics* **28** (2010), no. 3 397–409.
- [74] Y. Chen and X. Li, *Determining the number of factors in high-dimensional generalized latent factor models*, *Biometrika* **109** (2022), no. 3 769–782.
- [75] A. Onatski, *Testing hypotheses about the number of factors in large factor models*, *Econometrica* **77** (2009), no. 5 1447–1479.
- [76] S. Jung, M. H. Lee, and J. Ahn, *On the number of principal components in high dimensions*, *Biometrika* **105** (2018), no. 2 389–402.
- [77] S. C. Ahn and A. R. Horenstein, *Eigenvalue ratio test for the number of factors*, *Econometrica* **81** (2013), no. 3 1203–1227.
- [78] C. Lam and Q. Yao, *Factor modeling for high-dimensional time series: inference for the number of factors*, *The Annals of Statistics* **40** (2012), no. 2 694–726.
- [79] J. Fan, J. Guo, and S. Zheng, *Estimating number of factors by adjusted eigenvalues thresholding*, *Journal of the American Statistical Association* **117** (2022), no. 538 852–861.
- [80] M. Barigozzi and H. Cho, *Consistent estimation of high-dimensional factor models when the factor number is over-estimated*, . arXiv preprint arXiv:2002.08226.
- [81] T. Cai, Z. Ma, and Y. Wu, *Optimal estimation and rank detection for sparse spiked covariance matrices*, *Probability theory and related fields* **161** (2015), no. 3 781–815.
- [82] V. Q. Vu and J. Lei, *Minimax sparse principal subspace estimation in high dimensions*, .
- [83] T. T. Cai, D. Xia, and M. Zha, *Optimal differentially private pca and estimation for spiked covariance matrices*, *arXiv preprint arXiv:2401.03820* (2024).
- [84] B. Nadler, *Finite sample approximation results for principal component analysis: A matrix perturbation approach*, *The Annals of Statistics* **36** (2008), no. 6 2791–2817.
- [85] D. Paul and A. Aue, *Random matrix theory in statistics: A review*, *Journal of Statistical Planning and Inference* **150** (2014) 1–29.
- [86] E. P. Wigner, *Random matrices in physics*, *SIAM review* **9** (1967), no. 1 1–23.

- [87] L. Erdős, B. Schlein, and H.-T. Yau, *Local semicircle law and complete delocalization for wigner random matrices*, *Communications in Mathematical Physics* **287** (2009), no. 2 641–655.
- [88] L. Erdős, B. Schlein, and H.-T. Yau, *Semicircle law on short scales and delocalization of eigenvectors for wigner random matrices*, . arXiv preprint arXiv:0905.3330.
- [89] L. Erdős, *Universality of wigner random matrices: a survey of recent results*, *Russian Mathematical Surveys* **66** (2011), no. 3 507–626.
- [90] L. Erdős, H.-T. Yau, and J. Yin, *Rigidity of eigenvalues of generalized wigner matrices*, *Advances in Mathematics* **229** (2012), no. 3 1435–1515.
- [91] M. Capitaine, C. Donati-Martin, and D. Féral, *The largest eigenvalues of finite rank deformation of large wigner matrices: convergence and nonuniversality of the fluctuations*, *The Annals of Probability* **37** (2009), no. 1 1–47.
- [92] A. Guionnet, J. Ko, F. Krzakala, P. Mergny, and L. Zdeborová, *Spectral phase transitions in non-linear wigner spiked models*, . arXiv preprint arXiv:2310.14055.
- [93] A. Lee and J. O. Lee, *Fluctuations of the largest eigenvalues of transformed spiked wigner matrices*, . arXiv preprint arXiv:2502.04720.
- [94] J. H. Jung, H. W. Chung, and J. O. Lee, *Detection problems in the spiked matrix models*, . arXiv preprint arXiv:2301.05331.
- [95] P. Mergny, J. Ko, and F. Krzakala, *Spectral phase transition and optimal pca in block-structured spiked models*, . arXiv preprint arXiv:2403.03695.
- [96] L. Erdős and H.-T. Yau, *A dynamical approach to random matrix theory*, vol. 28. American Mathematical Soc., 2017.
- [97] A. Knowles and J. Yin, *Eigenvector distribution of wigner matrices*, *Probability Theory and Related Fields* **155** (2013), no. 3 543–582.
- [98] A. Bloemendal, L. Erdős, A. Knowles, H.-T. Yau, and J. Yin, *Isotropic local laws for sample covariance and generalized wigner matrices*, . arXiv preprint arXiv:1312.5694.
- [99] P. Bourgade and H.-T. Yau, *The eigenvector moment flow and local quantum unique ergodicity*, *Communications in Mathematical Physics* **350** (2017), no. 1 231–278.
- [100] L. Benigni and P. Lopatto, *Optimal delocalization for generalized wigner matrices*, *Advances in Mathematics* **396** (2022) 108109.

- [101] J. O. Lee and K. Schnelli, *Extremal eigenvalues and eigenvectors of deformed wigner matrices*, *Probability Theory and Related Fields* **164** (2016), no. 1 165–241.
- [102] M. Rudelson and R. Vershynin, *Delocalization of eigenvectors of random matrices with independent entries*, . arXiv preprint arXiv:1506.04012.
- [103] L. Shou and R. van Handel, *A localization–delocalization transition for nonhomogeneous random matrices*, *Journal of Statistical Physics* **191** (2024), no. 2 26.
- [104] I. Dumitriu and Y. Zhu, *Sparse general wigner-type matrices: Local law and eigenvector delocalization*, *Journal of Mathematical Physics* **60** (2019), no. 2.
- [105] J. Y. Hwang, J. O. Lee, and W. Yang, *Local law and tracy–widom limit for sparse stochastic block models*, . arXiv preprint arXiv:2008.02489.
- [106] E. P. Wigner, *On the distribution of the roots of certain symmetric matrices*, *Annals of Mathematics* **67** (1958), no. 2 325–327.
- [107] T. Tao, *Topics in random matrix theory*, vol. 132. American Mathematical Society, 2023.
- [108] N. J. Higham, *Functions of matrices: theory and computation*. SIAM, 2008.
- [109] F. Benaych-Georges and R. R. Nadakuditi, *The eigenvalues and eigenvectors of finite, low rank perturbations of large random matrices*, *Advances in Mathematics* **227** (2011), no. 1 494–521.
- [110] 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) 120–135.