



Estimation of the Number of Spiked Eigenvalues in a Covariance Matrix by Bulk Eigenvalue Matching Analysis

Zheng Tracy Ke^a, Yucong Ma^a, and Xihong Lin^{a, b} 

^aDepartment of Statistics, Harvard University, Cambridge, MA; ^bDepartment of Biostatistics, Harvard T.H. Chan School of Public Health, Boston, MA

ABSTRACT

The spiked covariance model has gained increasing popularity in high-dimensional data analysis. A fundamental problem is determination of the number of spiked eigenvalues, K . For estimation of K , most attention has focused on the use of *top* eigenvalues of sample covariance matrix, and there is little investigation into proper ways of using *bulk* eigenvalues to estimate K . We propose a principled approach to incorporating bulk eigenvalues in the estimation of K . Our method imposes a working model on the residual covariance matrix, which is assumed to be a diagonal matrix whose entries are drawn from a gamma distribution. Under this model, the bulk eigenvalues are asymptotically close to the quantiles of a fixed parametric distribution. This motivates us to propose a two-step method: the first step uses bulk eigenvalues to estimate parameters of this distribution, and the second step leverages these parameters to assist the estimation of K . The resulting estimator $\hat{K}$ aggregates information in a large number of bulk eigenvalues. We show the consistency of $\hat{K}$ under a standard spiked covariance model. We also propose a confidence interval estimate for K . Our extensive simulation studies show that the proposed method is robust and outperforms the existing methods in a range of scenarios. We apply the proposed method to analysis of a lung cancer microarray dataset and the 1000 Genomes dataset.

ARTICLE HISTORY

Received May 2020
Accepted May 2021

KEYWORDS

Empirical null; Factor model; Kaiser's criterion; Latent dimension; Machenko-Pastur distribution; Parallel analysis; Principal component analysis; Unsupervised learning

1. Introduction

The spiked covariance model (Johnstone 2001) has been widely used to model the covariance structure of high-dimensional data. In this model, the population covariance matrix has K large eigenvalues, called *spiked eigenvalues*, where K is presumably much smaller than the dimension. Estimation of K is of great interest in practice, as it helps determination of the latent dimension of data. For example, in a clustering model with K_0 clusters (Jin, Ke, and Wang 2017), the pooled covariance matrix has $(K_0 - 1)$ spiked eigenvalues; therefore, an estimate of K tells the number of clusters. Similarly, in Genome-Wide Association Studies (GWAS), the number of spiked eigenvalues of a genetic covariance matrix reveals the number of ancestry groups in the study (Patterson, Price, and Reich 2006). In high-dimensional covariance matrix estimation, K is often required as input for factor-based covariance estimation (Fan, Liao, and Mincheva 2013).

In this article, we assume the data vectors $X_1, X_2, \dots, X_n \in \mathbb{R}^p$ are independently generated from a multivariate distribution with covariance matrix $\Sigma \in \mathbb{R}^{p \times p}$, which has positive values $\mu_1 \geq \mu_2 \geq \dots \geq \mu_K$ and mutually orthogonal unit-norm vectors $\xi_1, \xi_2, \dots, \xi_K \in \mathbb{R}^p$ such that

$$\Sigma = \sum_{k=1}^K \mu_k \xi_k \xi_k^\top + D, \quad \text{where } D = \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_p^2). \quad (1)$$

Here, D is called the residual covariance matrix. The goal is to estimate K from $X_1, X_2, \dots, X_n$. We are primarily interested in the settings where K is finite and $p/n \rightarrow \gamma$, for a constant $\gamma > 0$. Throughout the article, we denote by $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$ the eigenvalues of Σ , and denote by $\hat{\lambda}_1 \geq \hat{\lambda}_2 \geq \dots \geq \hat{\lambda}_{n \wedge p}$ the nonzero eigenvalues of the sample covariance matrix.

In the literature, there are several approaches for estimating K . The first is the information criterion approach, which finds $\hat{K}$ that minimizes an objective of the form $L_n(K) + P_n(K)$, where $L_n(K)$ is a measure of goodness of fit and $P_n(K)$ is a penalty on K . An influential work is Bai and Ng (2002), who let $L_n(K)$ be the sum of squared residuals after fitting a K -factor model and studied a few choices of the penalty function $P_n(K)$. Other examples include Wax and Kailath (1985), where $L_n(K)$ is a function of the arithmetic and geometric means of the $(n-K)$ smallest eigenvalues. However, the information criterion approach requires the spiked eigenvalues to be sufficiently large. In Bai and Ng (2002), the spiked eigenvalues are at the order

of p , which is much larger than the necessary order. It has been recognized that correct estimation of K is possible even when the spiked eigenvalues are at the constant order (Baik, Ben Arous, and Peche 2005).

The second approach finds a big “gap” between eigenvalues of the sample covariance matrix. Recall that $\hat{\lambda}_k$ is the k th eigenvalue of the sample covariance matrix. Onatski (2009) introduced a test statistic, $\max_{K_0 < k \leq K_{\max}} (\hat{\lambda}_k - \hat{\lambda}_{k+1})/(\hat{\lambda}_{k+1} - \hat{\lambda}_{k+2})$, for testing against the null hypothesis $K = K_0$ and then applied it sequentially to estimate K . Cai, Han, and Pan (2020) proposed an iterative algorithm for estimating K that searches for a gap of $\gtrsim O(n^{-2/3})$ between eigenvalues. Passemier and Yao (2014) suggested estimating K by finding two consecutive gaps in eigenvalues. Such methods rely on sharp limiting distributions of the first K empirical eigenvalues, which theoretically requires a large magnitude of the spiked eigenvalues. Additionally, while utilizing eigengap is a neat idea in theory, its practical use faces challenges, since the actual eigengaps in many real datasets are slowly varying, without a clear cut.

The last approach estimates K by thresholding the empirical eigenvalues. For this approach, the key is to calculate a proper data-driven threshold. The threshold should reflect the “scaling” of the residual matrix $\mathbf{D}$. One idea is to first standardize the data matrix so that each variable has a unit variance and then use a scale-free threshold. Examples include the empirical Kaiser’s criterion (Braeken and Van Assen 2017) and parallel analysis (Horn 1965), where the scale-free threshold is determined by asymptotic behavior of the largest eigenvalue of sample covariance matrix associated with $\mathbf{X}_i \stackrel{\text{iid}}{\sim} N(0, \mathbf{I}_p)$. Another idea is to estimate $\mathbf{D}$ by the diagonal of the sample covariance matrix and then calculate the threshold via a deterministic algorithm (Dobriban 2015). The success of both ideas relies on regularity conditions to ensure that the low-rank part in Model (1) has a negligible effect on the diagonal of Σ ; for example, the population eigenvalues cannot be enormously large and the population eigenvectors have to satisfy “delocalization” conditions. Dobriban and Owen (2019) improved the algorithm in Dobriban (2015) by a recursive procedure to remove leading eigenvalues and eigenvectors, but their method still requires some “delocalization” conditions on eigenvectors. Other related work includes Onatski (2010), which used a convex combination of $\hat{\lambda}_{K_{\max}+1}$ and $\hat{\lambda}_{2K_{\max}+1}$ as the threshold, where $K_{\max}$ is a prespecified upper bound of K , and Fan, Guo, and Zheng (2020), which introduced an unbiased estimator for each of the first few eigenvalues of the population correlation matrix, and estimated K by thresholding these unbiased estimators at $1 + \sqrt{p/n}$.

To address the limitations of these methods, we propose a new estimator of K . Different from the existing work, our attention is largely focused on how to better use the *bulk* empirical eigenvalues in the estimation of K , especially those eigenvalues in the middle range

$$\{\hat{\lambda}_k : \alpha(n \wedge p) \leq k \leq (1 - \alpha)(n \wedge p)\},$$

for some constant $\alpha \in (0, 1/2)$.

It is well known in random matrix theory that these bulk eigenvalues are almost not affected by the low-rank part in Model (1) (see, e.g., Bloemendal et al. 2016). We can use these eigenvalues to gauge the “scaling” of $\mathbf{D}$ and determine an appropriate

threshold for top eigenvalues. To this end, we impose a working model on the diagonal matrix $\mathbf{D}$. Let $\text{Gamma}(a, b)$ denote the gamma distribution with shape parameter a and rate parameter b . Fixing $\sigma > 0$ and $\theta > 0$, we assume

$$\sigma_j^2 \stackrel{\text{iid}}{\sim} \text{Gamma}(\theta, \theta/\sigma^2), \quad 1 \leq j \leq p. \quad (2)$$

The mean and variance of $\text{Gamma}(\theta, \theta/\sigma^2)$ is σ^2 and σ^4/θ , respectively. As a result, the diagonal entries of $\mathbf{D}$ are centered around σ^2 , where the level of dispersion is controlled by θ . As $\theta \rightarrow \infty$, $\text{Gamma}(\theta, \theta/\sigma^2)$ converges to a point mass at σ^2 , and it yields $\mathbf{D} = \sigma^2 \mathbf{I}_p$. This case corresponds to the standard spiked covariance model which is frequently studied in the literature (Johnstone 2001; Donoho, Gavish, and Johnstone 2018). Combining Model (2) with Model (1), we now have a flexible spiked covariance model that includes the standard spiked covariance model as a special case.

Under Models (1) and (2), the empirical spectral distribution (ESD) converges to a limit, which is a fixed distribution with two parameters (σ^2, θ) (Silverstein 2009). Since the empirical eigenvalues are nothing but quantiles of the ESD, we expect that all the bulk eigenvalues are asymptotically close to the corresponding quantiles of the limit of ESD. We thus estimate (σ^2, θ) by minimizing the sum of squared differences between bulk eigenvalues and quantiles of the limiting distribution. Once $(\hat{\sigma}^2, \hat{\theta})$ are available, we borrow the idea of parallel analysis (Horn 1965) to decide a threshold for the top eigenvalues by Monte Carlo sampling. This gives rise to a new method for estimating K , which we call *bulk eigenvalue matching analysis* (BEMA). Analogous to the orators’ bema in Athens, our BEMA is a platform for gathering a large number of bulk eigenvalues and utilizing them efficiently in the estimation of K . Additional to the point estimator, we also propose a confidence interval for K .

Our method has an intuitive explanation in terms of a scree plot. Figure 1 shows the scree plot of a simulated example. There are multiple elbow points, and it is hard to decide where the true K is. The core idea of our method is to explore the “shape” of the scree plot in the middle range and fit it with a parametric curve; this curve is determined by the theoretical quantiles of the limit of ESD, governed by two parameters σ^2 and θ . Then, this curve can be extended to the left boundary of the scree plot to produce a threshold for top eigenvalues.

The goodness-of-fit check of Model (2) on real datasets can also be done via the scree plot. If the middle range of the scree plot can be well approximated by the estimated parametric curve, then it suggests that the model indeed fits the real data. In Section 6, we shall see that Model (2) is well suited to gene microarray data and GWAS data. We remark that assuming the diagonal entries of $\mathbf{D}$ are generated from a fixed distribution is only a mild assumption. Similar conditions appear in the literature (often implicitly as regularity conditions in the theory); for example, Dobriban and Owen (2019) and Fan, Guo, and Zheng (2020) assumed that the histogram of population eigenvalues of $\mathbf{D}$ converges to a fixed limit. We make one step ahead by assuming that this fixed distribution is a gamma distribution. At the first glance, restricting to the gamma family seems restrictive, but Model (2) is in fact much more flexible than expected. With only two parameters (σ^2, θ) , it can accommodate various kinds of real data and even misspecified models (see Section 5).

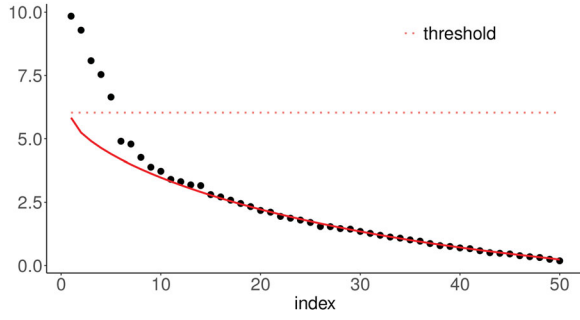


Figure 1. Illustration of BEMA via a scree plot. The red solid curve shows the quantiles of the theoretical limit of Empirical Spectral Distribution (ESD) under Models (1) and (2). It is a parametric curve with two parameters (σ^2, θ) , and by random matrix theory, it should fit the bulk eigenvalues well. BEMA first uses bulk eigenvalues to estimate (σ^2, θ) and then extends the estimated curve to the left boundary to get a threshold for top eigenvalues.

The special case of $\theta = \infty$ is of independent interest. It corresponds to the standard spiked covariance model (Johnstone 2001), where $\mathbf{D} = \sigma^2 \mathbf{I}_p$. This model has attracted a lot of attention (Baik, Ben Arous, and Peché 2005; Paul 2007; Donoho, Gavish, and Johnstone 2018). In this special case, BEMA reduces to a simpler algorithm. We conduct theoretical analysis under this model. First, we give an explicit error bound for estimating σ^2 . This is connected to the robust estimation of σ^2 in the literature of reconstruction of spiked covariance matrices (Donoho, Gavish, and Johnstone 2018; Shabalin and Nobel 2013). In our method, we obtain a new robust estimator of σ^2 as a byproduct, and we study it theoretically. Second, we prove the consistency of estimating K under minimal conditions. Our results impose no assumptions on the population eigenvectors $\xi_1, \dots, \xi_K$ and only require the spiked eigenvalues $\lambda_1, \dots, \lambda_K$ to be larger than a constant. In comparison, literature works often either require some regularity conditions on eigenvectors or need much larger spiked eigenvalues. We also provide theory for the general case of $\theta < \infty$, which has never been studied before.

The remaining of this article is organized as follows: In Section 2, we describe BEMA for the standard spiked covariance model (i.e., $\theta = \infty$); in this case, the idea is easier to understand and the algorithm is simpler. In Section 3, we describe BEMA for the general case. Section 4 states the theoretical properties. Section 5 and Section 6 provide simulation study results and real data analysis, respectively. Section 7 concludes the article. Proofs are relegated to the appendix.

2. BEMA for the Standard Spiked Covariance Model

In this section, we consider the standard spiked covariance model (Johnstone 2001), a special case of Models (1) and (2) with $\theta = \infty$. Since each σ_j^2 is equal to σ^2 , the model is rewritten as

$$\mathbf{\Sigma} = \sum_{k=1}^K \mu_k \xi_k \xi_k^\top + \sigma^2 \mathbf{I}_p. \quad (3)$$

The first K eigenvalues of $\mathbf{\Sigma}$ are $\lambda_k = \mu_k + \sigma^2$, and the remaining eigenvalues are σ^2 . The sample covariance matrix is $\mathbf{S} = \frac{1}{n} \sum_{i=1}^n (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^\top$, where $\bar{\mathbf{X}} = \frac{1}{n} \sum_{i=1}^n \mathbf{X}_i$. With probability 1, $\mathbf{S}$ has $n \wedge p$ distinct nonzero eigenvalues (Uhlhig 1994), denoted as $\hat{\lambda}_1 > \hat{\lambda}_2 > \dots > \hat{\lambda}_{n \wedge p}$.

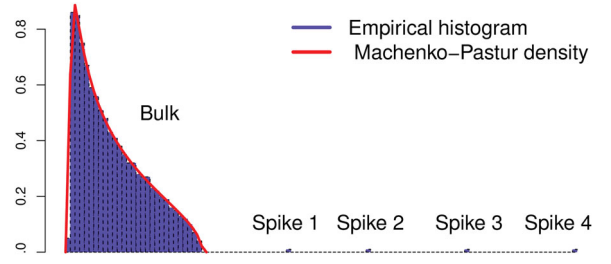


Figure 2. The asymptotic behavior of empirical eigenvalues. The histogram of bulk eigenvalues converges to an MP distribution, and K top eigenvalues are outside the support.

We first review some existing results about the asymptotic behavior of empirical eigenvalues.

Definition 1. Given a parameter $\gamma > 0$, the zero-excluded Machenko-Pastur (MP) distribution is defined by the density

$$f_\gamma(x; \sigma^2) = \frac{1}{2\pi\sigma^2} \frac{1}{x(\gamma \wedge 1)} \sqrt{(x - \sigma^2 h_-)(\sigma^2 h_+ - x)} \cdot 1\{\sigma^2 h_- < x < \sigma^2 h_+\}, \quad (4)$$

where $h_\pm = (1 \pm \sqrt{\gamma})^2$. We let $F_\gamma(x; \sigma)$ denote its cumulative distribution function.

When $\gamma \leq 1$, this definition is the same as the classical MP law; when $\gamma > 1$, it excludes the point mass at zero in the classical MP law. The zero-excluded ESD is given by $F_n(x) = \frac{1}{n \wedge p} \sum_{i=1}^{n \wedge p} 1\{\hat{\lambda}_i \leq x\}$. For convenience, we shall omit the word “zero-excluded” and still call them MP and ESD.

When $\mathbf{\Sigma}$ satisfies Equation (3), K is fixed and $p/n \rightarrow \gamma$ for a constant $\gamma \in (0, \infty)$, under mild regularity conditions, the following statements are true (Bloemendal et al. 2016):

- The ESD converges to the MP distribution with parameter γ ; more precisely, it holds that $\mathbb{E}[\sup_x |F_n(x) - F_\gamma(x)|] = O(n^{-1/2})$ (Götze et al. 2004).
- If $\mu_K \geq \sigma^2 \sqrt{\gamma} + n^{-1/3}$, then the first K empirical eigenvalues are located outside the support of the MP distribution with high probability.

See Figure 2 for an illustration via simulated data ($n = 1000$, $p = 500$).

Inspired by the asymptotic behavior of empirical eigenvalues, we propose a two-step approach to estimating K . In the first step, we use bulk eigenvalues to fit an MP distribution. The density $f_\gamma(x; \sigma^2)$ in Equation (4) has two parameters (γ, σ^2) , where γ can be approximated by $\gamma_n = p/n$. It reduces to considering $f_{\gamma_n}(x; \sigma^2)$, for all possible σ^2 . We aim to find $\hat{\sigma}^2$ such that $f_{\gamma_n}(x; \hat{\sigma}^2)$ is the best fit to the histogram of empirical eigenvalues. In the second step, we determine K by comparing top eigenvalues with the right boundary of the support of the estimated MP density, namely, $\hat{\sigma}^2(1 + \sqrt{\gamma_n})^2$.

Now, we describe the method in detail. First, consider the estimation of σ^2 . Fixing a constant $\alpha \in (0, 1/2)$, we take only a fraction of nonzero eigenvalues

$$\{\hat{\lambda}_k : \alpha(n \wedge p) \leq k \leq (1 - \alpha)(n \wedge p)\}.$$

Since K is fixed and $n \wedge p \rightarrow \infty$, any α guarantees that the first K eigenvalues are excluded. The choice of α does not matter.

Algorithm 1. BEMA for the standard spiked covariance model.

Input: Nonzero eigenvalues $\hat{\lambda}_1, \dots, \hat{\lambda}_{n \wedge p}$, $\alpha \in (0, 1/2)$ and $\beta \in (0, 1)$.

Output: An estimate of K .

Step 1: Write $\tilde{p} = n \wedge p$. For each $\alpha\tilde{p} \leq k \leq (1 - \alpha)\tilde{p}$, obtain q_k , the $(k/\tilde{p})$ -upper-quantile of the MP distribution associated with $\sigma^2 = 1$ and $\gamma_n = p/n$. Compute

$$\hat{\sigma}^2 = \frac{\sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} q_k \hat{\lambda}_k}{\sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} q_k^2}.$$

Step 2: Obtain $t_{1-\beta}$, the $(1 - \beta)$ -quantile of Tracy-Widom distribution. Estimate K by

$$\hat{K} = \#\{1 \leq k \leq \tilde{p} : \hat{\lambda}_k > \hat{\sigma}^2[(1 + \sqrt{\gamma_n})^2 + t_{1-\beta} \cdot n^{-\frac{2}{3}} \gamma_n^{-\frac{1}{6}} (1 + \sqrt{\gamma_n})^{\frac{4}{3}}]\}.$$

We usually set $\alpha = 0.2$, so that 60% of the nonzero eigenvalues in the middle range are used. Write for short $\tilde{p} = n \wedge p$. By definition, $\hat{\lambda}_k$ is the $(k/\tilde{p})$ -upper-quantile of the ESD. Let $q_k = q_k(\gamma_n)$ denote the $(k/\tilde{p})$ -upper-quantile of the MP distribution associated with $\gamma = \gamma_n$ and $\sigma^2 = 1$, that is,

$$q_k \text{ is the unique value such that } \int_{q_k}^{(1+\sqrt{\gamma_n})^2} f_{\gamma_n}(x; 1) dx = k/\tilde{p}. \quad (5)$$

These q_k 's can be easily computed (e.g., via the R package *RMTstat*). For an MP distribution with a general σ^2 , its $(k/\tilde{p})$ -upper-quantile equals to $\sigma^2 q_k$. Since the ESD is asymptotically close to the MP distribution, we expect that

$$\hat{\lambda}_k \approx \sigma^2 \cdot q_k.$$

It motivates us to use $\{(q_k, \hat{\lambda}_k)\}_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}}$ to fit a line without intercept, and this can be done by a simple least square. The slope of this line is an estimator of σ^2 .

Next, we use $\hat{\sigma}^2$ to determine a threshold for the top eigenvalues. A natural choice of threshold is $\hat{\sigma}^2(1 + \sqrt{\gamma_n})^2$, but it

has a considerable probability of over-estimating K . We slightly increase this threshold by taking an advantage of another result in random matrix theory. When $\mu_K > \sigma^2 \sqrt{\gamma}$, it is known that (Johnstone 2001; Bloemendal et al. 2016)

$$\frac{\hat{\lambda}_{K+1} - \sigma^2(1 + \sqrt{\gamma_n})^2}{\sigma^2 n^{-\frac{2}{3}} \gamma_n^{-\frac{1}{6}} (1 + \sqrt{\gamma_n})^{\frac{4}{3}}} \xrightarrow{d} \text{type-I Tracy-Widom distribution.} \quad (6)$$

We propose thresholding the top eigenvalues at

$$\hat{T} = \hat{\sigma}^2 \left[(1 + \sqrt{\gamma_n})^2 + t_{1-\beta} \cdot n^{-\frac{2}{3}} \gamma_n^{-\frac{1}{6}} (1 + \sqrt{\gamma_n})^{\frac{4}{3}} \right],$$

where $t_{1-\beta}$ denotes the $(1 - \beta)$ -quantile of the Tracy-Widom distribution. Then, the probability of over-estimating K is controlled by β .

Algorithm 1 has two tuning parameters (α, β) . The output of the algorithm is insensitive to α if α is not too small, and we set $\alpha = 0.2$ by default. β controls the probability of over-estimating K and is specified by the user. In theory, the ideal choice of β should satisfy that $\beta \rightarrow 0$ at a properly slow rate (see Section 4). In practice, choosing a moderate β often yields the best finite-sample performance. Our numerical experiments suggest that $\beta = 0.1$ is a good choice for most settings.

A simulation example. We illustrate Algorithm 1 on a simulation example. Fix $(n, p, K) = (1000, 500, 10)$. We generate $X_i \stackrel{\text{iid}}{\sim} N(0, \Sigma)$, where Σ is a diagonal matrix whose first K diagonals equal to 5.4 and the remaining diagonals equal to $\sigma^2 = 2$. In the left panel of Figure 3, we plot $\hat{\lambda}_k$ vs. q_k . Except for a few top eigenvalues, it fits well to a straight line crossing the origin. We use 300 bulk eigenvalues $\{\hat{\lambda}_k\}_{100 < k \leq 400}$ (the blue dots) to fit a regression line (the red dotted line). The slope of this line gives the estimate $\hat{\sigma}^2 = 2.04$. In the middle panel of Figure 3, we plot $\hat{\lambda}_k$ vs. k . The red solid line is the curve of $\hat{\sigma}^2 q_k$ vs. k . Although it is estimated using the blue dots only, we can extend this curve to the left boundary, which gives rise to the value $\hat{\sigma}^2(1 + \sqrt{\gamma_n})^2$. We then use this value and the Tracy-Widom distribution to calculate a threshold for the top eigenvalues. The estimator $\hat{K}$ equals to the number of top eigenvalues that exceed this threshold. The right panel of Figure 3 is a zoom-in of the

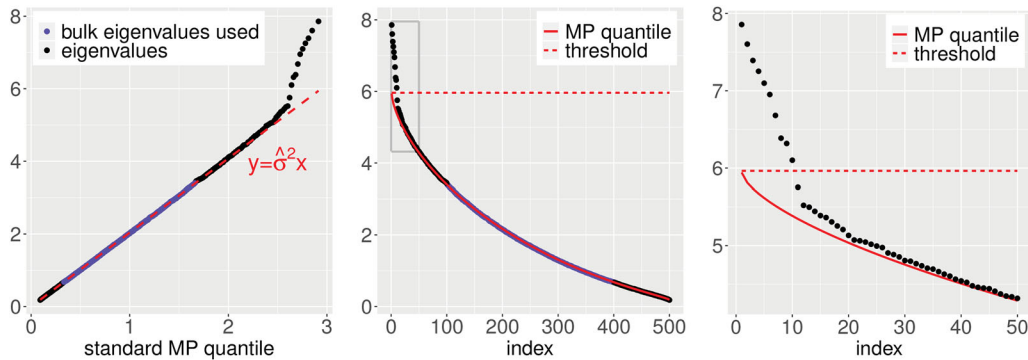


Figure 3. Illustration of BEMA for the standard spiked covariance model (simulated data, $n = 1000$, $p = 500$, $K = 10$). The left panel plots $\hat{\lambda}_k$ versus q_k , where q_k is the $(k/\tilde{p})$ -upper-quantile of the standard MP distribution. The dashed red line is the fitted regression line on bulk eigenvalues (blue dots), whose slope is an estimate of σ^2 . The middle panel plots $\hat{\lambda}_k$ versus k , which is the scree plot. The red solid curve is $\hat{\sigma}^2 q_k$ versus k . It fits the bulk eigenvalues (blue dots) very well. When this curve is extended to the left boundary, it hits $\hat{\sigma}^2(1 + \sqrt{\gamma_n})^2$. Our threshold for the top eigenvalues, which is the $(1 - \beta)$ -quantile of the Tracy-Widom distribution, is slightly larger than this value and shown by the dotted red line. The right panel zooms into the gray square area of the middle panel. It shows that 10 empirical eigenvalues exceeds the threshold, resulting in $\hat{K} = 10$.

middle panel. As k gets smaller (e.g., $k < 50$), the eigenvalues stay above the fitted MP quantile curve. This is because these $\hat{\lambda}_k$ are influenced by the spiked eigenvalues of Σ . Such eigenvalues are already excluded in the estimation of σ^2 . The right panel can also be viewed as a scree plot. Finding the elbow point of the scree plot is a common ad-hoc method for estimating K . In this plot, the elbow points are $\{6, 7, 10, 11\}$, hard to decide the true K . In contrast, our method correctly picks $\hat{K} = 10$.

Remark (Connection to the robust estimation of σ^2). As a byproduct, the BEMA algorithm yields a new estimator for σ^2 in the standard spiked covariance model, which can be useful for other problems such as reconstruction of spiked covariance matrices. Gavish and Donoho (2014) proposed a robust estimator of σ^2 , which is the ratio between the median of eigenvalues and the median of a standard MP distribution. Viewed in the Q-Q plot (left panel of Figure 3), their method is equivalent to using a single point to decide the slope. In comparison, our method uses a number of bulk eigenvalues to decide the slope and is thus more robust. Kritchman and Nadler (2009) proposed an estimator of σ^2 by solving a non-linear system of equations, and Shabalin and Nobel (2013) estimated σ^2 by minimizing the Kolmogorov–Smirnov distance between the ESD and its theoretical limit. In comparison, our estimator of σ^2 is from a simple least square and is much easier to compute. In Section 4, we also give an explicit error bound for our estimator.

3. BEMA for the General Spiked Covariance Model

We now consider the general case where the residual covariance matrix can have unequal diagonal entries. We shall modify Algorithm 1 to accommodate this setting. Rewrite Models (1) and (2) as

$$\Sigma = \sum_{k=1}^K \mu_k \xi_k \xi_k^\top + \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_p^2),$$

where $\sigma_k^2 \stackrel{\text{iid}}{\sim} \text{Gamma}(\theta, \theta/\sigma^2)$. (7)

Same as before, let $\hat{\lambda}_1 > \hat{\lambda}_2 > \dots > \hat{\lambda}_{n \wedge p}$ denote the nonzero eigenvalues of the sample covariance matrix. Below, in Section 3.1, we first state some well-known results from random matrix theory and motivate our methodology idea. In Section 3.2, we formally introduce the BEMA algorithm. In Section 3.3, we provide an asymptotic confidence interval for K .

3.1. The Asymptotic Behavior of Empirical Eigenvalues

Under Model (7), the asymptotic behavior of bulk eigenvalues and top eigenvalues exhibit some similarity to the case of standard spiked covariance model:

- The ESD converges to a fixed limit.
- The first K empirical eigenvalues stand out of the bulk.

However, the precise statement is more sophisticated.

We first consider the ESD. When K is finite and $p/n \rightarrow \gamma$, the ESD converges to a distribution $F_\gamma(x; \sigma^2, \theta)$. This distribution is parametrized by (σ^2, θ) , but it does not have an explicit form. It is defined implicitly by an equation of its Stieltjes

transform (Marcenko and Pastur 1967). Let $H_{\sigma^2, \theta}(t)$ be the CDF of $\text{Gamma}(\theta, \theta/\sigma^2)$. For each $z \in \mathbb{C}^+$, there is a unique $m = m(z; \gamma, \sigma^2, \theta) \in \mathbb{C}^+$ such that

$$z = -\frac{1}{m} + \gamma \int \frac{t}{1 + tm} dH_{\sigma^2, \theta}(t). \quad (8)$$

The density of $F_\gamma(x; \sigma^2, \theta)$, denoted by $f_\gamma(x; \sigma^2, \theta)$, satisfies that

$$f_\gamma(x; \sigma^2, \theta) = \lim_{y \rightarrow 0+} \left\{ \frac{1}{\pi(\gamma \wedge 1)} \Im(m(x + iy; \gamma, \sigma^2, \theta)) \right\},^1 \quad (9)$$

where $\Im(\cdot)$ denotes the imaginary part of a complex number.

We aim to estimate (σ^2, θ) by comparing the bulk eigenvalues with the corresponding quantiles of $F_\gamma(x; \sigma^2, \theta)$. In the special case of $\theta = \infty$, $F_\gamma(x; \sigma^2, \theta)$ reduces to the MP distribution. Therefore, we can compute its quantiles explicitly and estimate σ^2 by a simple least square. For the general case, we have to compute the quantiles of $F_\gamma(x; \sigma^2, \theta)$ numerically. There are two approaches, one is solving the density from Equations (8) and (9), and then computing the quantiles, and the other is using Monte Carlo simulations. We will describe them in Section 3.2.

Next, we consider the top eigenvalues. It requires a precise definition of “standing out” of the bulk. We use the distribution of $\hat{\lambda}_{K+1}$ under Model (7) as a benchmark, that is, $\hat{\lambda}_k$ needs to be much larger than a high-probability upper bound of $\hat{\lambda}_{K+1}$ in order to be called “standing out.” Fortunately, the behavior of $\hat{\lambda}_{K+1}$ has been studied in the literature of random matrix theory. We define the following null model, which is a special case of Model (7) with $K = 0$

$$\Sigma = \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_p^2),$$

where $\sigma_k^2 \stackrel{\text{iid}}{\sim} \text{Gamma}(\theta, \theta/\sigma^2)$. (10)

Let $\hat{\lambda}_1^*$ denote the largest eigenvalue of the sample covariance matrix under this null model. By eigenvalue sticking result (see Bloemendal et al. 2016, Knowles and Yin 2017 and a detailed discussion in Section 4.3), the distribution of $\hat{\lambda}_{K+1}$ is asymptotically close to the distribution of $\hat{\lambda}_1^*$. We now reframe the statement that “the first K empirical eigenvalues stand out” as follows: Under some regularity conditions, each of $\hat{\lambda}_1, \dots, \hat{\lambda}_K$ is significantly larger than $\hat{\lambda}_1^*$ associated with Model (10).

We aim to threshold the top eigenvalues by the $(1 - \beta)$ -quantile of the distribution of $\hat{\lambda}_1^*$, where β controls the probability of over-estimating K . In the special case of $\theta = \infty$, the distribution of $\hat{\lambda}_1^*$ converges to a Tracy–Widom distribution, so that we have a closed-form expression for the threshold. In the general case, we calculate this threshold by Monte Carlo simulation, where we simulate data from the null model to approximate the distribution of $\hat{\lambda}_1^*$. We relegate the details to Section 3.2.

3.2. The Algorithm of Estimating K

Same as before, the BEMA algorithm has two steps: Step 1 estimates (σ^2, θ) from bulk eigenvalues, and Step 2 calculates a threshold for the top eigenvalues.

¹The factor $1/(\gamma \wedge 1)$ is due to considering the zero-excluded ESD. If we consider the classical ESD, this factor should be $1/\gamma$.

Consider Step 1. Write $\tilde{p} = p \wedge n$ and $\gamma_n = p/n$. For a constant $\alpha \in (0, 1/2)$, we take the $(1 - 2\alpha)$ -fraction of bulk eigenvalues in the middle range, that is, $\{\hat{\lambda}_k : \alpha\tilde{p} \leq k \leq (1 - \alpha)\tilde{p}\}$. Each empirical eigenvalue $\hat{\lambda}_k$ is also the $(k/\tilde{p})$ -upper-quantile of the ESD. We recall that $F_{\gamma_n}(x; \sigma^2, \theta)$ is the theoretical limit of ESD as defined in Equations (8) and (9). Let $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \sigma^2, \theta)$ denote the $(k/\tilde{p})$ -upper-quantile of this distribution. We expect to see

$$\hat{\lambda}_k \approx \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \sigma^2, \theta).$$

It motivates the following estimator of (σ^2, θ) :

$$(\hat{\sigma}^2, \hat{\theta}) = \operatorname{argmin}_{(\sigma^2, \theta)} \left\{ \sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} [\hat{\lambda}_k - \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \sigma^2, \theta)]^2 \right\}. \quad (11)$$

We now describe how to solve Equation (11). This is a two-dimensional optimization. As long as we can evaluate the objective function for arbitrary (σ^2, θ) , we can solve it via a simple grid search. To further simplify the objective, we first get rid of σ^2 and reduce it to an optimization on θ only. Note that $\text{Gamma}(\theta, \theta/\sigma^2)$ is equivalent to $\sigma^2 \cdot \text{Gamma}(\theta, \theta)$. We can deduce from (8)-(9) that a similar connection holds between $F_{\gamma_n}(x; \sigma^2, \theta)$ and $F_{\gamma_n}(x; 1, \theta)$. Then, their quantiles satisfy

$$\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \sigma^2, \theta) = \sigma^2 \cdot \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \theta).$$

We rewrite Equation (11) as

$$\min_{\theta} H(\theta), \quad \text{where} \quad H(\theta) \equiv \min_{\sigma^2} \left\{ \sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} [\hat{\lambda}_k - \sigma^2 \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \theta)]^2 \right\}.$$

As long as we can compute $\bar{F}_{\gamma_n}^{-1}(y; 1, \theta)$ for any $\theta > 0$ and $y \in [0, 1]$, we can obtain $H(\theta)$ for each θ by least-square regression of the $\hat{\lambda}_k$'s on the $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \theta)$'s. Given $H(\theta)$, we can solve the optimization by a grid search on θ .

This is described in Step 1 of Algorithm 2. Suppose there is an available algorithm `GetQT` that computes $\bar{F}_{\gamma_n}^{-1}(y; 1, \theta)$ for any $\theta > 0$ and $y \in [0, 1]$. Fix a set of grid points $\{\theta_j\}_{j=1}^N$. For each θ_j , we first compute $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \theta_j)$ for all $\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}$. Given θ_j , the value of σ^2 that minimizes (11) is obtained by regressing $\{\hat{\lambda}_k\}_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}}$ on $\{\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \theta_j)\}_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}}$ with a least square. Let $\hat{\sigma}^2(\theta_j)$ denote this optimal value of σ^2 , and let v_j denote the objective in Equation (11) associated with $\{\theta_j, \hat{\sigma}^2(\theta_j)\}$. We select j^* so that v_j is minimized and set $\hat{\theta} = \theta_{j^*}$ and $\hat{\sigma}^2 = \hat{\sigma}^2(\theta_{j^*})$.

What remains is the design of an algorithm `GetQT`(y, γ_n, θ) to compute the y -upper-quantile of the distribution $F_{\gamma_n}(\cdot; 1, \theta)$ for arbitrary (θ, y) . We note that $F_{\gamma_n}(x; 1, \theta)$ only has an implicit definition through Equations (8) and (9). In the appendix, we propose two algorithms that serve for this purpose

- `GetQT1` takes advantage of the fact that $F_{\gamma_n}(x; \sigma^2, \theta)$ is also the theoretical limit of the ESD of the null model (10). This algorithm simulates data from Model (10) with $\sigma^2 = 1$ to get the Monte Carlo approximation of the target quantile.

Algorithm 2. BEMA for the general spiked covariance model.

Input: Nonzero eigenvalues $\hat{\lambda}_1, \dots, \hat{\lambda}_{n \wedge p}$, $\alpha \in (0, 1/2)$, $\beta \in (0, 1)$, a grid of values $0 < \theta_1 < \theta_2 < \dots < \theta_N$, an algorithm `GetQT`, and an integer $M \geq 1$.

Output: An estimate of K .

Step 1: Write $\tilde{p} = n \wedge p$ and $\gamma_n = p/n$. For each $1 \leq j \leq N$:

- For each $\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}$, run the algorithm `GetQT`($k/\tilde{p}, \gamma_n, \theta_j$) to obtain q_{kj} .
- Compute $\hat{\sigma}^2(\theta_j) = (\sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} q_{kj} \hat{\lambda}_k) / (\sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} q_{kj}^2)$.
- Let $v_j = \sum_{\alpha\tilde{p} \leq k \leq (1-\alpha)\tilde{p}} [\hat{\lambda}_k - \hat{\sigma}^2(\theta_j) \cdot q_{kj}]^2$.

Find $j^* = \operatorname{argmin}_{1 \leq j \leq N} v_j$. Let $\hat{\theta} = \theta_{j^*}$ and $\hat{\sigma}^2 = \hat{\sigma}^2(\theta_{j^*})$.

Step 2: For $1 \leq m \leq M$:

- Sample $d_j^* \sim \text{Gamma}(\hat{\theta}, \hat{\theta})$, independently for $1 \leq j \leq p$. Sample $X_i^*(j) \sim N(0, \hat{\sigma}^2 d_j^*)$, independently for $1 \leq i \leq n$ and $1 \leq j \leq p$.
- Compute the largest singular value of $n^{-1/2} X^*$, where $X^* = [X_1^*, X_2^*, \dots, X_n^*]^\top$. Let $\hat{\lambda}_{1(m)}^*$ be the square of this singular value.

Let $\hat{T}$ be the $(1 - \beta)$ -quantile of $\{\hat{\lambda}_{1(m)}^*\}_{1 \leq m \leq M}$. Output $\hat{K} = \#\{1 \leq k \leq \tilde{p} : \hat{\lambda}_k > \hat{T}\}$.

- `GetQT2` directly uses Equations (8) and (9) to solve the density $f_{\gamma_n}(x; 1, \theta)$: For any given θ , $H_{1,\theta}(t)$ is the CDF of $\text{Gamma}(\theta, \theta)$, which is known. Therefore, we can solve $f_{\gamma_n}(x; 1, \theta)$ from Equations (8) to (9) for a grid of x ; next, the values of $f_{\gamma_n}(x; 1, \theta)$ for arbitrary $x \in \mathbb{R}$ are obtained by linear extrapolation. Once we have the density function, we can compute the target quantiles.

The two `GetQT` algorithms have comparable numerical performance, but each has an advantage on running time in some cases; see the appendix for more discussions.

Consider Step 2. We estimate K by comparing each top eigenvalue with the $(1 - \beta)$ -quantile of the distribution of $\hat{\lambda}_1^*$ under the null model (10), with $(\hat{\sigma}^2, \hat{\theta})$ plugged in. The threshold is

$$\hat{T} = \begin{cases} (1 - \beta)\text{-quantile of the distribution of } \hat{\lambda}_1^* \text{ under the null model} \\ \Sigma = \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_p^2), \text{ where } \sigma_j^2 \stackrel{\text{iid}}{\sim} \text{Gamma}(\hat{\theta}, \hat{\theta}/\hat{\sigma}^2) \end{cases}. \quad (12)$$

The $\hat{T}$ here generalizes the threshold in Algorithm 1. The threshold in Algorithm 1 is a special case of $\hat{T}$ at $\hat{\theta} = \infty$, which happens to have an explicit formula.

We compute $\hat{T}$ via Monte Carlo simulations. We first draw Σ from the null model in Equation (12), and then draw the data matrix from multivariate normal distributions and compute the largest eigenvalue of the sample covariance matrix. By

repeating these steps multiple times, we obtain the sampling distribution of $\hat{\lambda}_1^*$ in Equation (12). This is described in Step 2 of Algorithm 2.

The BEMA algorithm has three tuning parameters (α, β, M) , where α controls the percentage of bulk eigenvalues used for estimating (σ^2, θ) and M is the number of Monte Carlo repetitions for approximating $\hat{T}$. The performance of the algorithm is relatively insensitive to (α, M) (see Section 5). We set $\alpha = 0.2$ and $M = 500$ by default. The parameter β controls the probability of over-estimating K . Theoretically, if the spiked eigenvalues are large enough, we should use a diminishing β (i.e., $\beta \rightarrow 0$ as $n \rightarrow \infty$) so that the probability of over-estimating K tends to zero. In practice, it often happens that the spiked eigenvalues are only moderately large. We thus need a moderate β to strike a balance between the probability of over-estimating K and the probability of under-estimating K . We leave it to the users to decide. It is analogous to the situation in false discovery rate control, where the users select the target false discovery rate. In our numerical experiments, we find that $\beta = 0.1$ is a good choice.

A simulation example. We illustrate Algorithm 2 using a simulation example. Fix $(n, p, K) = (1000, 200, 5)$ and $(\sigma^2, \theta) = (1, 10)$. We generate X_i iid from $N(0, \Sigma)$, where Σ satisfies model (7) with $\mu_k = 2.3$ for $1 \leq k \leq K$. The left panel of Figure 4 shows the plot of $\hat{\lambda}_k$ vs. the MP quantiles q_k . It does not fit a line crossing the origin, suggesting that Algorithm 1 does not work for this general covariance model. The middle panel contains the plot of $\hat{\lambda}_k$ vs. $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \hat{\theta})$, where $\hat{\theta}$ is from Algorithm 2. Except for a few top eigenvalues, it fits very well a line crossing the origin, suggesting that Algorithm 2 is successful in this setting. The estimated parameters are $(\hat{\sigma}^2, \hat{\theta}) = (1.02, 10.39)$, which is close to the true values. The right panel contains the plot of $\hat{\lambda}_k$ vs. k , and the fitted curve of $\hat{\sigma}^2 \cdot \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \hat{\theta})$ vs. k (solid red line). The threshold $\hat{T}$ is also shown by the dashed line. It yields $\hat{K} = 5$, which is the same as the ground truth.

Remark (Connection to parallel analysis). Parallel analysis (Horn 1965) is a popular method for estimating the number of spiked eigenvalues in real applications. It samples data from

a *null covariance model* that has no spiked eigenvalues, and estimates K by comparing the distribution of top empirical eigenvalues on simulated data to those actually observed from the original data. The most common version of parallel analysis first standardizes the data matrix so that each variable has a unit variance and then uses $\Sigma = I_p$ as the null model. Our algorithm has a similar spirit as parallel analysis, but we adopt a more sophisticated null covariance model, Model (10), and estimate parameters of this null model carefully from bulk eigenvalues.

Remark (Memory use of BEMA). The input of BEMA includes nonzero eigenvalues of the sample covariance matrix. These eigenvalues can be computed by eigen-decomposition on either the $p \times p$ matrix $X^T X$ or the $n \times n$ matrix XX^T . Therefore, the memory use depends on the minimum of n and p . In many real applications, p is very large but n is relatively small, and BEMA is still implementable under even strict memory constraints.

3.3. A Confidence Interval of K

By varying β in Algorithm 2, we get different estimators of K , where the over-shooting probability is controlled at different levels. We use these estimators to construct a confidence interval for K .

Definition 2 (Confidence interval of K). Denote the output of Algorithm 2 by $\hat{K}_\beta$ to indicate its dependence on β . Given any $\omega_0 \in (0, 1)$, we introduce the following $(1 - \omega_0)$ -confidence interval of K as $[\hat{K}_{\omega_0/2}, \hat{K}_{1-\omega_0/2}]$.

We explain why the confidence interval is asymptotically valid. Let $\hat{T} = \hat{T}_\beta$ be the threshold in Equation (12), and let $\hat{\lambda}_1^*$ be the largest eigenvalue of the sample covariance matrix when data are from the null model (10). We use $\mathbb{P}_0$ to denote the probability measure associated with Model (10). By definition of $\hat{T}_\beta$, $\mathbb{P}_0\{\hat{\lambda}_1^* \leq t\} \Big|_{t=\hat{T}_\beta} = 1 - \beta$. At the same time, the eigenvalue sticking result (Bloemendal et al. 2016; Knowles and Yin 2017) states that, under some regularity conditions, the distribution of

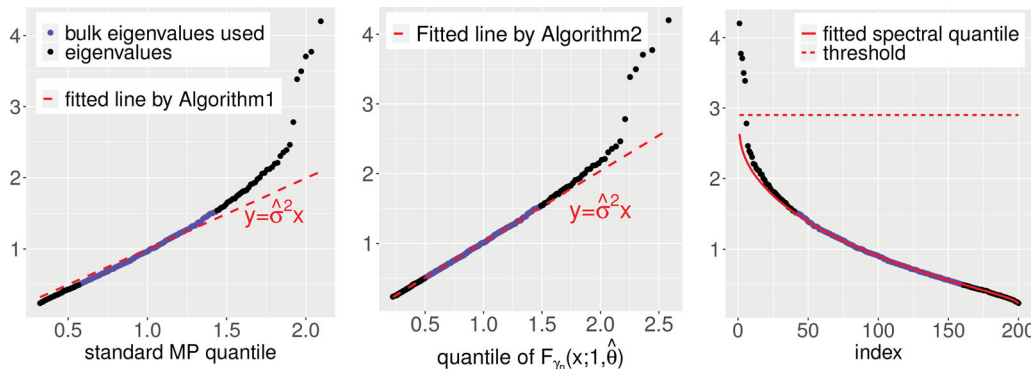


Figure 4. Illustration of BEMA for the general spiked covariance model. The left panel plots $\hat{\lambda}_k$ vs. q_k , where the q_k 's are quantiles of the standard MP distribution. It fits the regression line poorly, suggesting that Algorithm 1 is no longer working for this general model. The middle panel plots $\hat{\lambda}_k$ vs. $\bar{F}_{\gamma_n}^{-1}(x; 1, \hat{\theta})$, where $\hat{\theta}$ is an estimate of θ by Algorithm 2. The bulk eigenvalues (blue dots) fit the regression line very well. The right panel is the scree plot, where the red solid curve is $\bar{F}_{\gamma_n}^{-1}(x; \hat{\sigma}^2, \hat{\theta})$ vs. k . A threshold (red dotted line) is given by the 90%-quantile of the distribution of $\hat{\lambda}_1^*$ from a null model; see Equation (12). There are 5 empirical eigenvalues exceeding this threshold, which gives $\hat{K} = 5$.

$\hat{\lambda}_{K+1}$ is asymptotically close to the distribution $\hat{\lambda}_1^*$. It follows that

$$\begin{aligned} \mathbb{P}\{\hat{K}_{\omega_0/2} > K\} &\leq \mathbb{P}\{\hat{\lambda}_{K+1} > \hat{T}_{\omega_0/2}\} \approx \mathbb{P}_0\{\hat{\lambda}_1^* > t\} \Big|_{t=\hat{T}_{\omega_0/2}} \\ &= \omega_0/2, \\ \mathbb{P}\{\hat{K}_{1-\omega_0/2} < K\} &\leq \mathbb{P}\{\hat{\lambda}_K \leq \hat{T}_{1-\omega_0/2}\} \leq \mathbb{P}\{\hat{\lambda}_{K+1} \leq \hat{T}_{1-\omega_0/2}\} \\ &\approx \mathbb{P}_0\{\hat{\lambda}_1^* \leq t\} \Big|_{t=\hat{T}_{1-\omega_0/2}} = \omega_0/2. \end{aligned}$$

4. Theoretical Properties

We study in this section the theoretical properties of the proposed BEMA method. In Section 4.1, we focus on the standard spiked covariance model ($\theta = \infty$), where we derive the error rate of $\hat{\sigma}^2$ and the consistency of $\hat{K}$. In Section 4.2, we study the general spiked covariance model ($\theta < \infty$). This setting is much more complicated. It connects to an unsolved problem in random matrix theory, that is, how to get sharp asymptotic theory for eigenvalues when the limiting spectrum of Σ is unbounded and has convex decay in the tail. Only partial results are known (Kwak, Lee, and Park 2019). To overcome the technical difficulty, in our theoretical investigation, we approximate Model (7) by a proxy model where σ_j^2 are iid generated from a truncated Gamma distribution. Under this proxy model, we derive the rate of convergence for $(\hat{\sigma}^2, \hat{\theta})$ and the consistency of $\hat{K}$. In Section 4.3, we connect Model (7) to the proxy model and discuss the theory for Model (7).

Through this section, we assume $X_1, X_2, \dots, X_n$ are generated as follows:

Assumption 1. Let $Y = [Y_1, Y_2, \dots, Y_n]^\top \in \mathbb{R}^{n \times p}$ be a random matrix with independent but not necessarily identically distributed entries, where $\mathbb{E}[Y_i(j)] = 0$ and $\text{Var}(Y_i(j)) = 1$, for $1 \leq i \leq n$, $1 \leq j \leq p$. Given $\sigma_1, \sigma_2, \dots, \sigma_p > 0$, $\mu_1 \geq \mu_2 \geq \dots \geq \mu_K > 0$, and orthonormal vectors $\xi_1, \xi_2, \dots, \xi_p \in \mathbb{R}^p$, let $\Sigma = \sum_{k=1}^K (\sigma_k^2 + \mu_k) \xi_k \xi_k^\top + \sum_{j=K+1}^p \sigma_j^2 \xi_j \xi_j^\top$. We assume $X_i = \Sigma^{1/2} Y_i$, for $1 \leq i \leq n$.

Under this assumption, each X_i is a linear transform of a random vector Y_i that has independent entries. This is stronger than assuming $\text{Cov}(X_i) = \Sigma$ but is conventional in the literature.

Assumption 2. For each integer $m \geq 1$, there exists a universal constant $C_m > 0$ such that $\sup_{1 \leq i \leq n, 1 \leq j \leq p} \mathbb{E}[|Y_i(j)|^m] \leq C_m$.

This assumption can be further relaxed. For example, we do not actually need the inequality to hold for every $m \geq 1$ but only for $1 \leq m \leq M$, where M is a properly large integer (Bloemendal et al. 2016; Knowles and Yin 2017). We use the current assumption for convenience.

We will use the following notation frequently, which is conventional in random matrix theory:

Definition 3. Let U_n and V_n be two sequences of random variables indexed by n . We say that U_n is stochastically dominated by V_n , if for any $\epsilon > 0$ and $s > 0$ there exists $N = N(\epsilon, s)$ such that $\mathbb{P}(U_n > n^\epsilon V_n) \leq n^{-s}$ for all $n \geq N$. We write $U_n \prec V_n$.

4.1. The Standard Spiked Covariance Model

The standard spiked covariance model (Johnstone 2001) assumes $D = \sigma^2 I_p$. In this case, BEMA simplifies to Algorithm 1. It outputs $\hat{\sigma}^2$ and $\hat{K}$. We first give an error bound on estimating σ^2 .

Theorem 1 (Estimation error of $\hat{\sigma}^2$). Suppose $X_1, X_2, \dots, X_n$ satisfy Assumptions 1–2 with $\sigma_j^2 \equiv \sigma^2$. Suppose $K \geq 1$ is fixed and $p/n \rightarrow \gamma$ for a constant $\gamma > 0$. Let $\hat{\sigma}^2$ be the estimator of σ^2 by Algorithm 1, where the tuning parameter α is a constant in $(0, 1/2)$. Then, $|\hat{\sigma}^2 - \sigma^2| \prec n^{-1}$.

This result is connected to the robust estimation of σ^2 in a standard spiked covariance model (Gavish and Donoho 2014; Kritchman and Nadler 2009; Shabalin and Nobel 2013). In these work, there are only consistency results available (Donoho, Gavish, and Johnstone 2018) which say that $\hat{\sigma}^2 \rightarrow \sigma^2$ almost surely, but there are no explicit error rates. Using the recent advancement in random matrix theory on sharp large-deviation bounds for individual empirical eigenvalues (see Ke (2016) for a survey), we can leverage those results to obtain an explicit bound for $|\hat{\sigma}^2 - \sigma^2|$.

We then establish the consistency on estimating K .

Theorem 2 (Consistency of $\hat{K}$). Suppose $X_1, X_2, \dots, X_n$ satisfy Assumptions 1 and 2 with $\sigma_j^2 \equiv \sigma^2$. Suppose $K \geq 1$ is fixed, $p/n \rightarrow \gamma \in (0, \infty)$, and $\mu_K \geq \sigma^2(\sqrt{\gamma} + \tau_n)$, where $\tau_n \gg n^{-1/3}$. Let $\hat{K}$ be the estimator of K by Algorithm 1, where the tuning parameters are such that $\alpha \in (0, 1/2)$ is a constant and that $\beta \rightarrow 0$ at a properly slow rate. As $n \rightarrow \infty$, $\mathbb{P}\{\hat{K} = K\} = 1 - o(1)$.

We compare the conditions required for consistent estimation of K with those in other work. Let $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p$ denote the eigenvalues of Σ . In our model, $\lambda_k = \mu_k + \sigma^2$ for $1 \leq k \leq K$. The condition in Theorem 2 translates to

$$\lambda_K > \sigma^2(1 + \sqrt{\gamma} + \tau_n), \quad \tau_n \gg n^{-1/3}.$$

It is weaker than the conditions in Bai and Ng (2002) and Cai, Han, and Pan (2020), where the former requires $\lambda_K \asymp p$ and the latter needs $\lambda_K \rightarrow \infty$. Our condition on λ_K matches with the critical phase transition threshold in Baik, Ben Arous, and Peché (2005) and is hardly improvable. In fact, Fan, Guo, and Zheng (2020) showed that if $\lambda_K \leq \sigma^2(1 + \sqrt{\gamma})$ then there exists no consistent estimator of K . Dobriban and Owen (2019) imposed the same condition on λ_K , but they need stronger conditions on population eigenvectors. Their “delocalization” condition states as $\|\Xi \Lambda^{1/2}\|_\infty \rightarrow 0$, where $\Xi = [\xi_1, \dots, \xi_K]$, $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_K)$, and $\|\cdot\|_\infty$ is the maximum absolute row sum. It requires the eigenvectors to be incoherent (i.e., $\max_{1 \leq k \leq K} \|\xi_k\|_\infty$ is sufficiently small) and that the eigenvalues cannot be too large. Examples such as equal-correlation matrices (i.e., $\Sigma(i, j) = a$, for all $i \neq j$, where $a \in (0, 1)$ is a constant) are excluded. We do not need such a de-localization condition.²

²We remark that the comparison is for the standard spiked covariance model only. For this model, our method has the weakest conditions for consistent estimation of K . On the other hand, other methods apply to some other settings, which are not considered in the comparison.

The proof of [Theorem 2](#) is an application of the *eigenvalue sticking* theory (Bloemendal et al. 2016). It compares the distribution of empirical eigenvalues $\{\hat{\lambda}_k\}$ under the spiked covariance model with the distribution of empirical eigenvalues $\{\hat{\lambda}_k^*\}$ under the null model $\Sigma = \sigma^2 I_p$. The claim is that the distribution of $\hat{\lambda}_{K+s}$ is asymptotically close to the distribution of $\hat{\lambda}_s^*$, for a wide range of s . We use this result to study the thresholding step in Algorithm 1.

4.2. The Truncated Gamma-Based General Spiked Covariance Model

The general spiked covariance model (7) assumes σ_j^2 are iid drawn from $\text{Gamma}(\theta, \theta/\sigma^2)$. It differs from the conventional settings in random matrix theory because Σ is not a deterministic matrix and because the limiting spectral density of Σ does not have a compact support. Unfortunately, there is no existed random matrix theory that deals with this setting directly (Bao 2020). We thus approximate Model (2) by

$$\sigma_j^2 \stackrel{\text{iid}}{\sim} \text{TruncGamma}(\theta, \theta/\sigma^2, \sigma^2 T_1, \sigma^2 T_2), \quad 1 \leq j \leq p, \quad (13)$$

where $\text{TruncGamma}(\alpha, \beta, l, u)$ denotes the truncated Gamma distribution with rate and shape parameters α and β and truncations at l and u . When $(T_1, T_2) = (0, \infty)$, it reduces to Model (2). Given fixed $0 < T_1 < T_2 < \infty$, the limiting spectral density of Σ has a compact support, so that we can take advantage of the existing random matrix theory (Knowles and Yin 2017; Ding 2020). We first present the theory for Model (13) and then discuss how to extend it to $(T_1, T_2) = (0, \infty)$.

Fixing $0 < T_1 < T_2 < \infty$ and two intervals $\mathcal{J}_{\sigma^2} = [a, b] \subset (0, \infty)$ and $\mathcal{J}_\theta = [c, d] \subset (0, \infty)$, let $\mathcal{Q}(T_1, T_2, \mathcal{J}_{\sigma^2}, \mathcal{J}_\theta)$ be the family of distributions $\text{TruncGamma}(\theta, \theta/\sigma^2, \sigma^2 T_1, \sigma^2 T_2)$ satisfying that $\sigma^2 \in \mathcal{J}_{\sigma^2}$ and $\theta \in \mathcal{J}_\theta$. The following Lemma is a result of (Knowles and Yin 2017, theo. 3.12 and exam. 2.9), and its proof is omitted.

Lemma 1. Suppose $X_1, X_2, \dots, X_n$ satisfy [Assumptions 1](#) and [2](#) with σ_j^2 generated from Model (13). Suppose $K \geq 1$ is fixed and $p/n \rightarrow \gamma$ for a constant $\gamma \neq 1$. Suppose the truncated Gamma distribution in Equation (13) is from the family $\mathcal{Q}(T_1, T_2, \mathcal{J}_{\sigma^2}, \mathcal{J}_\theta)$, for fixed $(T_1, T_2, \mathcal{J}_{\sigma^2}, \mathcal{J}_\theta)$. Let $H_{\sigma^2, \theta, T_1, T_2}(t)$ be the CDF of $\text{TruncGamma}(\theta, \theta/\sigma^2, \sigma^2 T_1, \sigma^2 T_2)$. Define a distribution $F_{\gamma_n}(\cdot; \sigma^2, \theta, T_1, T_2)$ in the same way as in Equations (8) and (9), with $H_{\sigma^2, \theta}(t)$ replaced by $H_{\sigma^2, \theta, T_1, T_2}(t)$ and γ replaced by $\gamma_n = p/n$. Let $q_i \equiv \bar{F}_{\gamma_n}^{-1}(i/\tilde{p}; \sigma^2, \theta, T_1, T_2)$ be the $(i/\tilde{p})$ -upper-quantile of this distribution, where $\tilde{p} = n \wedge p$. As $n \rightarrow \infty$, for every $K < i \leq \tilde{p}$, we have $|\hat{\lambda}_i - q_i| < [i \wedge (\tilde{p} + 1 - i)]^{-1/3} n^{-2/3}$.

Given (T_1, T_2) , we estimate σ^2 and θ by

$$(\hat{\sigma}^2, \hat{\theta}) = \underset{(\sigma^2, \theta) \in \mathcal{J}_{\sigma^2} \times \mathcal{J}_\theta}{\operatorname{argmin}} \left\{ \sum_{\alpha \tilde{p} \leq i \leq (1-\alpha)\tilde{p}} [\hat{\lambda}_i - \bar{F}_{\gamma_n}^{-1}(i/\tilde{p}; \sigma^2, \theta, T_1, T_2)]^2 \right\}. \quad (14)$$

It can be solved by a slight modification of Step 1 of Algorithm 2. We note that (13) is equivalent to $\sigma_j^2/\sigma^2 \stackrel{\text{iid}}{\sim} \text{TruncGamma}$

$(\theta, \theta, T_1, T_2)$. Hence, the quantiles satisfy that $\bar{F}_{\gamma_n}^{-1}(i/\tilde{p}; \sigma^2, \theta, T_1, T_2) = \sigma^2 \cdot \bar{F}_{\gamma_n}^{-1}(i/\tilde{p}; 1, \theta, T_1, T_2)$. We first modify `GetQT` so that it outputs the quantiles of $F_{\gamma_n}(\cdot; 1, \theta, T_1, T_2)$ for any given θ . Next, we mimic Step 1 of Algorithm 2 to solve Equation (14), where we run a least square for every θ and then optimize over θ via a grid search. The details are relegated to the appendix.

Theorem 3 (Estimation error of $\hat{\sigma}^2$ and $\hat{\theta}$). Suppose the conditions of [Lemma 1](#) hold, where $K, \gamma, T_1, T_2, \mathcal{J}_{\sigma^2}$, and $\mathcal{J}_\theta$ are fixed. Let

$$\begin{aligned} \Phi(\theta) &= \Phi(\theta; T_1, T_2) \\ &= \frac{[\int_{T_1}^{T_2} x^{\theta+1} \exp(-\theta x) dx][\int_{T_1}^{T_2} x^{\theta-1} \exp(-\theta x) dx]}{[\int_{T_1}^{T_2} x^\theta \exp(-\theta x) dx]^2}. \end{aligned}$$

Suppose there exists a positive constant $\omega = \omega(T_1, T_2, \mathcal{J}_\theta)$ such that $\sup_{\theta \in \mathcal{J}_\theta} \Phi'(\theta) \leq -\omega$. Let $\hat{\sigma}^2$ and $\hat{\theta}$ be the estimators from (14), where the tuning parameter α satisfies $\alpha \tilde{p} > K$ and $\alpha \tilde{p} = O(n/\log(n))$. As $n \rightarrow \infty$, we have $|\hat{\sigma}^2 - \sigma^2| < n^{-1}$ and $|\hat{\theta} - \theta| < n^{-1}$.

[Theorem 3](#) assumes $\sup_{\theta \in \mathcal{J}_\theta} \Phi'(\theta) \leq -\omega$ for some constant $\omega > 0$. It is a regularity condition on $(\mathcal{J}_\theta, T_1, T_2)$. The next lemma shows that this condition is mild.

Lemma 2. For any fixed $\mathcal{J}_\theta = [c, d]$ and $\omega < d^{-2}$, there exist constants $0 < T_1^* < T_2^* < \infty$ such that $\sup_{\theta \in \mathcal{J}_\theta} \Phi'(\theta; T_1, T_2) \leq -\omega$ holds for all $T_1 \leq T_1^*$ and $T_2 \geq T_2^*$.

With the estimates $\hat{\sigma}^2$ and $\hat{\theta}$, we then slightly modify Step 2 of Algorithm 2 by thresholding all the empirical eigenvalues at

$$\hat{T}_\beta = \begin{cases} (1-\beta)\text{-quantile of the distribution of } \hat{\lambda}_1^* \text{ under the} \\ \text{null model} \\ \Sigma = \text{diag}(\sigma_1^2, \dots, \sigma_p^2), \text{ where } \sigma_j^2 \stackrel{\text{iid}}{\sim} \text{TruncGamma} \\ (\hat{\theta}, \hat{\theta}/\hat{\sigma}^2, \hat{\sigma}^2 T_1, \hat{\sigma}^2 T_2) \end{cases}. \quad (15)$$

This threshold can be computed via Monte Carlo simulations, similarly as in Step 2 of Algorithm 2. We estimate K by the number of empirical eigenvalues exceeding $\hat{T}$.

To establish the consistency of $\hat{K}$, we introduce the function

$$G(x) = -\frac{1}{x} + \gamma \int \frac{1}{t^{-1} + x} dH_{\sigma^2, \theta, T_1, T_2}(t). \quad (16)$$

By (Knowles and Yin 2017, exam. 2.9), $G(x)$ has 2 critical points $0 > x_1^* > x_2^*$ (the definition of critical points can be found in Knowles and Yin (2017)), and the distribution $F_\gamma(\cdot; \sigma^2, \theta, T_1, T_2)$ defined in [Lemma 1](#) has the support $[G(x_2^*), G(x_1^*)]$. The next theorem is proved in the appendix. It uses a result in Ding (2020) about the top empirical eigenvalues.

Theorem 4 (Consistency of $\hat{K}$). Suppose the conditions of [Lemma 1](#) and [Theorem 3](#) hold. Let x_1^* be the largest critical point of the function $G(x)$ in Equation (16). We assume $-1/(T_1 + \mu_K) \geq x_1^* + \tau$,³ where $\tau > 0$ is a constant and T_1 is a truncation

³In our model (see [Assumption 1](#)), the spiked eigenvalues of Σ are $\{\mu_k + \sigma_k\}_{1 \leq k \leq K}$. Therefore, $\mu_K + T_1$ is a lower bound of these spiked eigenvalues.

point in Equation (13). Let $\hat{K} = \#\{1 \leq i \leq (n \wedge p) : \hat{\lambda}_i > \hat{T}_\beta\}$, where $\hat{T}_\beta$ is as in Equation (15) with $\beta \rightarrow 0$ at a properly slow rate. As $n \rightarrow \infty$, $\mathbb{P}\{\hat{K} = K\} = 1 - o(1)$.

4.3. Remarks on Extension to the Gamma-Based General Spiked Covariance Matrix

We now discuss extension of the theoretical results to the Gamma-based general spiked covariance model (7), which is an extreme case of Model (13) at $T_1 = 0$ and $T_2 = \infty$. As mentioned earlier, this setting is unconventional because the eigenvalues of Σ are stochastic and the support of the limiting spectral density of Σ is unbounded.

First, we discuss the estimation of (σ^2, θ) . The accuracy of $(\hat{\sigma}^2, \hat{\theta})$ depends on whether we have similar large deviation bounds to those in Lemma 1. Our conjecture is that the stochasticity and unboundedness of the spectrum of Σ has a negligible effect on the eigenvalues deep into the bulk. To see why, we note that the classical result about weak convergence of ESD (Marcenko and Pastur 1967) does not need the limiting spectrum of Σ to have a compact support; therefore, the unboundedness is not an issue. The stochasticity is not an issue, either, because almost surely, the spectral distribution of Σ converges weakly to $\text{Gamma}(\theta, \theta/\sigma^2)$. We conclude that the weak convergence of ESD still holds. This further implies that the bulk eigenvalues still converge to the corresponding quantiles of the theoretical limit of ESD.

The open question is whether we have the rates of convergence as in Lemma 1. The stochasticity and unboundedness of the spectrum of Σ affect the rates of convergence of large eigenvalues. We thus do not expect Lemma 1 to hold for all i . Fortunately, the estimation of (σ^2, θ) in BEMA only involves bulk eigenvalues in the middle range, that is, $\alpha\tilde{p} \leq i \leq (1-\alpha)\tilde{p}$, where $\alpha \in (0, 1/2)$ is a constant. We conjecture that Lemma 1 continues to hold for these eigenvalues. If our conjecture is correct, then we can show similar results for $\hat{\sigma}^2$ and $\hat{\theta}$ as those in Theorem 3.

Next, we discuss the consistency of $\hat{K}$. The stochasticity and unboundedness of the spectrum of Σ together yields a significant change of the behavior of edge eigenvalues. This can be seen from a relevant setting in Kwak, Lee, and Park (2019)— Σ is a diagonal matrix whose diagonal entries are iid drawn from a density $\rho(t) \propto (1-t)^b f(t) \cdot 1\{l \leq t \leq 1\}$, where $b > 1$ and $l \in (0, 1)$ are constants and $f \in C^1([l, 1])$. This setting has no spike. They showed that the limiting distribution of the largest eigenvalue, $\hat{\lambda}_1^*$, is not a Tracy-Widom distribution; it is a Weibull distribution if $\gamma < \gamma_0$ and a Gaussian distribution if $\gamma > \gamma_0$, where γ_0 is a positive constant. Our model is even more complicated, where the Gamma density exhibits a similar convex decay on the right tail but has an unbounded support. We do not expect $\hat{\lambda}_1^*$ to follow a Tracy-Widom distribution any more.

However, this does not eliminate the consistency of $\hat{K}$. To prove consistency, we first need that the stochastic threshold (12) in BEMA well approximates the $(1-\alpha)$ -upper-quantile of $\hat{\lambda}_1^*$, where $\hat{\lambda}_1^*$ is the largest eigenvalue of the null model with no spike. This follows from the nature of Monte Carlo simulations, no matter whether $\hat{\lambda}_1^*$ converges to a Tracy-Widom distribution.

Furthermore, the implementation of Equation (12) does not need any knowledge of the limiting distribution of $\hat{\lambda}_1^*$.

To prove consistency, we also need to show that, under Model (7), when μ_K is appropriately large, (i) the distribution of $\hat{\lambda}_{K+1}$ is asymptotically close to the distribution of $\hat{\lambda}_1^*$ in the null model (this is the “eigenvalue sticking” argument), and (ii) each of $\hat{\lambda}_1, \hat{\lambda}_2, \dots, \hat{\lambda}_K$ is significantly larger than the $(1-\alpha)$ -upper-quantile of $\hat{\lambda}_1^*$. We conjecture that both (i) and (ii) are correct, provided that $\mu_K \gg \log(n)$. If our conjectures are correct, then we can obtain the consistency of $\hat{K}$ as in Theorem 4, under the slightly stronger condition that $\mu_K \gg \log(n)$.

The rigorous proofs of our conjectures require redevelopment of several fundamental results in random matrix theory for Model (7), such as the local law on bulk eigenvalues and the limiting behavior of edge eigenvalues (including the spiked and non-spiked ones). It is beyond the scope of this article, and we leave for future work.

5. Simulation Studies

We examine the performance of our methods in simulations. To differentiate between Algorithms 1 and 2, we call the former BEMA0 and the latter BEMA. BEMA0 is a simplified version of BEMA, specifically designed for the standard spiked covariance model. The tuning parameters are fixed as $(\alpha, \beta) = (0.2, 0.1)$ for BEMA0 and $(\alpha, \beta, M) = (0.2, 0.1, 500)$ for BEMA when not particularly specified.

In Section 4.2, we also introduced a modification of BEMA using the truncated Gamma-based spiked mode for technical needs in our theoretical studies. We showed that this algorithm has desirable theoretical properties. It however requires two additional tuning parameters (T_1, T_2) . Our simulation studies (not reported here) show that the performance of the modified BEMA is similar to that of BEMA, when T_1 is appropriately small and T_2 is appropriately large. For this reason, we use BEMA, instead of the modified BEMA, in the following simulation studies.

We compare our methods with a few methods in the literature, including the deterministic parallel analysis (DDPA) from Dobriban and Owen (2019), the empirical Kaiser’s criterion (EKC) from Braeken and Van Assen (2017), the information criteria IC_{p1} (Bai&Ng) from Bai and Ng (2002) and the eigen-gap detection (Pass&Yao) from Passemier and Yao (2014).

Simulation 1. This experiment is for the standard spiked covariance model, where we investigate the performance of BEMA0 and the confidence interval for K as described in Section 3.3. We generate data from $X_i \stackrel{\text{iid}}{\sim} N(0, \Sigma)$, $1 \leq i \leq n$, where Σ satisfies Model (3) with

$$\mu_1 = \mu_2 = \dots = \mu_K = \rho \cdot \sigma^2 \sqrt{p/n}, \quad \text{for some } \rho > 0.$$

The value of ρ controls the magnitude of spiked eigenvalues. $\rho \leq 1$ is the region where consistent estimation of K is impossible (Baik, Ben Arous, and Peche 2005; Fan, Guo, and Zheng 2020). We examine the performance of BEMA0 in the region of $\rho > 1$.

Fix $K = 5$ and $\sigma^2 = 1$. We consider three settings, where (n, p) are (10,000, 1000), (1500, 5000), and (1500, 1500),

respectively. They cover different cases of size relationship between p and n . The eigenvector matrix Ξ is drawn uniformly from the Stiefel manifold (which is the collection of all $p \times K$ matrices that have orthonormal columns). For each of the three settings, we vary the value of ρ and report the average of $\hat{K}$ and upper/lower boundary of a 95% confidence interval, based on 100 repetitions; the results are in the top three panels of Figure 5. We also report the probability of correctly estimating K (correct rate) and the coverage probability of the 95% confidence interval (coverage rate); see the bottom three panels of Figure 5.

It agrees with our theoretical understanding that $\rho = 1$ is the critical phase transition point. When ρ slightly departs from 1, the coverage rate starts to increase from 0% and quickly reaches the target of 95%. The increase of the correct rate is slightly slower, but it reaches 100% before $\rho = 1.5$, for all three settings. Our theory suggests that the correct rate is asymptotically 100% as long as $\rho > 1$, but in the finite-sample performance we need a larger ρ to attain a 100% correct rate. Furthermore, as ρ increases, the estimated $\hat{K}$ increases from 0 to 5, with a sharp change at around $\rho = 1$. The length of the 95% confidence interval initially decreases with ρ and then stays almost constant.

Simulation 2. In this simulation, we compare BEMA0 and BEMA with other methods. We consider both the standard spiked covariance model (3) and the general spiked covariance model (7). BEMA0 and BEMA are designed for these two settings, respectively. We note that BEMA can also be applied to Model (3), which simply ignores the prior knowledge of equal

diagonal in the residual covariance matrix. We thereby also include BEMA in the numerical comparison on the standard spiked covariance model.

Given $(n, p, K, \lambda, \theta)$, we generate data $\mathbf{X}_i \stackrel{\text{iid}}{\sim} N(0, \Sigma)$, $1 \leq i \leq n$, where Σ satisfies Model (7) with $\sigma^2 = 1$ and $\mu_k = \lambda$, for $1 \leq k \leq K$. The eigenvector matrix Ξ is drawn uniformly from the Stiefel manifold. We allow θ to take the value of ∞ ; when $\theta = \infty$, it indicates that Σ follows the standard spiked covariance model (3). We consider 8 different settings which cover a wide range of parameter values. The results are shown in Table 1, where the average $\hat{K}$ and the probability of correctly estimating K (correct rate) are reported based on 500 repetitions.

We have a few observations. First, in the standard spiked covariance model ($\theta = \infty$, top four rows of Table 1), BEMA0 has the best performance. Interestingly, BEMA has nearly comparable performance. The reason is that the algorithm will automatically output a very large $\hat{\theta}$, so that the estimator is similar to that of knowing $\theta = \infty$. This suggests that we do not have to choose between BEMA0 and BEMA in practice. We can always use BEMA, even when the data come from the standard spiked covariance model. On the other hand, BEMA0 is conceptually simpler and computationally much faster, hence, it is still the better choice if we are confident that the standard spiked covariance model holds.

Second, in the general spiked covariance model (bottom four rows of Table 1), BEMA outperforms DDPA, EKC and Pass&Yao in all settings, and outperforms Bai&Ng in two out of four

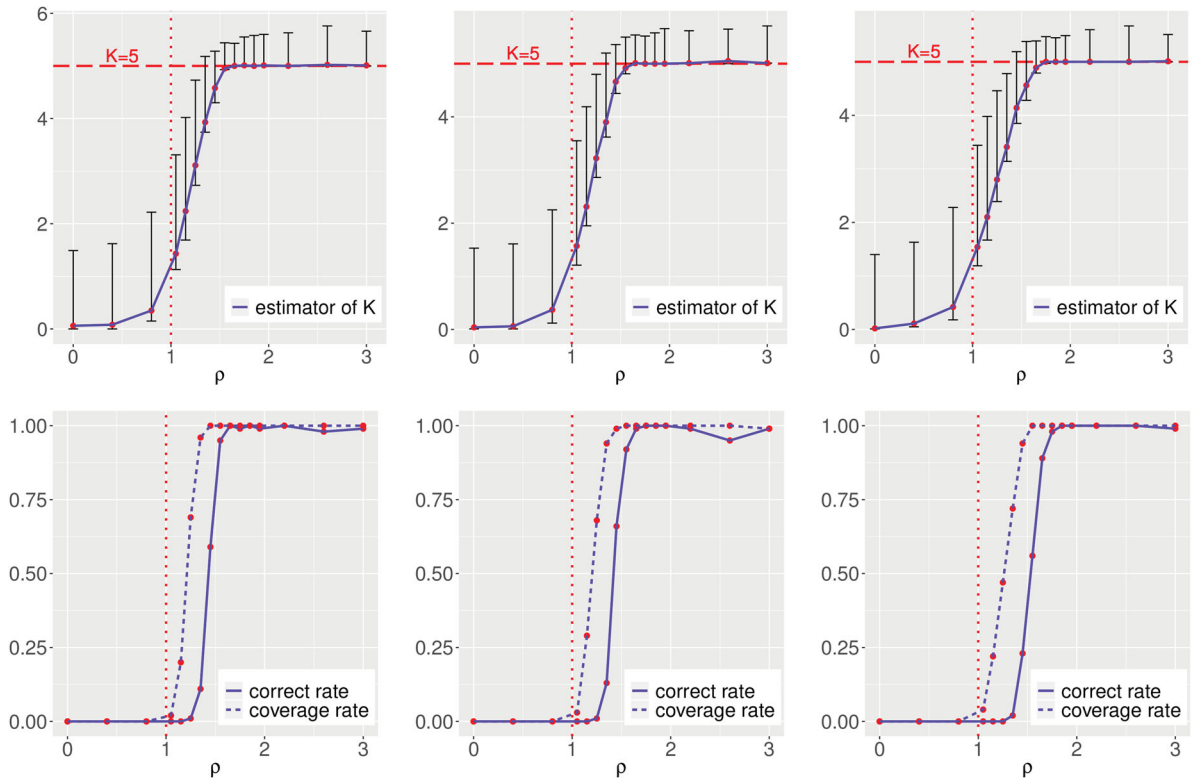


Figure 5. Simulation 1: The performance of BEMA0 in a standard spiked model. $K = 5$, and (n, p) take the value of $(10,000, 1000)$, $(1500, 5000)$, and $(1500, 1500)$ (from left to right). The top three panels show the estimator $\hat{K}$ along with the 95% confidence upper/lower bound, where each quantity is the average of 100 repetitions. The bottom three panels show the probability of correctly estimating K (correct rate) and the coverage probabilities of the 95% confidence intervals (coverage rates). In each panel, the x-axis is the value of ρ (see the text for definition), controlling the magnitude of spiked eigenvalues. Our theory states that BEMA0 gives a consistent estimator of K when ρ slightly exceeds 1. This is confirmed by these simulations.

Table 1. Simulation 2: Comparison of different methods in the standard/general spiked model.

$(n, p, K, \lambda, \theta)$	BEMAO	BEMA	DDPA	EKC	Bai&Ng	Pass&Yao
(100, 500, 5, 9, ∞)	4.996 (99.6%)	4.982 (98.2%)	6.102 (41%)	5.552 (57.8%)	0 (0%)	4.904 (92%)
(100, 500, 5, 49, ∞)	5 (100%)	5 (100%)	6.328 (38%)	6.4 (27.4%)	5 (100%)	5.012 (98.8%)
(500, 100, 5, 1.5, ∞)	5 (100%)	4.93 (93.0%)	6.1 (45.6%)	5.016 (98.4%)	0 (0%)	2.784 (43.8%)
(500, 100, 5, 3, ∞)	5 (100%)	5 (100%)	5.92 (45.4%)	5.056 (94.4%)	0 (0%)	4.432 (84.4%)
(100, 500, 5, 15, 3)	–	5.182 (85.2%)	9.222 (20.8%)	5.974 (40.2%)	0.078 (0%)	5.292 (73.2%)
(100, 500, 5, 50, 3)	–	5.142 (88.4%)	9.214 (20.8%)	9.852 (8.6%)	5 (100%)	5.362 (70.4%)
(500, 100, 5, 4.5, 3)	–	4.748 (81.2%)	57.954 (25.4%)	5.588 (49.0%)	3.392 (39%)	7.624 (5%)
(500, 100, 5, 6, 3)	–	5.018 (98.2%)	43.734 (38.8%)	6.244 (18.4%)	5.002 (99.8%)	8.098 (4.2%)

NOTES: In these settings, all the spiked eigenvalues are equal to λ , and the eigenvectors are randomly generated from the Stiefel manifold. The top four rows ($\theta = \infty$) correspond to the standard spiked model, and the bottom four rows correspond to the general spiked model. The number in each cell is the average $\hat{K}$ over 500 repetitions, and the number in brackets is the probability of correctly estimating K (correct rate).

Table 2. Simulation 3: Comparison of different methods in the standard/general spiked model, when the eigenvectors are “delocalized.”

(n, p, K, s_1, θ)	BEMAO	BEMA	DDPA	EKC	Bai&Ng	Pass&Yao
(100, 500, 1, 1, ∞)	0.988 (96%)	0.956 (95.2%)	1.086 (88.6%)	1.07 (88.8%)	0 (0%)	0.934 (91.8%)
(100, 500, 1, 3, ∞)	1.012 (98.8%)	1.008 (99.2%)	1.138 (87%)	1.146 (86.4%)	1 (100%)	1.036 (96.8%)
(500, 100, 1, 3, ∞)	1.020 (98%)	1 (100%)	1.152 (85.6%)	1.056 (94.4%)	0 (0%)	1.018 (98.2%)
(500, 100, 1, 6, ∞)	1.014 (98.6%)	1 (100%)	1.124 (88.6%)	1.12 (88%)	1 (100%)	1.014 (98.8%)
(100, 500, 1, 2, 10)	–	1.096 (90.6%)	1.2 (82.6%)	1.102 (90.4%)	0.388 (38.8%)	1.084 (92.6%)
(100, 500, 1, 6, 10)	–	1.104 (89.8%)	1.226 (79%)	1.608 (54.2%)	1 (100%)	1.054 (95%)
(500, 100, 1, 6, 3)	–	1.114 (89.2%)	1.062 (95.4%)	1.226 (78.2%)	1.008 (99.4%)	3.93 (6.2%)
(500, 100, 1, 12, 3)	–	1.124 (88.0%)	1.042 (97.4%)	3.782 (0.8%)	1.006 (99.4%)	3.672 (9.8%)

NOTES: Here, s_1 controls the magnitude of spiked eigenvalues, where $s_1^2(p/n)$ plays the role of λ in Simulation 2. The top four rows ($\theta = \infty$) correspond to the standard spiked model, and the bottom four rows correspond to the general spiked model. The number in each cell is the average $\hat{K}$, and the number in brackets is the probability of correctly estimating K (correct rate).

settings. BEMA is the only method whose correct rate is above 80% in *all* settings.

DDPA requires a delocalization condition. Let Ξ be the $p \times K$ matrix of eigenvectors, and let Λ be the diagonal matrix consisting of spiked eigenvalues. The delocalization condition is $\|\Xi \Lambda^{1/2}\|_\infty \rightarrow 0$. It prevents eigenvectors from having large entries. This condition is not satisfied here, explaining the unsatisfactory performance of DDPA. Bai&Ng requires that the spikes are sufficiently large. The larger p/n , the higher requirement of spikes. When $p/n = 5$ and $\lambda = 49$ or when $p/n = 0.2$ and $\lambda = 6$, Bai&Ng has a nearly 100% correct rate. However, as λ decreases, the correct rate drops very quickly. EKC uses a thresholding scheme that gives smaller thresholds to lower ranked eigenvalues (e.g., the threshold for $\hat{\lambda}_2$ is smaller than the threshold for $\hat{\lambda}_1$). This method often over-estimates K , especially when all the spikes are large (e.g., Row 6 of Table 1). Pass&Yao is developed for the standard spiked model. It has an unsatisfactory performance in the general spiked model (bottom four rows of Table 1).

Simulation 3. In this simulation, we change the generation process of eigenvectors to satisfy the “delocalization condition” (Dobriban and Owen 2019). This condition means $\|\Xi \Lambda^{1/2}\|_\infty$ is sufficiently small, where Ξ is the $p \times K$ matrix consisting of eigenvectors and Λ is the diagonal matrix consisting of spiked eigenvalues.

We adapt the simulation settings in Dobriban and Owen (2019) to our general spiked model. Given (n, p, K, θ) and $s_1, \dots, s_K > 0$, we generate $X_i \stackrel{\text{iid}}{\sim} N(0, \Sigma)$, $1 \leq i \leq n$, where $\Sigma = \mathbf{B}\mathbf{B}^\top + \mathbf{D}$. The matrix $\mathbf{D} = \text{diag}(\sigma_1^2, \sigma_2^2, \dots, \sigma_p^2)$ is generated in the same way as in Model (7), and $\mathbf{B}$ is a $p \times K$ matrix obtained by first generating a $p \times K$ matrix

with independent $N(0, 1)$ entries and then renormalizing each column to have an ℓ^2 -norm equal to $s_{k\sqrt{p/n}}$. Under this setting, the L_∞ -norm of each population eigenvector is only $O(p^{-1/2}\sqrt{\log(p)})$, so the “delocalization” condition is satisfied. We fix $K = 1$ and let (n, p, s_1, θ) vary. The results are shown in Table 2.

Compared with Simulation 2, the performance of DDPA is significantly better. BEMAO and BEMA continue to perform well, indicating that their performance is insensitive to the generating process of eigenvectors. This is consistent with our theoretic understanding. In Section 4, we have seen that the success of BEMAO and BEMA requires no conditions on eigenvectors.

Simulation 4. In this simulation, we investigate the case of model misspecification. We still assume that Σ is a low-rank matrix plus a residual covariance matrix $\mathbf{D}$. However, we no longer let $\mathbf{D}$ be a diagonal matrix. Below, we consider three misspecified models, where $\mathbf{D}$ is a Toeplitz matrix, a block-wise diagonal matrix, and a sparse matrix, respectively.

- In the first model, $\mathbf{D}(i, j) = (1 + |i - j|)^{-t}$, for $1 \leq i, j \leq p$. Here, $\mathbf{D}$ is a Toeplitz matrix with polynomial decays in the off-diagonal. The larger t , the closer to a diagonal matrix.
- In the second model, $\mathbf{D}(i, i) = 1$ for $1 \leq i \leq p$, and $\mathbf{D}(2j - 1, 2j) = \mathbf{D}(2j, 2j - 1) = b$ for $1 \leq j \leq p/2$. $\mathbf{D}$ is a block-wise diagonal matrix which has many 2×2 diagonal blocks. The smaller b , the closer to a diagonal matrix.
- In the third model, $\mathbf{D}(i, i) = 1$ for $1 \leq i \leq p$, and $\mathbf{D}(i, j) = \mathbf{D}(j, i) \sim c \cdot \text{Bernoulli}(0.1)$ for $i \neq j$. The matrix $\mathbf{D}$ has approximately $0.1p$ nonzero entries in each row. The smaller c , the closer to a diagonal matrix.

Table 3. Simulation 4: Comparison of different methods in three misspecified models, where the residual covariance matrix $\mathbf{D}$ is a Toeplitz matrix, a block diagonal matrix, and a sparse matrix, respectively. $(n, p, K) = (500, 100, 1)$.

λ	Residual covariance	BEMAO	BEMA	DDPA	EKC	Bai&Ng	Pass&Yao
6	Toeplitz($t=4$)	1.104 (89.6%)	1 (100%)	1.422 (65.4%)	1.36 (67.4%)	1 (100%)	1.06 (94.8%)
3	Toeplitz($t=2$)	9.352 (0%)	1.12 (88.6%)	100 (0%)	15.148 (0%)	0 (0%)	2.46 (24.6%)
6	block diagonal($b=0.1$)	1.344 (66.8%)	1 (100%)	2.378 (31.6%)	1.854 (33.6%)	1 (100%)	1.038 (96.6%)
3	block diagonal($b=0.2$)	3.764 (0%)	1 (100%)	100 (0%)	6.602 (0%)	0 (0%)	1.12 (89.8%)
6	sparse($c=0.05$)	1.784 (30.2%)	1.016 (98.4%)	5.024 (9.4%)	2.474 (9.4%)	1 (100%)	1.084 (91.6%)
3	sparse($c=0.08$)	3.348 (0%)	1.036 (96.4%)	97.752 (0%)	5.18 (0%)	0 (0%)	1.58 (47.4%)

NOTES: The spiked eigenvalue is equal to λ . For each misspecified model, we consider two settings, where $\mathbf{D}$ is closer to a diagonal matrix in the first setting (rows 1, 3, 5) than in the second setting (rows 2, 4, 6). The number in each cell is the average $\hat{K}$, and the number in brackets is the probability of correctly estimating K (correct rate).

Table 4. Simulation 5: The robustness of BEMAO and BEMA under non-Gaussian data and different values of α .

Distribution	(λ, θ)	BEMAO (0.1)	BEMAO (0.2)	BEMAO (0.3)	BEMA (0.1)	BEMA (0.2)	BEMA (0.3)
Gaussian	(1.5, ∞)	5 (100%)	5 (100%)	5 (100%)	4.95 (95%)	4.93 (93%)	4.904 (90.4%)
Random sign	(1.5, ∞)	4.996 (99.6%)	4.996 (99.6%)	4.998 (99.8%)	4.972 (97.2%)	4.96 (96%)	4.94 (94%)
Laplace	(1.5, ∞)	4.998 (99.8%)	4.998 (99.8%)	4.998 (99.8%)	4.914 (91.4%)	4.9 (90%)	4.88 (88%)
Gaussian	(4.5, 3)	—	—	—	4.518 (69%)	4.748 (81.2%)	4.76 (81%)
Random sign	(4.5, 3)	—	—	—	4.678 (78.4%)	4.818 (85%)	4.9 (85.4%)
Laplace	(4.5, 3)	—	—	—	4.352 (56.8%)	4.634 (73.8%)	4.656 (74.8%)

NOTES: Data are generated from the factor model with Gaussian/random-sign/Laplace factors and noise. $K = 5$, and all the spiked eigenvalues are equal to λ . BEMAO and BEMA are implemented with $\alpha \in \{0.1, 0.2, 0.3\}$ (denoted as BEMAO (α)/BEMA (α) in the table). The number in each cell is the average $\hat{K}$, and the number in brackets is the probability of correctly estimating K (correct rate).

The low-rank part of Σ is generated in the same way as before: We let all μ_k equal to λ and let the eigenvector matrix Ξ be drawn uniformly from the Stiefel manifold, which allows Ξ to have orthonormal columns. Fix $(n, p, K) = (500, 100, 1)$. The results are shown in Table 3.

For each misspecified model, we consider two settings, where $\mathbf{D}$ is closer to a diagonal matrix in the first setting (Rows 1, 3, and 5 of Table 3) than in the second one (Rows 2, 4, and 6 of Table 3). Every method performs better in the first case, suggesting that the diagonal assumption on $\mathbf{D}$ is indeed critical. In comparison, BEMA is least sensitive to a nondiagonal $\mathbf{D}$. In Rows 2, 4, and 6 of Table 3, the correct rate of BEMA is still above 80%, while the correct rate of some other methods is only 0%. Pass&Yao is the second least sensitive to a nondiagonal $\mathbf{D}$.

To try to understand this phenomenon, we first note that one can always apply an orthogonal transformation to data vectors $\mathbf{X}_1, \dots, \mathbf{X}_n$, so that the post-transformation data follow a different spiked covariance model whose residual covariance matrix $\tilde{\mathbf{D}}$ is a diagonal matrix containing the eigenvalues of $\mathbf{D}$. This orthogonal transformation is unknown in practice. However, if a method uses the empirical eigenvalues *only*, it does not matter whether or not we know this orthogonal transformation, because any orthogonal transformation does not change eigenvalues of the sample covariance matrix and thus it does not change the estimator of K . It implies that, for methods that only use eigenvalues, we can treat the misspecified model as if $\mathbf{D}$ is replaced by the diagonal matrix $\tilde{\mathbf{D}}$. Therefore, the surprising robustness of BEMA can be interpreted as the capability of the gamma model (2) in approximating the eigenvalue structure in $\mathbf{D}$. The flexibility of this gamma model comes from the parameter θ . In comparison, such strong robustness is not observed for BEMAO, where θ is fixed as ∞ .

The method of DDPA uses empirical eigenvectors in the procedure, thus, it is more sensitive to the diagonal assumption of $\mathbf{D}$. EKC uses eigenvalues only, but its thresholding scheme is too

conservative. In these misspecified models, some bulk empirical eigenvalues can get large; EKC gives too small thresholds to non-leading eigenvalues and yields over-estimation of K .

Simulation 5. In this simulation, we test the robustness of our proposed methods against the choice of α and the distributional assumption on data generation. Fix $(n, p, K) = (500, 100, 5)$. We generate $\mathbf{X}_i = \Xi \omega_i + \epsilon_i$, where $\Xi \in \mathbb{R}^{p \times K}$ is uniformly drawn from the Stiefel manifold, ω_i are iid drawn from a multivariate zero-mean distribution with covariance matrix $\lambda \mathbf{I}_K$, ϵ_i are iid drawn from a multivariate zero-mean distribution with covariance matrix $\mathbf{D}$, and $\mathbf{D}$ is generated in the same way as in Model (7) with $\sigma^2 = 1$ and $\theta \in \{\infty, 3\}$. We consider three settings where the entries of ω_i and ϵ_i are Gaussian, random sign, or Laplace variables (centered and rescaled to match the required variance), respectively. The results are in shown Table 4.

For the standard spiked covariance model (top 3 rows of Table 4), the results are very similar for different distributions. For the general spiked covariance model (bottom 3 rows of Table 4), the performance of BEMA increases/decreases when the data have lighter/heavier tails, but the difference is within a reasonable range. Our theory only requires a mild distributional assumption (Assumption 2), which is validated by this simulation.

The choice of α decides the fraction of bulk eigenvalues used to estimate (σ^2, θ) . The larger α , we restrict to a narrower range of eigenvalues deep into the bulk. The performance of BEMA is similar for $\alpha \in \{0.2, 0.3\}$ and slightly worse for $\alpha = 0.1$. In the asymptotic theory, α can be chosen as any constant, but for good finite-sample performance we need $(\tilde{p}\alpha - K)$ to be properly large, where $\tilde{p} = n \wedge p$. In practice, if $\tilde{p}$ is extremely large, the choice of α has a negligible effect; if $\tilde{p}$ is only moderately large, we recommend choosing a large α so that we are confident that $\tilde{p}\alpha$ is significantly larger than K .

Table 5. Simulation 6: The performance of BEMA0 and BEMA with different values of β .

Method	(λ, θ)	$\beta = 0.01$	$\beta = 0.05$	$\beta = 0.1$	$\beta = 0.2$	$\beta = 0.3$	$\beta = 0.5$
BEMA0	$(1.5, \infty)$	4.994 (99.4%)	5 (100%)	5 (100%)	5 (100%)	5 (100%)	5.006 (99.4%)
BEMA	$(1.5, \infty)$	4.712 (72.6%)	4.888 (89%)	4.93 (93%)	4.966 (96.6%)	4.982 (98.2%)	4.996 (99.6%)
BEMA	$(6, 3)$	4.734 (83.6%)	4.978 (97.2%)	5.018 (98.2%)	5.056 (94.4%)	5.082 (92%)	5.188 (83%)

NOTES: The spiked eigenvalues are all equal to λ , and the eigenvectors are randomly generated from the Stiefel manifold. BEMA0 and BEMA are implemented with $\beta \in \{0.01, 0.05, 0.1, 0.2, 0.3, 0.5\}$. The number in each cell is the average $\hat{K}$, and the number in brackets is the probability of correctly estimating K (correct rate).

Simulation 6. In this simulation, we tested the dependency of our methods upon the choice of β . Fix $(n, p, K) = (500, 100, 5)$.

We generate data $X_i \stackrel{\text{iid}}{\sim} N(0, \Sigma)$, $1 \leq i \leq n$, where Σ satisfies Model (7) with $\sigma^2 = 1$ and $\mu_k = \lambda$, for $1 \leq k \leq K$. The eigenvector matrix Ξ is drawn uniformly from the Stiefel manifold. The results are shown in Table 5.

Note that β has an explicit meaning: It is the (asymptotic) upper bound for the probability of over-estimating K . If the spike eigenvalues are all sufficiently large, we can safely choose a very small β (e.g., $\beta = 0.01$). However, if the spike eigenvalues are only moderately large, we need to choose an appropriate β to tradeoff the probability of over-estimation and the probability of under-estimation. As seen in Table 5, $\beta \in [0.05, 0.3]$ empirically gives nice and stable results.

6. Real Applications

We apply BEMA to two real datasets. We compare our method with EKC (Braeken and Van Assen 2017), Bai&Ng (Bai and Ng 2002), DDPA and its variants (Dobriban and Owen 2019), and Pass&Yao (Passemier and Yao 2014). DDPA has three versions: DPA is a deterministic implementation of parallel analysis (Horn 1965); DDPA is an improvement of DPA aiming to resolve the issue of “eigenvalue shadowing,” that is, an extremely large spiked eigenvalue shadows the other spiked eigenvalues and causes an under-estimation of K ; DDPA+ is a robust version of DDPA recommended for real data analysis. We include all three versions in comparison.

6.1. The Lung Cancer Data

The Lung Cancer dataset was collected and cleaned by Gordon et al. (2002). The original dataset contains the expression data of 12,533 genes and 181 subjects. The subjects divide into two groups, the diseased group and the normal group. Jin and Wang (2016) processed this dataset by removing genes that are not differentially expressed across subject groups and resulted in a new data matrix with $(p, n) = (251, 181)$. The selection of these 251 “influential genes” used no information of true groups, including the number of groups. We use this processed data matrix, because the original data matrix contains too many features (genes) that are irrelevant to the clustering structure, where no method gives meaningful results. It was argued in Jin and Wang (2016) that this data matrix follows a clustering model. As a result, the covariance matrix has $(K_0 - 1)$ spiked eigenvalues, where K_0 is the number of clusters. Here, the ground-truth is $K_0 = 2$, that is, the true number of spiked eigenvalues is $K = 1$.

We apply BEMA with $(\alpha, \beta, M) = (0.2, 0.1, 500)$, that is, 60%(= $1 - 2\alpha$) of the bulk eigenvalues in the middle range

are used to estimate model parameters, the probability of over-estimating K is controlled by 0.1, and 500 Monte Carlo samples are used to determine the ultimate threshold for eigenvalues. The BEMA algorithm outputs $(\hat{\theta}, \hat{\sigma}^2) = (0.288, 0.926)$. In Figure 6(a), we check the goodness of fit. If the proposed spiked covariance model (7) is suited for the data, then we expect to see $\hat{\lambda}_k \approx \hat{\sigma}^2 \cdot \bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \hat{\theta})$, except for a few small k . The left panel of Figure 6(a) plots $\hat{\lambda}_k$ vs. $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; 1, \hat{\theta})$, suggesting a good fit to a line crossing the origin. The right panel contains the scree plot, that is, $\hat{\lambda}_k$ vs. k . We also plot the curve of $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \hat{\sigma}^2, \hat{\theta})$ vs. k . This curve is a good fit to the scree plot in the middle range. These plots suggest that Model (7) is well-suited for this dataset.

The estimator of K by BEMA is $\hat{K} = 1$, which is exactly the same as the ground truth. This is the output of the algorithm by setting $\beta = 0.1$. Using the argument in Section 3.3, this is also a confidence lower bound for K . By setting $\beta = 0.9$ in the algorithm, we get a confidence upper bound which is 4. This gives an 80% confidence interval for K as $[1, 4]$. Figure 6(b) contains the scatterplots of the left singular vectors of X , colored by the true group label. The first singular vector clearly contains information for separating two groups, but other singular vectors also contain some information. This explains why the confidence upper bound is larger than 1.

The comparison with other methods is summarized in Table 6. The behavior of EKC is consistent with our observation in simulations. In this dataset, the eigenvalues of the residual covariance matrix vary widely (this can be seen from the estimated θ by BEMA, $\hat{\theta} = 0.288$, which is far from ∞), and EKC gives too small threshold to non-leading eigenvalues. The behavior of Bai&Ng is different from what we observe in simulations. Note that we have to use the effective p after the data processing by Jin and Wang (2016), where the dimension reduces from 12,533 to 251. As a result, the penalty in Bai&Ng is weaker than that in simulations, and so the method significantly over-estimates K . Pass&Yao also over-estimates K . Among DDPA and its variants, DPA performs the best. A possible reason is that DPA does not use empirical eigenvectors and is more stable than DDPA and DDPA+.

Different from all other methods, BEMA not only outputs an estimator of K but also yields a fitted model, $\text{Gamma}(\hat{\theta}, \hat{\theta}/\hat{\sigma}^2) = \text{Gamma}(0.288, 0.311)$, for eigenvalues of the residual covariance matrix. This can be useful for many other statistical inference tasks.

6.2. The 1000 Genomes Data

The 1000 Genomes Phase 3 whole genome sequencing dataset (1000 Genomes Project Consortium 2015) consists of the

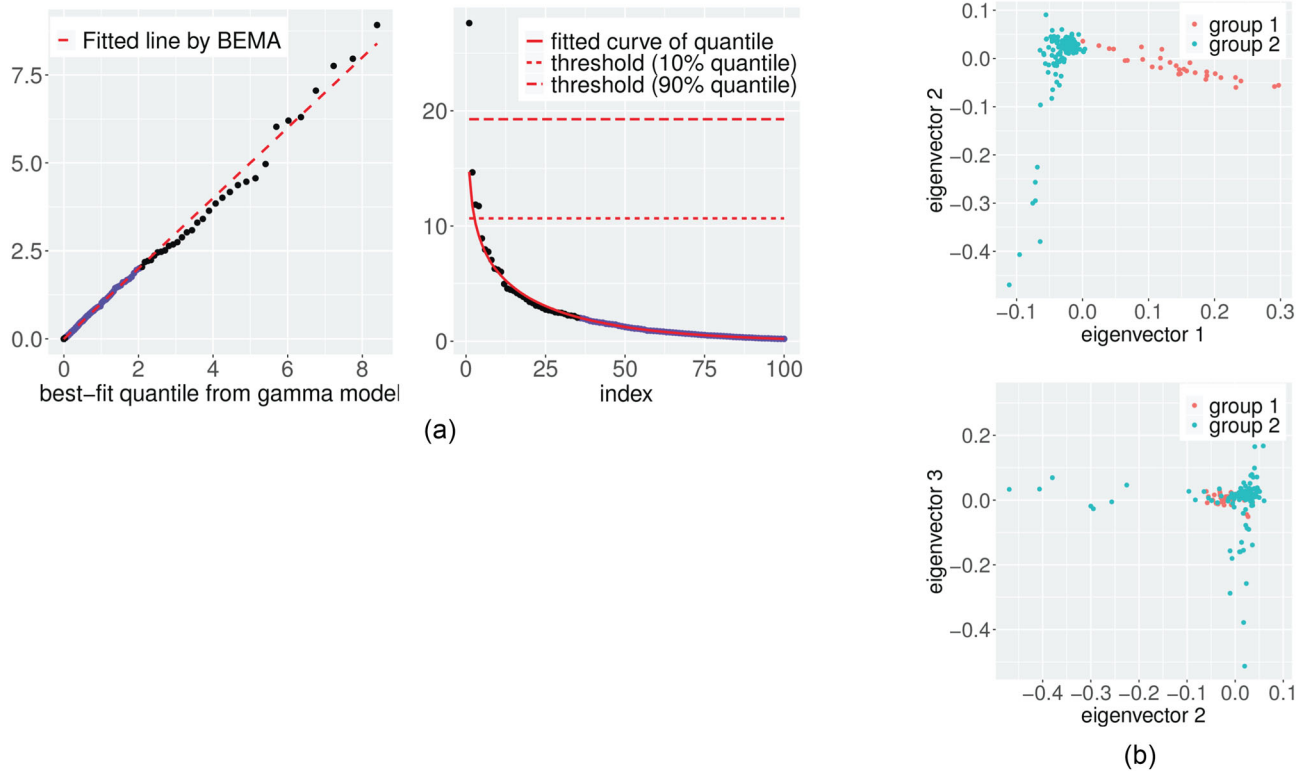


Figure 6. Results for the Lung Cancer data. (a) The goodness of fit of BEMA on the Lung Cancer data. The left panel plots $\hat{\lambda}_k$ vs. $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \hat{\sigma}^2, \hat{\theta})$ (which is quantile of the theoretical limit of ESD with estimated θ), where the first 4 eigenvalues are removed for better visualization. It fits well a line crossing the origin. The right panel plots $\hat{\lambda}_k$ vs. k , where the red solid curve is $\bar{F}_{\gamma_n}^{-1}(k/\tilde{p}; \hat{\sigma}^2, \hat{\theta})$ vs. k . The curve fits the bulk eigenvalues (blue dots). These two plots together suggest that the spiked covariance model (7) is suitable for this dataset. (b) The plots of singular vectors of X .

Table 6. Comparison of different estimators of K using two real datasets: the lung cancer gene expression data and the 1000 Genome data of genome-wide common genetic variants.

	BEMA	BEMA0	EKC	Bai&Ng	Pass&Yao	DDPA	DPA	DDPA+	Truth
Lung cancer data	1	27	56	180	8	180	1	11	1
1000 Genomes data	28	67	2503	4	28	85	20	4	25

NOTE: For BEMA and BEMA0, the choices of tuning parameters are described in the text. In the appendix, we report the results with various choices of tuning parameters, which are very stable.

genotypes of 2504 subjects for over 84.4 million variants. We restrict the analysis to common variants with minor allele frequencies greater than 0.01. There are 26 self-reported ethnicity groups, coming from five super-populations: African (AFR), Ad Mixed American (AMR), East Asian (EAS), European (EUS), and South Asian (SAS).

In view of high linkage disequilibrium (LD) among some variants, which can distort the eigenvector and eigenvalue structure (Patterson, Price, and Reich 2006), we first performed LD pruning. We used an independent pair-wise LD pruning, with window size 1000, step size 50 and a threshold 0.02 for R-squared. Restricting to LD pruned markers, we obtain a data matrix with $p = 24,248$ and $n = 2504$. The number of spiked eigenvalues equals to the number of true ancestry groups minus one (Patterson, Price, and Reich 2006). We treat the self-reported ethnicity groups as the ground truth, which gives $K = 25$.

We apply BEMA with $(\alpha, \beta, M) = (0.1, 0.1, 500)$. First, we check the goodness of fit. BEMA outputs $(\hat{\theta}, \hat{\sigma}^2) =$

$(4.256, 0.377)$. Figure 7(a) shows the Q-Q plot and the scree plot, with reference curves from the BEMA fitting. The meaning of these plots is the same as described in Section 6.1 and is also explained in the caption of this figure, which we do not repeat here. The conclusion is that our proposed spiked covariance model (7) is an excellent fit to this dataset.

The estimated model for eigenvalues of the residual covariance matrix is $\text{Gamma}(\hat{\theta}, \hat{\theta}/\hat{\sigma}^2) = \text{Gamma}(4.256, 11.3)$. We note that the variance of the genotype on each SNP is $2q(1 - q)$, where q is the null minor Allele frequency (MAF) of this SNP. We thus interpret the BEMA fitting as follows: After the ancestry effect is removed, the null MAFs q_j (on LD pruned SNPs) satisfy that $2q_j(1 - q_j) \stackrel{\text{iid}}{\sim} \text{Gamma}(4.256, 11.3)$. The mean and standard deviation of this gamma distribution is 0.377 and 0.18, respectively.

Next, we look at the estimation of K . The BEMA algorithm outputs $\hat{K} = 28$, which is very close to the ground truth $K = 25$. The 98% confidence interval of K is $[27, 31]$.

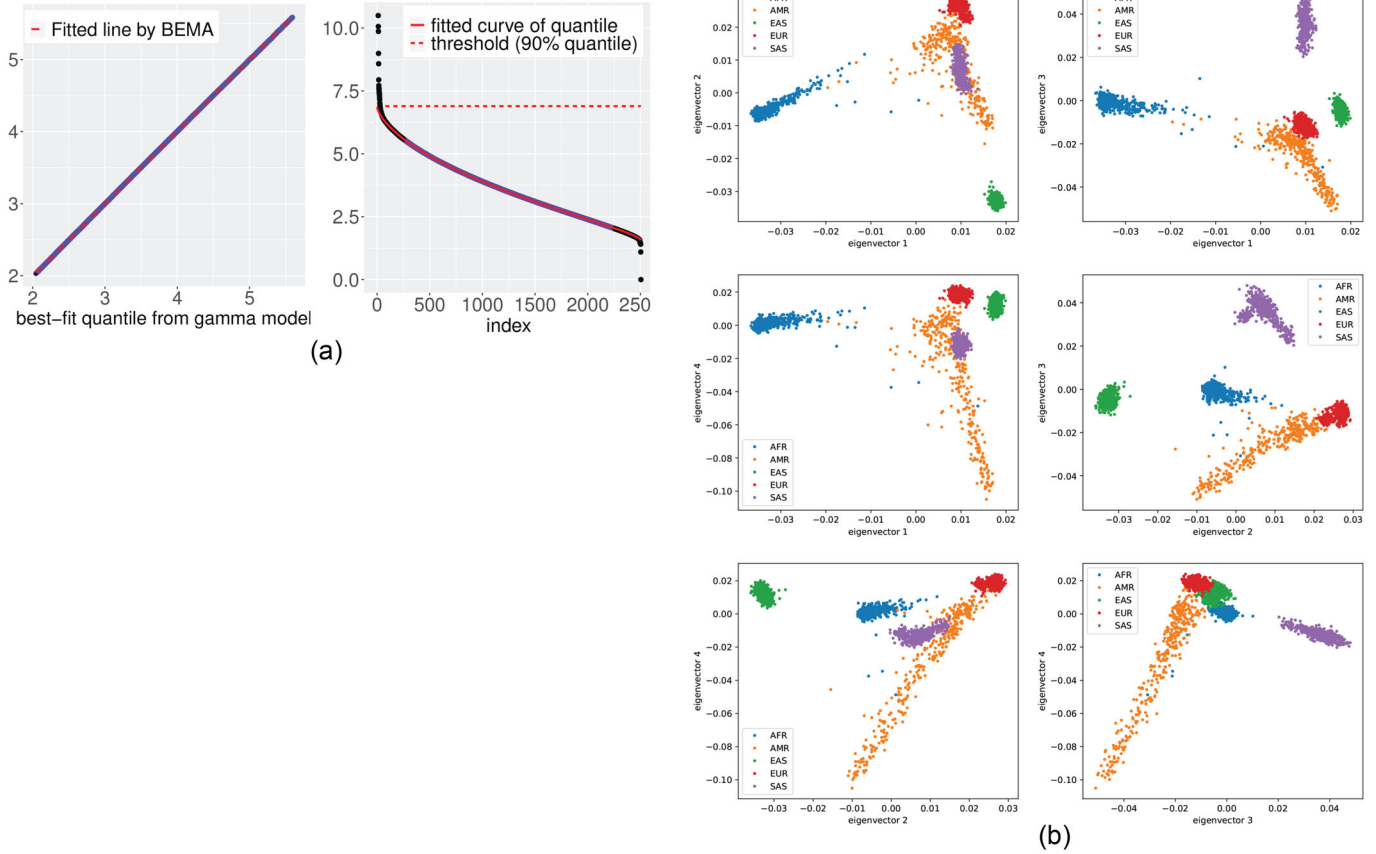


Figure 7. Results for analysis of the 1000 Genomes data. (a) The goodness of fit of BEMA on the 1000 Genomes data. The left panel is the plot of $\hat{\lambda}_k$ vs. $\bar{F}_{\gamma_n}^{-1}(k/\bar{p}; 1, \hat{\theta})$ (which is quantile of the theoretical limit of ESD with estimated θ) for $\alpha n \leq k \leq (1 - \alpha)n$. It fits well a line crossing the origin. The right panel plots $\hat{\lambda}_k$ vs. k , where the red solid curve is the curve of $\bar{F}_{\gamma_n}^{-1}(k/\bar{p}; 1, \hat{\theta})$ vs. k ; for better visualization, this curve is only plotted for $11 \leq k \leq 2504$. It fits well the bulk eigenvalues (blue dots). These two plots suggest that the spiked covariance model (7) is suitable for this dataset. (b) The plots of singular vectors of X .

A comparison with other methods is summarized in Table 6. EKC and DDPA significantly over-estimate K , and Bai&Ng and DDPA+ significantly under-estimate K . DPA gives $\hat{K} = 20$, which is relatively close to the ground truth. BEMA and Pass&Yao both give $\hat{K} = 28$, which is closest to the ground truth. Pass&Yao assumes that all σ_j^2 are equal. In this dataset, BEMA estimates the standard deviation of σ_j^2 to be 0.18, which is relatively small. This explains why Pass&Yao also performs well.

Last, we validate the results by investigating the singular vectors of X . We first measure the association between each singular vector and the true ethnicity labels by the Rayleigh quotient (Horn and Johnson 2012). Let $\hat{\eta}_k \in \mathbb{R}^n$ be the k th left singular vector of the centralized data matrix. We treat its entries as n data points and compute the ratio of between-cluster-variance and within-cluster-variance, denoted as RQ_k . A larger RQ_k indicates that $\hat{\eta}_k$ is more correlated with the true ethnicity labels. Figure 8(a) plots RQ_k vs. k . The first a few singular vectors have very high association with the ethnicity labels. These singular vectors capture the super population structure. The pairwise scatterplots of the first 4 singular vectors are contained in Figure 7(b), which show clearly that super populations are well separated on these singular vectors. Besides, the first few singular vectors, the remaining singular vectors capture more of the sub-structure within each super population. Figure 8(b)

is the parallel coordinate plot. In Figure 8(c), we regenerate parallel coordinate plots by restricting to each super population. Within the super population AMR, there is still separation of ethnicity groups for k as large as 27. This explains why BEMA outputs a $\hat{K}$ that is slightly larger than the ground truth.

7. Discussion

We propose a new method for estimating the number of spiked eigenvalues in a large covariance matrix. The novelty of our method lies in a systematic approach to incorporating bulk eigenvalues in the estimation of K . Under a working model which assumes the diagonal entries of the residual covariance matrix are iid drawn from a Gamma distribution, we fit a parametric curve on bulk eigenvalues. The estimated parameters of this curve are then used to decide a threshold for top eigenvalues and produce an estimator of K . We study the theoretical properties of our method under a standard spiked covariance model, and show that our estimator requires weaker conditions for consistent estimation of K compared with the existing methods. We examine the performance of our method using both simulated data and two real datasets. Our empirical results show that the proposed method outperforms other competitors in a variety of scenarios.

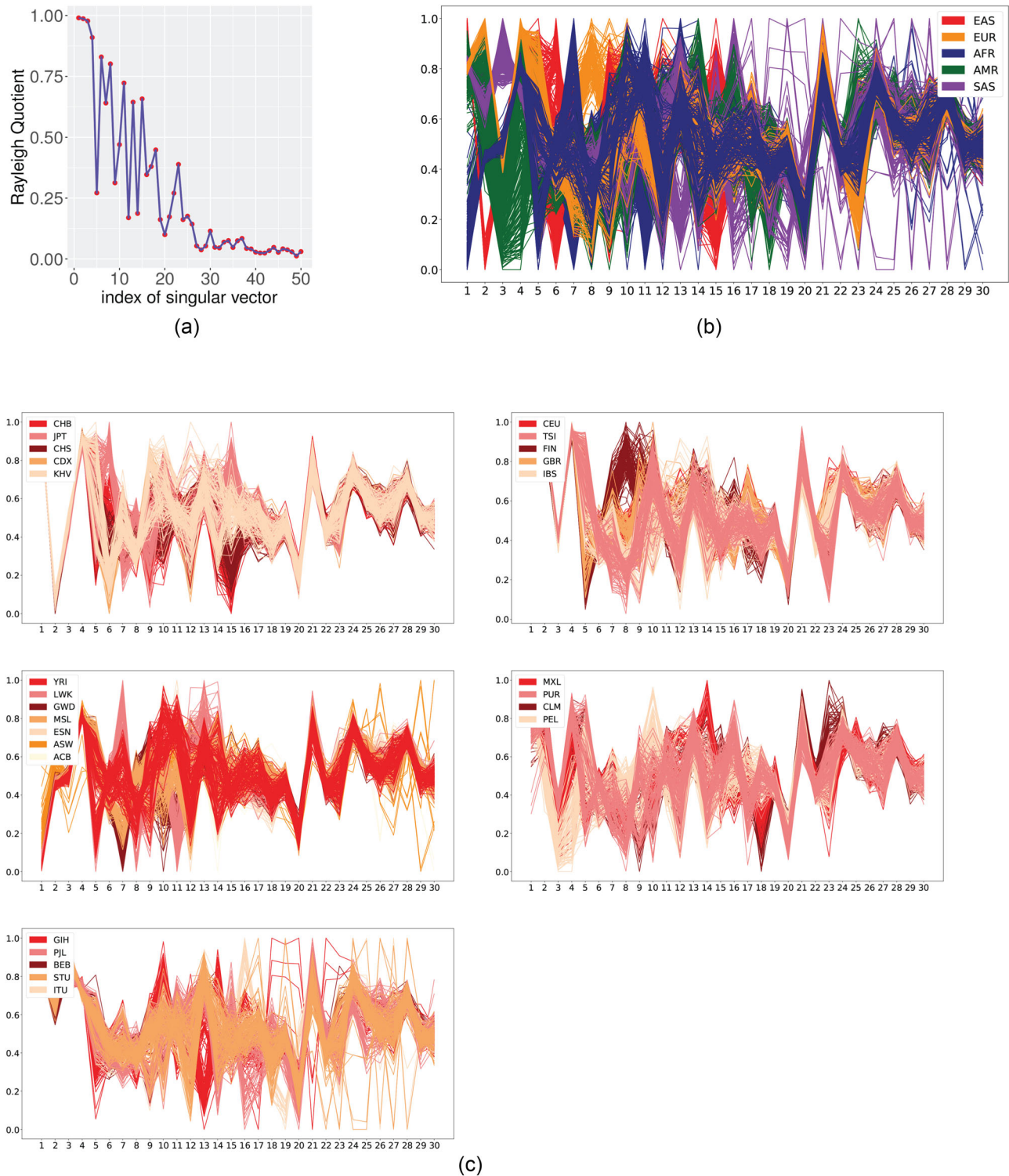


Figure 8. Interpretation of results for the 1000 Genomes data. (a) The association between singular vectors of X and the true ethnicity labels. (b) The parallel coordinate plot of singular vectors, color-coded by five super-populations. (c) The parallel coordinate plots of singular vectors for each super-population, color-coded by the ethnicity groups within each super-population. The five super-populations are EAS (top left), EUR (top right), AFR (middle left), AMR (middle right), and SAS (bottom left). The sub-population labels used in the legends of can be found in 1000 Genomes Project Consortium (2015).

Our approach is conceptually connected to the empirical null (Efron 2004) in multiple testing. The empirical null imposes a working model (e.g., a normal distribution) on Z -scores of individual null hypotheses and estimates the parameters of this distribution from a large number of Z -scores. The fitted null model is then used to correct p -values and help identify the

non-null hypotheses. Similarly, we impose a working model (i.e., a Gamma distribution) on non-spiked population eigenvalues and estimate the parameters of this distribution from a large number of bulk empirical eigenvalues. The fitted null model is then used to assist estimation of K . From this perspective, our method can be regarded as a *conceptual* application of the

empirical null approach to eigenvalues. Meanwhile, our setting is much more complicated than that in multiple testing. The bulk eigenvalues are highly correlated, and their marginal distribution has no explicit form. These impose great challenges on algorithm design and theoretical analysis.

For the theoretical study, we first analyze the special case of $\theta = \infty$. This corresponds to the well-known standard spiked covariance model (Johnstone 2001), which has attracted many theoretical interests. Our theory contributes to this literature with an explicit error bound on estimating σ^2 and consistency theory on estimating K . The theoretical study for a general θ that corresponds to the setting of heterogeneous residual variances is of great interest but is technically challenging. Instead, we study a proxy model where the population eigenvalues are iid drawn from a truncated Gamma distribution. Under this model we derive error bounds for $(\hat{\sigma}^2, \hat{\theta})$ and prove the consistency of $\hat{K}$ with mild conditions. The analysis uses advanced results in random matrix theory (Bloemendal et al. 2016; Knowles and Yin 2017; Ding 2020).

The method can be extended in multiple directions. Here, we assume that the diagonal entries of the residual covariance matrix are from a Gamma distribution. It can be generalized to other parametric distributions. In Section 4.2, we have already seen a variant of our method by using a truncated Gamma distribution, which assumption helps eliminate extremely large variances for the residuals. We can also use a mixture of Gamma distributions to accommodate heterogeneous feature groups. Our main algorithm can be easily adapted to such cases. When the distribution family is unknown, we may combine our method with the techniques in nonparametric density estimation. The thresholding scheme in our method can also be modified. We currently apply a single threshold to all eigenvalues. Alternatively, we may use different thresholds for different eigenvalues. One proposal is to use the $(1 - \beta)$ -quantile of the distribution of $\hat{\lambda}_k^*$ in the null model (12) as a threshold for $\hat{\lambda}_k$. We leave these extensions to future work.

In the numerical experiments, our method exhibits robustness to model misspecification. It is suggested by Simulation 4 of Section 5 that our method continues to work when the residual covariance matrix is a Toeplitz matrix, or a block-wise diagonal matrix, or a sparse matrix. A theoretical understanding to this phenomenon will be useful. As stated in Section 5, we have observed empirically that there always exist (σ^2, θ) such that the theoretical limit of ESD induced by the Gamma model (2) can reasonably approximate the theoretical limit of ESD induced by a Toeplitz or block-wise diagonal or sparse covariance matrix. It remains an interesting question on how to justify it theoretically. We leave it to future work.

Acknowledgments

The authors thank to Rounak Dey and Derek Shyr for their help on downloading and pruning the 1000 Genomes dataset, and Zhigang Bao and Xiukai Ding for helpful pointers on random matrix theory.

Funding

This work was supported by the National Institutes of Health grants R35-CA197449, U01-HG009088, U19-CA203654 (Lin), and National Science Foundation grant DMS-1925845 (Ke).

Supplementary material

The supplementary material contains the description of two GetQT algorithms, the proof of main theorems and lemmas, and the results of BEMA on real datasets with varying α .

Software

The code for implementing BEMA is available at <https://github.com/ZhengTracyKe/BEMA>

ORCID

Xihong Lin  <http://orcid.org/0000-0001-7067-7752>

References

- Bai, J., and Ng, S. (2002), “Determining the Number of Factors in Approximate Factor Models,” *Econometrica*, 70, 191–221. [1,8,10,14]
- Baik, J., Ben Arous, G., and Peche, S. (2005), “Phase Transition of the Largest Eigenvalue for Nonnull Complex Sample Covariance Matrices,” *The Annals of Probability*, 33, 1643–1697. [2,3,8,10]
- Bao, Z. (2020), Personal Communications. [9]
- Bloemendal, A., Knowles, A., Yau, H.-T., and Yin, J. (2016), “On the Principal Components of Sample Covariance Matrices,” *Probability Theory and Related Fields*, 164, 459–552. [2,3,4,5,7,8,9,18]
- Braeken, J., and Van Assen, M. A. (2017), “An Empirical Kaiser Criterion,” *Psychological Methods*, 22, 450. [2,10,14]
- Cai, T. T., Han, X., and Pan, G. (2020), “Limiting Laws for Divergent Spiked Eigenvalues and Largest Nonspiked Eigenvalue of Sample Covariance Matrices,” *Annals of Statistics*, 48, 1255–1280. [2,8]
- Ding, X. (2020), “Spiked Sample Covariance Matrices With Possibly Multiple Bulk Components,” *Random Matrices: Theory and Applications*, 2150014. [9,18]
- Dobriban, E. (2015), “Efficient Computation of Limit Spectra of Sample Covariance Matrices,” *Random Matrices: Theory and Applications*, 4, 1550019. [2]
- Dobriban, E., and Owen, A. B. (2019), “Deterministic Parallel Analysis: An Improved Method for Selecting Factors and Principal Components,” *Journal of the Royal Statistical Society, Series B*, 81, 163–183. [2,8,10,12,14]
- Donoho, D. L., Gavish, M., and Johnstone, I. M. (2018), “Optimal Shrinkage of Eigenvalues in the Spiked Covariance Model,” *The Annals of Statistics*, 46, 1742. [2,3,8]
- Efron, B. (2004), “Large-Scale Simultaneous Hypothesis Testing: The Choice of a Null Hypothesis,” *Journal of the American Statistical Association*, 99, 96–104. [17]
- Fan, J., Guo, J., and Zheng, S. (2020), “Estimating Number of Factors by Adjusted Eigenvalues Thresholding,” *Journal of the American Statistical Association*, 1–33. [2,8,10]
- Fan, J., Liao, Y., and Mincheva, M. (2013), “Large Covariance Estimation by Thresholding Principal Orthogonal Complements,” *Journal of the Royal Statistical Society, Series B*, 75, 603–680. [1]
- Gavish, M., and Donoho, D. L. (2014), “The Optimal Hard Threshold for Singular Values is $4/\sqrt{3}$,” *IEEE Transactions on Information Theory*, 60, 5040–5053. [5,8]
- 1000 Genomes Project Consortium, A. (2015), “A Global Reference for Human Genetic Variation,” *Nature*, 526, 68. [14,17]
- Gordon, G. J., Jensen, R. V., Hsiao, L.-L., Gullans, S. R., Blumenstock, J. E., Ramaswamy, S., Richards, W. G., Sugarbaker, D. J., and Bueno, R. (2002), “Translation of Microarray Data Into Clinically Relevant Cancer Diagnostic Tests Using Gene Expression Ratios in Lung Cancer and Mesothelioma,” *Cancer Research*, 62, 4963–4967. [14]
- Götze, F., Tikhomirov, A. (2004), “Rate of Convergence in Probability to the Marchenko-Pastur Law,” *Bernoulli*, 10, 503–548. [3]
- Horn, J. L. (1965), “A Rationale and Test for the Number of Factors in Factor Analysis,” *Psychometrika*, 30, 179–185. [2,7,14]

- Horn, R. A., and Johnson, C. R. (2012), *Matrix Analysis*, Cambridge: Cambridge University Press. [16]
- Jin, J., Ke, Z. T., and Wang, W. (2017), "Phase Transitions for High Dimensional Clustering and Related Problems," *The Annals of Statistics*, 45, 2151–2189. [1]
- Jin, J., and Wang, W. (2016), "Influential Features PCA for High Dimensional Clustering," *The Annals of Statistics*, 44, 2323–2359. [14]
- Johnstone, I. M. (2001), "On the Distribution of the Largest Eigenvalue in Principal Components Analysis," *The Annals of Statistics*, 29, 295–327. [1,2,3,4,8,18]
- Ke, Z. T. (2016), "Detecting Rare and Weak Spikes in Large Covariance Matrices," arXiv no. 1609.00883. [8]
- Knowles, A., and Yin, J. (2017), "Anisotropic Local Laws for Random Matrices," *Probability Theory and Related Fields*, 169, 257–352. [5,7,8,9,18]
- Kritchman, S., and Nadler, B. (2009), "Non-Parametric Detection of the Number of Signals: Hypothesis Testing and Random Matrix Theory," *IEEE Transactions on Signal Processing*, 57, 3930–3941. [5,8]
- Kwak, J., Lee, J. O., and Park, J. (2019), "Extremal Eigenvalues of Sample Covariance Matrices With General Population," arXiv no. 1908.07444. [8,10]
- Marcenko, V. A., and Pastur, L. A. (1967), "Distribution of Eigenvalues for Some Sets of Random Matrices," *Mathematics of the USSR-Sbornik*, 1, 457–483. [5,10]
- Onatski, A. (2009), "Testing Hypotheses About the Number of Factors in Large Factor Models," *Econometrica*, 77, 1447–1479. [2]
- (2010), "Determining the Number of Factors From Empirical Distribution of Eigenvalues," *The Review of Economics and Statistics*, 92, 1004–1016. [2]
- Passemier, D., and Yao, J. (2014), "Estimation of the Number of Spikes, Possibly Equal, in the High-Dimensional Case," *Journal of Multivariate Analysis*, 127, 173–183. [2,10,14]
- Patterson, N., Price, A. L., and Reich, D. (2006), "Population Structure and Eigenanalysis," *PLoS Genetics*, 2, e190. [1,15]
- Paul, D. (2007), "Asymptotics of Sample Eigenstructure for a Large Dimensional Spiked Covariance Model," *Statistica Sinica*, 17, 1617–1642. [3]
- Shabalin, A. A., and Nobel, A. B. (2013), "Reconstruction of a Low-Rank Matrix in the Presence of Gaussian Noise," *Journal of Multivariate Analysis*, 118, 67–76. [3,5,8]
- Silverstein, J. W. (2009), "The Stieltjes Transform and Its Role in Eigenvalue Behavior of Large Dimensional Random Matrices," in *Random Matrix Theory and Its Applications*, Lecture Notes Series, Singapore: Institute for Mathematical Sciences, National University of Singapore, 18, 1–25. <https://www.worldscientific.com/worldscibooks/10.1142/7285> [2]
- Uhlig, H. (1994), "On Singular Wishart and Singular Multivariate Beta Distributions," *The Annals of Statistics*, 22, 395–405. [3]
- Wax, M., and Kailath, T. (1985), "Detection of Signals by Information Theoretic Criteria," *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 33, 387–392. [1]