

Optimal singular value shrinkage for operator norm loss: Extending to non-square matrices

William Leeb*

Abstract

We correct a formula of Gavish and Donoho for singular value shrinkage with operator norm loss for non-square matrices. We also observe that in the classical regime, optimal shrinkage for any Schatten loss converges to the best linear predictor.

1 Introduction

Low-rank matrix denoising is the task of estimating a low-rank matrix $\mathbf{X}$ from a noisy observed matrix of the form $\mathbf{Y} = \mathbf{X} + \mathbf{G}$, where $\mathbf{G}$ is a matrix of noise. The method of *singular value shrinkage*, which keeps the singular vectors of $\mathbf{Y}$ while deflating the singular values to remove the effects of noise, is a popular and well-studied approach to matrix denoising [15], [8], [7], [3], [12], [4], [13], [17], [10], [11], [2], [6].

In foundational work on this topic, Gavish and Donoho [8] observe that the optimal singular value shrinker depends crucially on the choice of loss function between $\mathbf{X}$ and the estimated matrix $\hat{\mathbf{X}}$, and provide a powerful framework for deriving asymptotically optimal singular value shrinkers for many loss functions in the spiked model [9]. However, the derivation of the optimal shrinker for operator norm loss $\|\hat{\mathbf{X}} - \mathbf{X}\|_{\text{op}}$ appears to contain a mistake. In this note, we derive the correct formula, and observe that the optimal shrinker matches the shrinker proposed in [8] only in the special case of square matrices. We also show that when $\mathbf{X}$'s columns are iid random vectors, then the optimal shrinker for any Schatten quasinorm loss converges to the best linear predictor of each column in the limiting regime of aspect ratio zero.

*School of Mathematics, University of Minnesota, Twin Cities. Minneapolis, MN.

2 Preliminaries

2.1 Model and estimation problem

We observe a matrix $\mathbf{Y} = \mathbf{X} + \mathbf{G}$ of size p -by- n , where $\mathbf{G}$ has entries which are iid $N(0, 1/n)$ and $\mathbf{X}$ is rank r with singular value decomposition $\mathbf{X} = \sum_{k=1}^r t_k \mathbf{u}_k \mathbf{v}_k^T$. Here, $\mathbf{u}_1, \dots, \mathbf{u}_r$ and $\mathbf{v}_1, \dots, \mathbf{v}_r$ are orthonormal vectors in $\mathbb{R}^p$ and $\mathbb{R}^n$, respectively, and $t_1 > \dots > t_r > 0$ are the singular values of $\mathbf{X}$. This model is generally referred to as a *spiked model* [9]. We study the spiked model as both n and $p = p(n)$ grow to infinity, and their ratio converges to a parameter $\gamma = \lim_{n \rightarrow \infty} p(n)/n$, called the *aspect ratio*. We assume that $t_1, \dots, t_r$ remain fixed as $p, n \rightarrow \infty$.

As in [8], our goal is estimating $\mathbf{X}$ from $\mathbf{Y}$ by a singular value shrinkage estimator of the form $\widehat{\mathbf{X}}^{\mathbf{q}} = \sum_{k=1}^r q_k \hat{\mathbf{u}}_k \hat{\mathbf{v}}_k^T$, where $\hat{\mathbf{u}}_1, \dots, \hat{\mathbf{u}}_r$ and $\hat{\mathbf{v}}_1, \dots, \hat{\mathbf{v}}_r$ are the top r singular vectors of $\mathbf{Y}$, and $\mathbf{q} = (q_1, \dots, q_r)$ is the vector of $\widehat{\mathbf{X}}^{\mathbf{q}}$'s singular values. [8] considers loss functions $\mathcal{L}_{p,n}(\widehat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X})$ that are orthogonally-invariant, block-decomposable, and continuous as a function of $\widehat{\mathbf{X}}^{\mathbf{q}}$. The typical examples of such loss functions are the Schatten loss functions, of the form $\mathcal{L}_{p,n}^{(\ell)}(\widehat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X}) = \|\widehat{\mathbf{X}}^{\mathbf{q}} - \mathbf{X}\|_{S_\ell}^\ell$, where $\|\Delta\|_{S_\ell}$ is the Schatten ℓ -quasinorm, defined for $\ell \in (0, \infty)$ by

$$\|\Delta\|_{S_\ell} = \left(\sum_{k=1}^{\min\{p,n\}} \lambda_k^\ell \right)^{1/\ell}, \quad (2.1)$$

where $\lambda_1, \dots, \lambda_{\min\{p,n\}}$ Δ 's singular values; and for $\ell = \infty$ by $\|\Delta\|_{S_\infty} = \|\Delta\|_{\text{op}}$. We use the notation $\mathcal{L}_{p,n}$ and $\mathcal{L}_\infty$, without superscripts, for operator norm loss ($\ell = \infty$).

For such a loss function and any choice of $\mathbf{q} = (q_1, \dots, q_r)$, the asymptotic loss $\mathcal{L}_\infty^{(\ell)}(\mathbf{q}) = \lim_{p,n \rightarrow \infty} \mathcal{L}_{p,n}^{(\ell)}(\widehat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X})$ is well-defined almost surely. The task is then to find the $\mathbf{q}^*$ minimizing $\mathcal{L}_\infty(\mathbf{q})^{(\ell)}$. The paper [8] introduces a method for this problem, which we review in Section 2.3.

2.2 Asymptotics of the spiked model

The spiked model has been well-studied in the statistics and random matrix literature. Writing the SVD of $\mathbf{Y}$ as $\mathbf{Y} = \sum_{k=1}^{\min(p,n)} \sigma_k \hat{\mathbf{u}}_k \hat{\mathbf{v}}_k^T$, we can summarize the relevant results as follows:

Proposition 2.1. *For $1 \leq k \leq r$, the k^{th} squared singular value of $\mathbf{Y}$ converges almost surely to the following deterministic limit:*

$$\sigma_k^2 = \begin{cases} (t_k^2 + 1) \left(1 + \frac{\gamma}{t_k^2}\right), & \text{if } t_k > \gamma^{1/4}, \\ (1 + \sqrt{\gamma})^2, & \text{if } t_k \leq \gamma^{1/4}, \end{cases} \quad (2.2)$$

For $1 \leq j, k \leq r$, the squared cosines between the j^{th} and k^{th} singular vectors of $\mathbf{X}$ and $\mathbf{Y}$ converge almost surely to the following limits:

$$\lim_{p \rightarrow \infty} \langle \hat{\mathbf{u}}_j, \mathbf{u}_k \rangle^2 = c_{jk}^2 = \begin{cases} \frac{1-\gamma/t_k^4}{1+\gamma/t_k^2}, & \text{if } j = k \text{ and } t_k > \gamma^{1/4}, \\ 0, & \text{if } j \neq k \text{ or } t_k \leq \gamma^{1/4}, \end{cases} \quad (2.3)$$

and

$$\lim_{n \rightarrow \infty} \langle \hat{\mathbf{v}}_j, \mathbf{v}_k \rangle^2 = \tilde{c}_{jk}^2 = \begin{cases} \frac{1-\gamma/t_k^4}{1+\gamma/t_k^2}, & \text{if } j = k \text{ and } t_k > \gamma^{1/4}, \\ 0, & \text{if } j \neq k \text{ or } t_k \leq \gamma^{1/4}. \end{cases} \quad (2.4)$$

The proof of this result in the case of Gaussian noise may be found in [16]; generalizations to other noise models appear in [1]. For each singular value σ_k of $\mathbf{Y}$ with $\sigma_k > 1 + \sqrt{\gamma}$, we may estimate t_k by inverting formula (2.2) and using $t_k > \gamma^{1/4}$:

$$t_k = \sqrt{\frac{\sigma_k^2 - 1 - \gamma + \sqrt{(\sigma_k^2 - 1 - \gamma)^2 - 4\gamma}}{2}}. \quad (2.5)$$

From t_k , the cosines c_k and $\tilde{c}_k$ are estimable by directly applying formulas (2.3) and (2.4).

2.3 The shrinkage framework of [8]

In [8], Gavish and Donoho use the behavior of the spiked model described in Proposition 2.1 to construct a framework for deriving asymptotically optimal singular value shrinkers. We review this framework here as it applies to Schatten ℓ -loss. The key observation is that when $0 < \ell < \infty$, the Schatten ℓ -loss $\mathcal{L}_{p,n}^{(\ell)}(\hat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X})$ converges asymptotically to the following expression:

$$\lim_{p,n \rightarrow \infty} \mathcal{L}_{p,n}^{(\ell)}(\hat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X}) = \mathcal{L}_{\infty}^{(\ell)}(\mathbf{q}) = \sum_{k=1}^r \left\| \begin{pmatrix} t_k & 0 \\ 0 & 0 \end{pmatrix} - q_k \begin{pmatrix} c_k \tilde{c}_k & c_k \tilde{s}_k \\ s_k \tilde{c}_k & s_k \tilde{s}_k \end{pmatrix} \right\|_{S_{\ell}}^{\ell}, \quad (2.6)$$

whereas for $\ell = \infty$ (the case of operator norm loss), we have

$$\lim_{p,n \rightarrow \infty} \mathcal{L}_{p,n}(\hat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X}) = \mathcal{L}_{\infty}(\mathbf{q}) = \max_{1 \leq k \leq r} \left\| \begin{pmatrix} t_k & 0 \\ 0 & 0 \end{pmatrix} - q_k \begin{pmatrix} c_k \tilde{c}_k & c_k \tilde{s}_k \\ s_k \tilde{c}_k & s_k \tilde{s}_k \end{pmatrix} \right\|_{\text{op}}. \quad (2.7)$$

Here, $s_k = \sqrt{1 - c_k^2}$ and $\tilde{s}_k = \sqrt{1 - \tilde{c}_k^2}$.

Consequently, the asymptotically optimal singular values q_k^* can be solved for as follows:

$$q_k^* = \operatorname{argmin}_{q_k} \left\| \begin{pmatrix} t_k & 0 \\ 0 & 0 \end{pmatrix} - q_k \begin{pmatrix} c_k \tilde{c}_k & c_k \tilde{s}_k \\ s_k \tilde{c}_k & s_k \tilde{s}_k \end{pmatrix} \right\|_{S_\ell}. \quad (2.8)$$

Explicit formulas for q_k^* when $\ell = 1, 2, \infty$ (corresponding to nuclear norm loss, Frobenius norm loss, and operator norm loss, respectively) are derived in [8]. However, the solution presented there for operator norm loss is not correct. In Section 3, specifically Theorem 3.1, we will present the correct formula for the optimal q_k^* that solves the minimization problem (2.8) for $\ell = \infty$.

3 The optimal singular values

In this section, we derive the optimal singular values q_k^* and the resulting asymptotic operator norm loss $\mathcal{L}_\infty(\mathbf{q}^*)$. The main result is the following:

Theorem 3.1. *The optimal singular value shrinkage estimator $\widehat{\mathbf{X}}^{\mathbf{q}^*}$ of $\mathbf{X}$ has singular values*

$$q_k^* = \begin{cases} t_k \sqrt{\frac{t_k^2 + \min\{1, \gamma\}}{t_k^2 + \max\{1, \gamma\}}}, & \text{if } \sigma_k > 1 + \sqrt{\gamma} \\ 0, & \text{otherwise} \end{cases}, \quad (3.1)$$

where $1 \leq k \leq r$. The loss $\|\widehat{\mathbf{X}}^{\mathbf{q}^*} - \mathbf{X}\|_{\text{op}}$ converges almost surely to

$$\mathcal{L}_\infty(\mathbf{q}^*) = \max_{1 \leq k \leq r} t_k \sqrt{1 - \min\{c_k^2, \tilde{c}_k^2\}} = t_1 \sqrt{1 - \min\{c_1^2, \tilde{c}_1^2\}}. \quad (3.2)$$

Remark 1. The optimal q_k^* and loss $\mathcal{L}_\infty(\mathbf{q}^*)$ may be consistently estimated from the observed singular values σ_k of $\mathbf{Y}$ using (2.2), (2.3), and (2.4).

Remark 2. The paper [8] proposes the shrinker $\widehat{\mathbf{X}}^{\mathbf{q}}$ with singular values $q_k = t_k$. From Theorem 3.1, this is only optimal when $\gamma = 1$, and is suboptimal for all other γ (the mistake in [8] appears to occur in Section VI, part B, page 2146). In Figure 1, we plot both shrinkers $q = q^*$ and $q = t$ and their asymptotic losses as functions of the observed singular value σ .

Theorem 3.1 is a consequence of the following result:

Lemma 3.2. *Let $t > 0$, and $0 \leq c \leq 1$, $0 \leq \tilde{c} \leq 1$. For any $q \in \mathbb{R}$, define the 2-by-2 matrix*

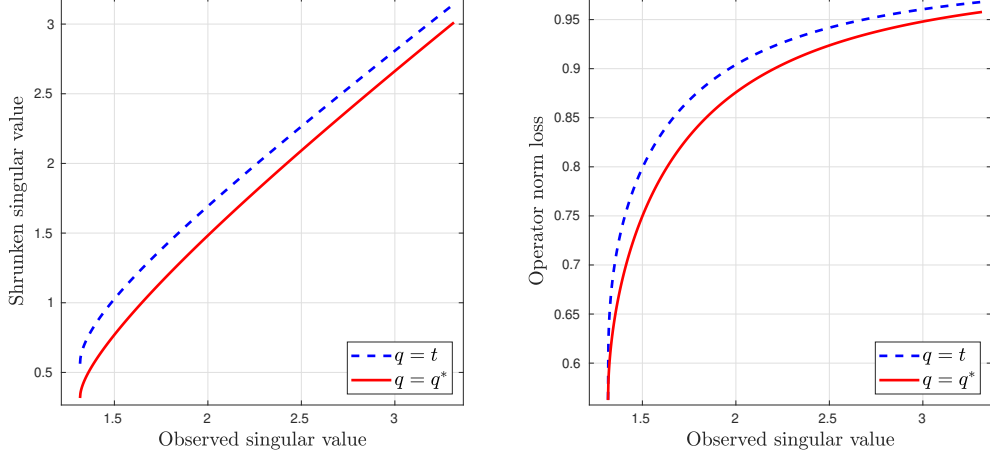


Figure 1: Singular values and asymptotic errors, rank 1 signal, $\gamma = 0.1$. Left: The singular values $q = q^*$ and $q = t$ as functions of the observed singular value σ . Right: The asymptotic operator norm losses as functions of σ .

$\mathbf{D}(q)$ by

$$\mathbf{D}(q) = \begin{pmatrix} t & 0 \\ 0 & 0 \end{pmatrix} - q \begin{pmatrix} c\tilde{c} & c\tilde{s} \\ s\tilde{c} & s\tilde{s} \end{pmatrix}, \quad (3.3)$$

where $s = \sqrt{1 - c^2}$ and $\tilde{s} = \sqrt{1 - \tilde{c}^2}$. Define $F(q) = \|\mathbf{D}(q)\|_{\text{op}}^2$. Then the minimizer q^* of F is

$$q^* = t \cdot \frac{\min\{c, \tilde{c}\}}{\max\{c, \tilde{c}\}}, \quad (3.4)$$

when c and $\tilde{c}$ are not both 0; and q^* may be taken as any value with $|q^*| \leq t$ if $c = \tilde{c} = 0$. Furthermore, $F(q^*) = t^2 \cdot \max\{s^2, \tilde{s}^2\}$.

Remark 3. Lemma 3.2 is applicable for general parameters $t > 0$ and $0 \leq c \leq 1$, $0 \leq \tilde{c} \leq 1$, even when they do not satisfy the relationships (2.2), (2.3) and (2.4).

We now prove Theorem 3.1 assuming Lemma 3.2. If $\sigma_k > 1 + \sqrt{\gamma}$, or equivalently $t_k > \gamma^{1/4}$, Lemma 3.2 says that the optimal singular value is $q_k^* = t_k \min\{c_k, \tilde{c}_k\} / \max\{c_k, \tilde{c}_k\}$. Using the expressions (2.3) and (2.4), formula (3.1) follows immediately; indeed, when $\gamma \leq 1$ we have

$$q_k^* = t_k \frac{\min\{c_k, \tilde{c}_k\}}{\max\{c_k, \tilde{c}_k\}} = t_k \frac{\tilde{c}_k}{c_k} = t_k \sqrt{\frac{t_k^4 + \gamma t_k^2}{t_k^4 + t_k^2}} = t_k \sqrt{\frac{t_k^2 + \min\{1, \gamma\}}{t_k^2 + \max\{1, \gamma\}}}, \quad (3.5)$$

and similarly when $\gamma \geq 1$. When $\sigma \leq 1 + \sqrt{\gamma}$, or equivalently $t_k \leq \gamma^{1/4}$, both c_k and $\tilde{c}_k$ are zero. Consequently, $q_k^* = 0$ is optimal.

All that remains is to show $\max_{1 \leq k \leq r} t_k \sqrt{1 - \min\{c_k^2, \tilde{c}_k^2\}} = t_1 \sqrt{1 - \min\{c_1^2, \tilde{c}_1^2\}}$. When $\gamma \leq 1$,

we have $\tilde{c} \leq c$. Then

$$t^2(1 - \tilde{c}^2) = t^2 \frac{t^2 + \gamma}{t^4 + t^2} = \frac{t^2 + \gamma}{t^2 + 1} \quad (3.6)$$

is an increasing function of t , and is equal to $\sqrt{\gamma}$ (the largest possible error for any component $t_k \leq \gamma^{1/4}$) when $t = \gamma^{1/4}$. Consequently, the maximum error is achieved at $t = t_1$. A nearly identical argument works when $\gamma \geq 1$. This completes the proof of Theorem 3.1.

It remains to prove Lemma 3.2.

Proof of Lemma 3.2. First, suppose $c = \tilde{c} = 0$. Then

$$\mathbf{D}(q) = \begin{pmatrix} t & 0 \\ 0 & -q \end{pmatrix}, \quad (3.7)$$

and so $F(q) = \|\mathbf{D}(q)\|_{\text{op}}^2 = \max\{t^2, q^2\}$. Consequently, any q with $|q| \leq t$ minimizes $F(q)$, and since $s = \tilde{s} = 1$, $F(q) = t^2 \cdot \max\{s^2, \tilde{s}^2\}$ for such q .

Next, assume that at least one of c and/or $\tilde{c}$ is not 0. The proof when $c = \tilde{c}$ is identical to the proof for optimal shrinkage of eigenvalues for covariance estimation contained in [5]; so we will assume that $c \neq \tilde{c}$. Because $F(q)$ is convex, we must show that $F'(q^*) = 0$. We will use the expression for the operator norm as a function of q derived in [8]. We have:

$$\begin{aligned} F(q) &= \frac{q^2 + t^2 - 2qtc\tilde{c} + \sqrt{(q^2 + t^2 - 2qtc\tilde{c})^2 - 4(tqss\tilde{s})^2}}{2} \\ &= \frac{A(q) + \sqrt{A(q)^2 - 4B(q)^2}}{2}, \end{aligned} \quad (3.8)$$

where $A(q) = q^2 + t^2 - 2qtc\tilde{c}$ and $B(q) = -tqss\tilde{s}$. The function $F(q)$ is differentiable whenever $A(q)^2 - 4B(q)^2 > 0$. Furthermore, at points where F is differentiable, its derivative is given by

$$F'(q) = \frac{1}{2} \left(A'(q) + \frac{A(q)A'(q) - 4B(q)B'(q)}{\sqrt{A(q)^2 - 4B(q)^2}} \right). \quad (3.9)$$

Suppose, without loss of generality, that $\tilde{c} < c$. Then the proposed minimizer of F is $q^* = t\tilde{c}/c$; we will show that $F'(q^*) = 0$. First, we have $A'(q) = 2q - 2tc\tilde{c}$. Consequently

$$A'(q^*) = 2t\frac{\tilde{c}}{c} - 2tc\tilde{c} = 2t\frac{\tilde{c}}{c}(1 - c^2) = 2t\frac{\tilde{c}}{c}s^2. \quad (3.10)$$

We also have

$$A(q^*) = t^2 \left(1 + \frac{\tilde{c}^2}{c^2} - 2\tilde{c}^2 \right), \quad (3.11)$$

and so

$$A(q^*)A'(q^*) = 2t^3 s^2 \frac{\tilde{c}}{c} \left(1 + \frac{\tilde{c}^2}{c^2} - 2\tilde{c}^2 \right). \quad (3.12)$$

Next, observe that for all q , $B'(q) = -ts\tilde{s}$. We also have $B(q^*) = -t^2 \frac{\tilde{c}}{c} s\tilde{s}$, and so

$$B(q^*)B'(q^*) = t^3 s^2 \frac{\tilde{c}}{c} (1 - \tilde{c}^2). \quad (3.13)$$

Combining (3.12) and (3.13), we obtain:

$$\begin{aligned} A(q^*)A'(q^*) - 4B(q^*)B'(q^*) &= 2t^3 s^2 \frac{\tilde{c}}{c} \left(1 + \frac{\tilde{c}^2}{c^2} - 2\tilde{c}^2 - 2 + 2\tilde{c}^2 \right) \\ &= 2t^3 s^2 \frac{\tilde{c}}{c} \left(\frac{\tilde{c}^2}{c^2} - 1 \right) \\ &= 2t^3 s^2 \frac{\tilde{c}}{c^3} (\tilde{c}^2 - c^2). \end{aligned} \quad (3.14)$$

Next, we observe that we may write

$$A(q^*)^2 - 4B(q^*)^2 = t^4 \left(1 - \frac{\tilde{c}^2}{c^2} \right)^2, \quad (3.15)$$

and consequently,

$$\sqrt{A(q^*)^2 - 4B(q^*)^2} = \frac{t^2}{c^2} (c^2 - \tilde{c}^2), \quad (3.16)$$

where we have used the fact that $c > \tilde{c}$. Note that (3.16) implies that F is differentiable at q^* whenever $c \neq \tilde{c}$.

Combining (3.14) and (3.16), we get:

$$\frac{A(q^*)A'(q^*) - 4B(q^*)B'(q^*)}{\sqrt{A(q^*)^2 - 4B(q^*)^2}} = -2t \frac{\tilde{c}}{c} s^2. \quad (3.17)$$

Consequently,

$$\begin{aligned}
F'(q^*) &= \frac{1}{2} \left(A'(q^*) + \frac{A(q^*)A'(q^*) - 4B(q^*)B'(q^*)}{\sqrt{A(q^*)^2 - 4B(q^*)^2}} \right) \\
&= \frac{1}{2} \left(2t \frac{\tilde{c}}{c} s^2 - 2t \frac{\tilde{c}}{c} s^2 \right) \\
&= 0.
\end{aligned} \tag{3.18}$$

Next, we evaluate $F(q^*)$. From (3.11) and (3.16), we get

$$\begin{aligned}
F(q^*) &= \frac{1}{2} \left(t^2 \left(1 + \frac{\tilde{c}^2}{c^2} - 2\tilde{c}^2 \right) + \frac{t^2}{c^2} (c^2 - \tilde{c}^2) \right) \\
&= \frac{1}{2} \left(t^2 + t^2 \frac{\tilde{c}^2}{c^2} - 2t^2 \tilde{c}^2 + t^2 - t^2 \frac{\tilde{c}^2}{c^2} \right) \\
&= t^2(1 - \tilde{c}^2),
\end{aligned} \tag{3.19}$$

which is the desired result. $\square$

4 Convergence to the best linear predictor

In this section, we consider the setting where the columns of $\mathbf{X}$ are iid random vectors from a distribution in $\mathbb{R}^p$ with mean zero. We will write each column of $\mathbf{X}$ as $X_j/\sqrt{n}$, where X_j is a random vector of the following form:

$$X_j = \sum_{k=1}^r t_k z_{jk} \mathbf{u}_k, \tag{4.1}$$

where the z_{jk} are mean zero, unit variance sub-Gaussian random variables, and the $\mathbf{u}_k$ are the orthonormal principal components of X_j . The assumption that X_j has mean zero is easily removed by subtracting the sample mean from each X_j .

Remark 4. In the new setting, each t_k is the standard deviation of X_j along the principal component $\mathbf{u}_k$, not the singular value of $\mathbf{X}$; and $\mathbf{u}_1, \dots, \mathbf{u}_r$ are not the left singular vectors of $\mathbf{X}$. However, in the large n limit, the singular values of $\mathbf{X}$ converge almost surely to $t_1, \dots, t_r$, and the left singular vectors of $\mathbf{X}$ almost surely make zero angle with, respectively, $\mathbf{u}_1, \dots, \mathbf{u}_r$.

It is known [14] that the best linear predictor of X_j from Y_j has the following form:

$$\hat{X}_j^{\text{BLP}} = \sum_{k=1}^r \frac{t_k^2}{t_k^2 + 1} \langle Y_j, \mathbf{u}_k \rangle \mathbf{u}_k. \tag{4.2}$$

The next result shows that optimal singular value shrinkage with *any* of the loss functions considered in [8] converges to the best linear predictor when $\gamma = 0$. We will let $\widehat{\mathbf{X}}^{\mathbf{q}^*} = [\widehat{X}_j^{\mathbf{q}^*}, \dots, \widehat{X}_j^{\mathbf{q}^*}]/\sqrt{n}$ and $\widehat{\mathbf{X}}^{\text{BLP}} = [\widehat{X}_j^{\text{BLP}}, \dots, \widehat{X}_j^{\text{BLP}}]/\sqrt{n}$.

Theorem 4.1. *Let $\mathcal{L}_{p,n}^{(\ell)}(\widehat{\mathbf{X}}, \mathbf{X})$ be the Schatten ℓ -loss, and suppose $\mathbf{q}^* = \operatorname{argmin}_{\mathbf{q}} \mathcal{L}_{\infty}^{(\ell)}(\widehat{\mathbf{X}}^{\mathbf{q}}, \mathbf{X})$ is the vector of asymptotically optimal singular values. Then in the limit $n \rightarrow \infty$ and $p/n \rightarrow 0$,*

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{j=1}^n \|\widehat{X}_j^{\mathbf{q}^*} - \widehat{X}_j^{\text{BLP}}\|^2 = \lim_{n \rightarrow \infty} \|\widehat{\mathbf{X}}^{\mathbf{q}^*} - \widehat{\mathbf{X}}^{\text{BLP}}\|_{\text{F}}^2 = 0, \quad (4.3)$$

where the limit holds almost surely.

Theorem 4.1 is a consequence of the following result, which is a special case of Theorems V.2 and V.3 in [13]:

Lemma 4.2. *Let $\widehat{\mathbf{X}}^{\mathbf{q}}$ be any singular value shrinker, with singular values $q_1, \dots, q_r$. Define the linear predictor $\widetilde{X}_j^{\mathbf{q}}$ by:*

$$\widetilde{X}_j^{\mathbf{q}} = \sum_{k=1}^r \frac{q_k}{\sigma_k} \langle Y_j, \mathbf{u}_k \rangle \mathbf{u}_k, \quad (4.4)$$

where $\sigma_1, \dots, \sigma_r$ are given by (2.2). Then

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{j=1}^n \|\widehat{X}_j^{\mathbf{q}} - \widetilde{X}_j^{\mathbf{q}}\|^2 = \lim_{n \rightarrow \infty} \|\widehat{\mathbf{X}}^{\mathbf{q}} - \widetilde{\mathbf{X}}^{\mathbf{q}}\|_{\text{F}}^2 = 0, \quad (4.5)$$

where the limit holds almost surely as $n \rightarrow \infty$ and $p/n \rightarrow 0$.

In particular, if we treat $q_k = q_k(\gamma)$ as a function of γ , and

$$\lim_{\gamma \rightarrow 0} q_k(\gamma) = t_k \sqrt{\frac{t_k^2}{t_k^2 + 1}}, \quad (4.6)$$

then

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{j=1}^n \|\widehat{X}_j^{\mathbf{q}} - X_j^{\text{BLP}}\|^2 = \lim_{n \rightarrow \infty} \|\widehat{\mathbf{X}}^{\mathbf{q}} - \mathbf{X}^{\text{BLP}}\|_{\text{F}}^2 = 0. \quad (4.7)$$

Remark 5. Lemma 4.2 states that singular value shrinkage converges to a linear predictor of the form (4.4) when $\gamma = 0$ (the “classical” regime of zero aspect ratio). So long as (4.6) is satisfied, the limiting linear filter is optimal. By contrast, any shrinker that does not satisfy (4.6), including $q_k = t_k$, will converge to a suboptimal linear filter in the $\gamma = 0$ regime.

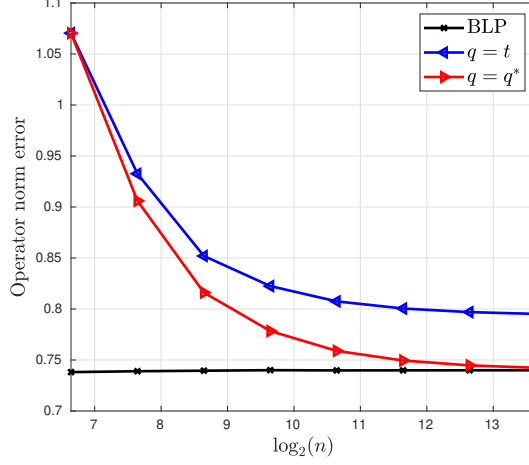


Figure 2: Operator norm errors for $p = 100$ and increasing n for the BLP (4.2) and the shrinkers $\hat{\mathbf{X}}^q$, $q = q^*$ and $q = t$. The signal is rank 1 with $t = 1.1$. Errors are averaged over 4000 runs.

To illustrate Theorem 4.1 and Lemma 4.2 numerically, we draw n iid observations from a spiked model in $\mathbb{R}^{100}$, for increasing values of $n \geq 100$. We take the rank $r = 1$ and $t = t_1 = 1.1$ (to ensure the signal is detectable for all $n \geq 100$). We apply the best linear predictor $\hat{X}_j^{\text{BLP}}$ (which assumes the principal component $\mathbf{u} = \mathbf{u}_1$ is known), optimal shrinkage $\hat{X}_j^{q^*}$, and the suboptimal shrinker $\hat{X}_j^t$. In Figure 2, we plot the average operator norm error over 4000 runs of the experiment, as a function of $\log_2(n)$. The error for the BLP is approximately constant, since the BLP does not vary with the sample size n . As n grows, the error for optimal shrinkage approaches that of the BLP as predicted by Theorem 4.1, because $\hat{X}_j^{q^*}$ converges to the BLP $\hat{X}_j^{\text{BLP}}$. By contrast, the error for the shrinker of [8] converges to a strictly larger value, since according to Lemma 4.2 $\hat{X}_j^t$ converges to the suboptimal linear predictor $\tilde{X}_j^t$.

We now give the proof of Theorem 4.1:

Proof of Theorem 4.1. Since each component is treated separately, we drop the subscripts. When $\gamma = 0$, formulas (2.3) and (2.4) reduce to $c = 1$ and $\tilde{c} = t^2/(t^2 + 1)$. From (2.8), the optimal singular value q^* must then be equal to

$$q^* = \underset{q}{\operatorname{argmin}} \left\| \begin{pmatrix} t & 0 \\ 0 & 0 \end{pmatrix} - q \begin{pmatrix} \tilde{c} & \tilde{s} \\ 0 & 0 \end{pmatrix} \right\|_{S_\ell} = \underset{q}{\operatorname{argmin}} \sqrt{(t - q\tilde{c})^2 + (q\tilde{s})^2} = t\tilde{c} = t\sqrt{\frac{t^2}{t^2 + 1}}. \quad (4.8)$$

Applying Lemma 4.2, the proof of Theorem 4.1 is complete. $\square$

5 Conclusion

We have considered the problem of estimating a low-rank matrix $\mathbf{X}$ from noisy observations $\mathbf{Y} = \mathbf{X} + \mathbf{G}$, where we measure the error by operator norm loss $\|\mathbf{X} - \widehat{\mathbf{X}}\|_{\text{op}}$. We have proven (Theorem 3.1) that the optimal singular value shrinker has singular values of the form (3.1). For square matrices ($\gamma = 1$), the optimal singular values agree with those proposed in [8], though the two methods differ when γ differs from 1. We have also shown (Theorem 4.1) that when the columns of $\mathbf{X}$ are iid vectors in $\mathbb{R}^p$, then in the classical regime ($\gamma = 0$) the optimal shrinker for any Schatten loss, including operator norm loss, converges to the best linear predictor.

Acknowledgements

I thank Edgar Dobriban for helpful feedback on an earlier version of this manuscript. I acknowledge support from the NSF BIGDATA program IIS 1837992 and BSF award 2018230.

References

- [1] Florent Benaych-Georges and Raj Rao Nadakuditi. The singular values and vectors of low rank perturbations of large rectangular random matrices. *Journal of Multivariate Analysis*, 111:120–135, 2012.
- [2] Jérémie Bigot, Charles Deledalle, and Delphine Féral. Generalized SURE for optimal shrinkage of singular values in low-rank matrix denoising. *Journal of Machine Learning Research*, 18:1–50, 2017.
- [3] Sourav Chatterjee. Matrix estimation by universal singular value thresholding. *The Annals of Statistics*, 43(1):177–214, 2015.
- [4] Edgar Dobriban, William Leeb, and Amit Singer. Optimal prediction in the linearly transformed spiked model. *Annals of Statistics*, 48(1):491–513, 2020.
- [5] David L. Donoho, Matan Gavish, and Iain M. Johnstone. Optimal shrinkage of eigenvalues in the spiked covariance model. *Annals of Statistics*, 46(4):1742–1778, 2018.
- [6] Matan Gavish and David L. Donoho. Minimax risk of matrix denoising by singular value thresholding. *The Annals of Statistics*, 42(6):2413–2440, 2014.
- [7] Matan Gavish and David L. Donoho. The optimal hard threshold for singular values is $4/\sqrt{3}$. *IEEE Transactions on Information Theory*, 60(8):5040–5053, 2014.

- [8] Matan Gavish and David L. Donoho. Optimal shrinkage of singular values. *IEEE Transactions on Information Theory*, 63(4):2137–2152, 2017.
- [9] Iain M Johnstone. On the distribution of the largest eigenvalue in principal components analysis. *Annals of Statistics*, 29(2):295–327, 2001.
- [10] Julie Josse and Sylvain Sardy. Adaptive shrinkage of singular values. *Statistics and Computing*, 26:715724, 2016.
- [11] Julie Josse and Stefan Wager. Bootstrap-based regularization for low-rank matrix estimation. *Journal of Machine Learning Research*, 17:1–29, 2016.
- [12] William Leeb. Matrix denoising for weighted loss functions and heterogeneous signals. *SIAM Journal on Mathematics of Data Science*, 3(3):987–1012, 2021.
- [13] William Leeb and Elad Romanov. Optimal spectral shrinkage and PCA with heteroscedastic noise. *IEEE Transactions on Information Theory*, 67(5):3009–3037, 2021.
- [14] D. J. C. MacKay. Deconvolution. In *Information Theory, Inference and Learning Algorithms*, pages 550–551. Cambridge University Press, Cambridge, UK, 2004.
- [15] Raj Rao Nadakuditi. OptShrink: An algorithm for improved low-rank signal matrix denoising by optimal, data-driven singular value shrinkage. *IEEE Transactions on Information Theory*, 60(5):3002–3018, 2014.
- [16] Debashis Paul. Asymptotics of sample eigenstructure for a large dimensional spiked covariance model. *Statistica Sinica*, 17(4):1617–1642, 2007.
- [17] Andrey A. Shabalin and Andrew B. Nobel. Reconstruction of a low-rank matrix in the presence of Gaussian noise. *Journal of Multivariate Analysis*, 118:67–76, 2013.