

ROBUST SPARSE COVARIANCE ESTIMATION BY THRESHOLDING TYLER’S M-ESTIMATOR

BY JOHN GOES^{1,*}, GILAD LERMAN^{1,**} AND BOAZ NADLER²

¹*School of Mathematics, University of Minnesota, *johngoes@umn.edu; **lerman@umn.edu*

²*Department of Computer Science and Applied Mathematics, Weizmann Institute of Science, boaz.nadler@weizmann.ac.il*

Estimating a high-dimensional sparse covariance matrix from a limited number of samples is a fundamental task in contemporary data analysis. Most proposals to date, however, are not robust to outliers or heavy tails. Toward bridging this gap, in this work we consider estimating a sparse shape matrix from n samples following a possibly heavy-tailed elliptical distribution. We propose estimators based on thresholding either Tyler’s M-estimator or its regularized variant. We prove that in the joint limit as the dimension p and the sample size n tend to infinity with $p/n \rightarrow \gamma > 0$, our estimators are minimax rate optimal. Results on simulated data support our theoretical analysis.

1. Introduction. The covariance matrix Σ of a p -dimensional random variable X is a central object in statistical data analysis. Given n samples $\{x_i\}_{i=1}^n$, an accurate estimate of Σ is needed in many tasks including PCA, clustering and discriminant analysis (Anderson (2003), Mardia, Kent and Bibby (1979)). The sample covariance matrix, which is the standard estimator for Σ , is quite accurate when X is sub-Gaussian and $p \ll n$.

In several contemporary applications, however, the number of samples n and the dimension p are comparable, and the data may be heavy tailed. To accurately estimate the covariance matrix when n and p are comparable, additional assumptions, such as its approximate sparsity are typically made. Over the past decade, several sparse covariance matrix estimators were proposed and analyzed (Bickel and Levina (2008), Cai and Liu (2011), El Karoui (2008), Lam and Fan (2009), Rothman, Levina and Zhu (2009)). In addition, minimax rates for estimating high-dimensional sparse covariance matrices were established (Cai and Zhou (2012a, 2012b), Cai, Ren and Zhou (2016)).

With respect to heavy-tailed data, a popular model which we consider in this work is the elliptical distribution (Cambanis, Huang and Simons (1981), Fang, Kotz and Ng (1990), Frahm (2004), Kelker (1970)). An elliptical distribution is characterized by a $p \times p$ shape or scatter matrix S_p , which is proportional to the population covariance matrix, when the latter exists. When the elliptical distribution is heavy tailed, the sample covariance is a poor estimate of the population covariance (Falk (2002)). Moreover, the elliptical distribution might be so heavy tailed as to not even have finite second moments, in which case its population covariance does not exist. Yet due to the structure of the elliptical distribution, even with heavy tails it is nonetheless possible to accurately estimate its shape matrix. This is useful in various applications, since the shape matrix preserves the directional properties of the distribution, such as its principal components.

Following Huber’s pioneering work (Huber and Ronchetti (2009)), over the past decades several robust estimators of the covariance and shape matrix were proposed and theoretically studied; see Dümbgen, Nordhausen and Schuhmacher (2016), Dümbgen, Pauly and Schweizer (2015), Kent and Tyler (1991), Maronna (1976), Maronna and Yohai (2017) and

Received June 2017; revised November 2018.

MSC2010 subject classifications. Primary 62H12; secondary 62G35, 62G20.

Key words and phrases. Covariance matrix estimation, spectral norm, elliptical distribution, sparsity, Tyler’s M-estimator, thresholding.

references therein. For elliptical distributions, [Tyler \(1987a\)](#) proposed a robust M-estimator for the scatter matrix S_p and an iterative scheme to compute it. Tyler’s M-estimator has found widespread use in various applications involving heavy-tailed data. However, as it is defined only for $p < n$, in recent years several regularized variants, applicable also for $p > n$ were proposed and analyzed ([Abramovich and Spencer \(2007\)](#), [Chen, Wiesel and Hero \(2011\)](#), [Ollila and Tyler \(2014\)](#), [Pascal, Chitour and Quek \(2014\)](#), [Sun, Babu and Palomar \(2014\)](#), [Wiesel \(2012\)](#)). The spectral properties of Maronna’s M-estimators and specifically Tyler’s M-estimator and its regularized variants, in high dimensions as $n, p \rightarrow \infty$ with $p/n \rightarrow \gamma$ were studied by [Dümbgen \(1998\)](#), [Couillet, Pascal and Silverstein \(2014, 2015\)](#), [Couillet, Kammoun and Pascal \(2016\)](#), [Couillet and McKay \(2014\)](#), [Zhang, Cheng and Singer \(2016\)](#), among others. For a recent survey on Tyler’s M-estimator and its variants, see [Wiesel and Zhang \(2014\)](#).

In this paper, we study the combination of heavy-tailed data with a “large p –large n ” setting. As formulated in [Section 2](#), we consider robust estimation of the shape matrix of an elliptical distribution, assuming it is approximately sparse. We address the following two challenges: (i) design a computationally efficient and statistically accurate estimator of the shape matrix S_p , that is adaptive to its unknown sparsity parameters; (ii) provide theoretical guarantees on its accuracy, in the large p large n regime.

We make the following contributions. First, in [Section 3](#) we propose simple and computationally efficient estimators for the sparse shape matrix of an elliptical distribution. These are based on thresholding either Tyler’s M-estimator (TME) or its regularized variant. Second, we provide theoretical guarantees on their accuracy in the limit $n, p \rightarrow \infty$ with $p/n \rightarrow \gamma$. [Theorems 1 and 2](#) show that the estimator $\hat{E}$ based on thresholding either TME for $\gamma < 1$ or its regularized variant for any $\gamma \in (0, \infty)$, converges in spectral norm to a sparse shape matrix S_p at rate $\|\hat{E} - S_p\| = O_P((\log p/n)^{(1-q)/2})$, where q is the sparsity parameter of S_p . Estimating a sparse shape matrix under a heavy-tailed elliptical distribution is thus possible with the same asymptotic error rate as estimating a sparse covariance matrix under sub-Gaussian distributions. Moreover, our estimators are rate optimal, as this rate coincides with the minimax rate for sparse covariance estimation with sub-Gaussian data ([Cai and Zhou \(2012a\)](#)).¹

Our proofs follow the approach of [Bickel and Levina \(2008\)](#), with required modifications given that we analyze Tyler’s M-estimators. [Theorem 1](#), which concerns the TME and is thus valid for $p < n$, is proven in [Section 5](#). The proof is relatively simple and heavily relies on [Zhang, Cheng and Singer \(2016\)](#), who studied the spectral properties of Tyler’s M-estimator when $n, p \rightarrow \infty$. [Theorem 2](#) considers the thresholded regularized TME, and is thus applicable also for $p > n$. As detailed in [Section 6](#), its proof is far more involved, and combines a careful analysis of the regularized TME with several results in random matrix theory. [Section 7](#) presents simulation results that support our theoretical analysis. With an eye toward practitioners, given that regularization is common also when $p < n$, we focus on the regularized TME. With Gaussian data, our thresholded TME estimator is as accurate as thresholding the sample covariance. In contrast, in the presence of heavy tails it is far more accurate. We also illustrate its potential utility in handling outliers. In addition, our estimator is quite fast to compute in practice, requiring only few seconds on a standard PC, say for $p = n = 1000$.

Our work is related to several recent papers, that also considered sparse shape or covariance matrix estimation with heavy-tailed data. [Han, Lu and Liu \(2014\)](#) considered a generalization of the elliptical distribution, denoted as the pair-elliptical distribution, with moderate

¹Technically, the minimax rate was proven under the assumption that $p/n^\beta \rightarrow c$ with $\beta > 1$; see [Remark 5](#) in [Cai and Zhou \(2012a\)](#). However, from personal communication with Professors Cai and Zhou, the same minimax rate should hold also when $\beta = 1$.

tails so the population covariance matrix exists. They proposed an estimator for the population covariance and derived finite sample approximation bounds, which depend on various properties of the distribution. For well-behaved elliptical distributions with an exactly sparse covariance matrix, their estimator is minimax rate optimal under the Frobenius norm. Soloveychik and Wiesel (2014) considered estimating a covariance matrix from a convex subset of all positive semidefinite matrices. They added a convex regularization term to the TME and solved the resulting optimization problem by a semidefinite program (SDP). They proved the existence of their estimator and its asymptotic consistency for fixed dimension p and $n \rightarrow \infty$. However, their SDP-based method is computationally demanding even for moderate values of n and p . Sun, Babu and Palomar (2016) considered a wider nonconvex class of matrices, and derived an SDP-based algorithm with lower time complexity.

Chen, Gao and Ren (2018) considered an elliptical distribution, corrupted by an ϵ -contamination model. They proposed several estimators for the shape matrix of the elliptical distribution, based on a generalization of Tukey's depth function. Under a notion of sparsity different from the one considered here, they proved their estimator is minimax rate optimal when $n, p \rightarrow \infty$ and $(\log p)/n \rightarrow 0$. However, from a practical perspective this depth function estimator has a significant limitation—it is intractable to compute. Balakrishnan et al. (2017) considered an ϵ -contamination model for a Gaussian distribution with sparse covariance matrix Σ , such that $\|\Sigma - I\|_0 \leq s$ for a fixed $s \geq 0$. They proposed a polynomial-time algorithm to robustly estimate Σ under this model and established an upper bound on its error under Frobenius norm, assuming $n, p \rightarrow \infty$ and $(\log p)/n \rightarrow c \geq 0$. Our work in contrast provides a computationally efficient and rate optimal estimator for an approximately sparse shape matrix of a potentially heavy-tailed elliptical distribution in the high-dimensional setting $p, n \rightarrow \infty$ with $p/n \rightarrow \gamma$. Finally, Avella-Medina et al. (2018) developed rate optimal robust sparse covariance estimators for heavy-tailed distributions via a different approach than the one presented here, based on various robust pilot estimators. Further discussion and directions for future research appear in Section 8.

2. Problem setting. With precise definitions below, given n i.i.d. observations from an elliptical distribution, the problem we study is how to estimate its $p \times p$ shape matrix S_p . Of particular interest to us is the high-dimensional regime, where both p, n are large and comparable. Following previous works, to be able to accurately estimate the shape matrix in this regime we assume that it is approximately sparse. For completeness, we first introduce some notation, briefly review the elliptical distribution and the class of approximately sparse shape matrices we consider.

Notation. We denote vectors by bold lowercase letters as in $\mathbf{v}$, and matrices by bold uppercase letters as in $\mathbf{A}$. For a vector $\mathbf{v} \in \mathbb{R}^n$, $\|\mathbf{v}\|$ is its Euclidean norm, $\|\mathbf{v}\|_\infty = \max_i |v_i|$, and $B_R(\mathbf{u}) = \{\mathbf{v} \in \mathbb{R}^n \mid \|\mathbf{v} - \mathbf{u}\|_\infty \leq R\}$. The unit sphere in $\mathbb{R}^p$ is denoted $\mathbb{S}^{p-1}$. The identity matrix is $\mathbf{I}$ and $\mathbf{0}$ and $\mathbf{1}$ are the vectors of zeros and ones respectively, with dimensions clear from the context. For a matrix $\mathbf{A} = (a_{ij})$, $\|\mathbf{A}\|$ denotes its spectral norm, $\|\mathbf{A}\|_F$ its Frobenius norm, $\|\mathbf{A}\|_{\max} = \max_{i,j} |a_{ij}|$ and $\|\mathbf{A}\|_\infty = \max_i \sum_j |a_{ij}|$. We denote the set of $p \times p$ symmetric positive semidefinite and definite matrices by S_+^p and S_{++}^p , respectively. We say that an event occurs with high probability (abbreviated w.h.p.), if its probability is at least $1 - C \exp(-cp)$ for constants $c, C > 0$ independent of p .

Elliptical distribution and its shape matrix. A random vector $\mathbf{x} \in \mathbb{R}^p$ follows an elliptical distribution with location vector $\boldsymbol{\mu}$ if it has the form

$$(1) \quad \mathbf{x} = \boldsymbol{\mu} + u S_p^{\frac{1}{2}} \boldsymbol{\xi} = \boldsymbol{\mu} + u \mathbf{z},$$

where $\boldsymbol{\xi}$ is drawn uniformly from $\mathbb{S}^{p-1}$, $S_p \in S_{++}^p$ and u is an arbitrary random or deterministic nonzero scalar, independent of $\boldsymbol{\xi}$.

In equation (1), S_p is not unique, as it can be arbitrarily scaled with u absorbing the inverse scaling factor. Without loss of generality, we thus fix

$$\text{tr}(S_p) = p,$$

and refer to S_p as the *shape matrix*. This normalization is natural in the sense that the mean variance of the p coordinates of $\mathbf{z}$ is one. If the distribution is elliptical and the population covariance Σ exists, then $\Sigma = cS_p$ for some constant $c > 0$; see, for example, [Soloveychik and Wiesel \(2014\)](#).

An important property of the elliptical distribution is that if $\mathbf{x}_1, \mathbf{x}_2$ are independent random vectors from (1), then $\mathbf{x}_1 - \mathbf{x}_2$ has an elliptical distribution with the same shape matrix S_p but with a zero location vector $\boldsymbol{\mu} = \mathbf{0}$. When the goal is to estimate the shape matrix S_p , this allows to remove the typically unknown location vector by a symmetrization principle ([Dümbgen \(1998\)](#)). Specifically, all $\mathbf{x}_i - \mathbf{x}_j$ are elliptically distributed with location vector $\boldsymbol{\mu} = \mathbf{0}$, and one may estimate the shape matrix using all of these pairwise differences ([Dümbgen \(1998\)](#), [Sirkiä, Taskinen and Oja \(2007\)](#)). As discussed by [Nordhausen and Tyler \(2015\)](#), such a procedure is beneficial also for nonelliptical distributions. The resulting $O(n^2)$ pairs are, however, dependent which may complicate the analysis of the resulting estimator. For simplicity, we shall thus assume to have initially observed $2n$ i.i.d. samples $\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{2n}$ from model (1) and in what follows consider the n differences $\mathbf{x}_i = \tilde{\mathbf{x}}_{2i} - \tilde{\mathbf{x}}_{2i-1}$ for $i = 1, \dots, n$ which form an i.i.d. sample from the elliptical distribution (1) with location vector $\boldsymbol{\mu} = \mathbf{0}$.

Approximate sparsity of the shape matrix. Following [Bickel and Levina \(2008\)](#), we consider the following class of row/column approximately sparse shape matrices with fixed parameters $0 \leq q \leq 1$, $M > 0$ and $s_p > 0$:

$$\mathcal{U}(q, s_p, M) = \left\{ \mathbf{A} \in S_{++}^p : a_{ii} \leq M, \sum_{j=1}^p |a_{ij}|^q \leq s_p, 1 \leq i \leq p \right\}.$$

Problem statement. Let $\{\mathbf{x}_i\}_{i=1}^n$ be n i.i.d. samples from the model (1) with location vector $\boldsymbol{\mu} = \mathbf{0}$ and a sparse shape matrix $S_p \in \mathcal{U}(q, s_p, M)$. We consider the following two problems: (i) without explicit knowledge of q, s_p and M , design a computationally efficient and statistically accurate estimator of the shape matrix S_p ; (ii) provide theoretical guarantees on its accuracy, in the asymptotic limit as $p, n \rightarrow \infty$ with $p/n \rightarrow \gamma \in (0, \infty)$.

3. Sparse shape matrix estimation. If the elliptical distribution is sub-Gaussian, then thresholding the sample covariance matrix, proposed by [Bickel and Levina \(2008\)](#) and [El Karoui \(2008\)](#), yields an accurate estimate of S_p up to a multiplicative scaling. As illustrated in Section 7, however, in the presence of heavy tails, thresholding the sample covariance may give a poor estimate of the shape matrix.

To handle heavy tails, we propose the following approach: compute Tyler's M-estimator (TME) or its regularized variant, and threshold it. In Section 3.1, we review TME and its regularized variant. We prove that computing the latter is computationally efficient. Section 3.2 presents our proposed estimators. A theoretical analysis of their accuracy appears in Section 3.3.

3.1. TME and its regularized variant. TME, proposed by [Tyler \(1987a\)](#) for elliptical distributions with a known location vector, which w.l.o.g. is assumed to be $\mathbf{0}$, is a $p \times p$ matrix $\hat{\Sigma}$ which satisfies the nonlinear equation

$$(2) \quad \frac{p}{n} \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\mathbf{x}_i^T \hat{\Sigma}^{-1} \mathbf{x}_i} = \hat{\Sigma}.$$

Here, samples $\mathbf{x}_i$ lying at the origin are ignored as they provide no information on the scatter matrix, and n is the number of samples not at the origin. As solutions to (2) can be multiplied by an arbitrary constant, Tyler (1987a) considered the normalization $\text{tr}(\hat{\Sigma}) = p$, and suggested to solve equation (2) by the following iterations, starting from an arbitrary $\hat{\Sigma}_1 \in S_{++}^p$:

$$\hat{\Sigma}_{k+1} = p \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\mathbf{x}_i^T \hat{\Sigma}_k^{-1} \mathbf{x}_i} / \text{tr} \left(\sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\mathbf{x}_i^T \hat{\Sigma}_k^{-1} \mathbf{x}_i} \right).$$

Kent and Tyler (1988), Theorems 1 and 2, showed that if any linear subspace in $\mathbb{R}^p$ of dimension $1 \leq d \leq p-1$ contains less than nd/p of the data samples, then there exists a unique solution to equation (2), and the above iterations converge to it. With n i.i.d. observations from an elliptical distribution, no samples lie at the origin and this condition holds with probability 1.

TME enjoys several important properties: First, it may equivalently be defined as the minimizer of the following cost function, over all positive definite matrices with the constraint $\text{tr}(\mathbf{R}) = p$,

$$(3) \quad L(\mathbf{R}) = \frac{p}{n} \sum_{i=1}^n \log(\mathbf{x}_i^T \mathbf{R}^{-1} \mathbf{x}_i) + \log(\det(\mathbf{R})).$$

Equation (3) implies that $\hat{\Sigma}$ is the maximum likelihood estimator of the shape matrix of the angular central Gaussian distribution (Tyler (1987b)) and of the generalized elliptical distribution (Frahm and Jaekel (2010)). Second, for data i.i.d. from a continuous elliptical distribution with p fixed, it is the “most robust” estimator of the shape matrix as $n \rightarrow \infty$ (Tyler (1987a), Remark 3.1). TME outperforms the sample covariance in a variety of applications, including finance (Frahm and Jaekel (2007)), anomaly detection in wireless sensor networks (Chen, Wiesel and Hero (2011)), antenna array processing (Ollila and Koivunen (2003)) and radar detection (Ollila and Tyler (2012)).

As the TME does not exist when $p > n$, several regularized variants have been proposed and analyzed (Abramovich and Spencer (2007), Chen, Wiesel and Hero (2011), Pascal, Chitour and Quek (2014), Sun, Babu and Palomar (2014), Wiesel (2012)). Even when $p \leq n$, it is common to add small regularization to the TME. Following Sun, Babu and Palomar (2014), we consider the following regularized TME $\hat{\Sigma}(\alpha)$, defined as the solution of

$$(4) \quad \hat{\Sigma}(\alpha) = \frac{1}{1+\alpha} \frac{p}{n} \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\mathbf{x}_i^T \hat{\Sigma}(\alpha)^{-1} \mathbf{x}_i} + \frac{\alpha}{1+\alpha} \mathbf{I},$$

where α is a regularization parameter. If $\alpha = 0$, equation (4) reverts to equation (2). While regularization toward general target matrices is possible (Wiesel (2012)), here for simplicity we consider only regularization toward the identity. In contrast to the TME of equation (2), where solutions can be multiplied by an arbitrary positive scalar, any solution to equation (4) satisfies $\text{tr}(\hat{\Sigma}(\alpha)^{-1}) = p$, for any value of α (Pascal, Chitour and Quek (2014), Proposition III.1).

Sun, Babu and Palomar (2014), Theorem 11 and Proposition 13, derived a sufficient and necessary condition for existence of a unique positive definite matrix which solves equation (4). Again, ignoring samples at the origin, the condition is that any linear subspace in $\mathbb{R}^p$ of dimension $1 \leq d \leq p-1$ contains less than $(1+\alpha)nd/p$ of the samples. Since $\alpha > 0$, this condition is weaker than for the original TME. In particular, with data i.i.d. from a continuous distribution, equation (4) has a unique solution for $\alpha > \max(0, p/n - 1)$; see also Pascal, Chitour and Quek (2014), Theorem III.1. With n i.i.d. samples from an elliptical distribution, these conditions hold with probability 1.

Sun, Babu and Palomar (2014), Proposition 18, further showed that starting from any positive definite initial guess, the following iterations

$$(5) \quad \hat{\Sigma}_{k+1}(\alpha) = \frac{1}{1+\alpha} \frac{p}{n} \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\mathbf{x}_i^T \hat{\Sigma}_k(\alpha)^{-1} \mathbf{x}_i} + \frac{\alpha}{1+\alpha} \mathbf{I}$$

converge to the unique solution. Various properties of TME and its regularized variant, in the limit as $p, n \rightarrow \infty$ with $p/n \rightarrow \gamma$, were proven by Couillet, Kammoun and Pascal (2016), Couillet and McKay (2014), Dümbgen (1998), Zhang, Cheng and Singer (2016).

The following lemma, proven in the Appendix, shows that if α is sufficiently large and $\hat{\Sigma}(\alpha)$ exists, then the iterations (5), starting with $\hat{\Sigma}_1(\alpha) = \alpha \mathbf{I} / (1 + \alpha)$, have a uniform linear convergence rate from the first iteration. As far as we know, this result is new and is of independent interest.

To state the lemma, let $e_k = \|\hat{\Sigma}(\alpha) - \hat{\Sigma}_k(\alpha)\|$ be the error after k iterations, $\tilde{X}$ be the $p \times n$ matrix whose columns are $\{\mathbf{x}_i / \|\mathbf{x}_i\|\}_{i=1}^n$ and let

$$C(\tilde{X}) = \frac{p}{n} \|\tilde{X} \tilde{X}^T\| = \frac{p}{n} \left\| \sum_{i=1}^n \frac{\mathbf{x}_i \mathbf{x}_i^T}{\|\mathbf{x}_i\|^2} \right\|.$$

Note that for a given data set, $C(\tilde{X})$ is fixed and can be computed a priori.

LEMMA 1. *Let $\{\mathbf{x}_i\}_{i=1}^n$ be a data set in $\mathbb{R}^p$ with constant $C(\tilde{X})$ and let $0 < R < 1$. Suppose that $\alpha > \max((3 + R^{-1})C(\tilde{X}) - 1, 0)$ and let $\hat{\Sigma}(\alpha)$ be a solution of (4). Then the iterations of equations (5), starting from $\hat{\Sigma}_1(\alpha) = \frac{\alpha}{1+\alpha} \mathbf{I}$, uniformly and linearly converge to $\hat{\Sigma}(\alpha)$ with the ratio R . That is,*

$$(6) \quad e_{k+1} \leq R e_k \leq R^k e_1 \quad \text{for all } k \geq 1.$$

A straightforward calculation yields the bound $C(\tilde{X}) \geq p/n$. Hence, the above assumptions on α imply that $\alpha > \max(0, p/n - 1)$, and consequently guarantee the existence and uniqueness of $\hat{\Sigma}(\alpha)$ in our setting.

Lemma 1 implies that for sufficiently large α , calculating $\hat{\Sigma}(\alpha)$ is computationally efficient. Specifically, for accuracy ϵ with convergence ratio R , at most $\lceil \log_{R^{-1}}(\epsilon^{-1}) \rceil$ iterations are needed. If $n > p$, then at each iteration, the matrix inversion costs $O(p^3)$ operations and the other operations are $O(np^2)$. For $n < p$, one may first perform an SVD of the data and calculate $\hat{\Sigma}(\alpha)$ restricted to the subspace $W = \text{Span}(\mathbf{x}_i)$ whose dimension is at most n . This suffices, since for any $\mathbf{v} \perp W$, by definition $\hat{\Sigma}(\alpha)\mathbf{v} = \alpha/(1 + \alpha)\mathbf{v}$. Each iteration then costs at most $O(n^3)$. Therefore, the total cost of computing $\hat{\Sigma}(\alpha)$ within accuracy ϵ is $O(\log(\epsilon^{-1})(n + p) \min(n, p)^2)$.

Our theoretical analysis below studies the regularized TME as $p, n \rightarrow \infty$ and $p/n \rightarrow \gamma \in (0, \infty)$, but with a fixed value of α . The next lemma shows that for data sampled from an elliptical distribution, with high probability $C(\tilde{X})$ is bounded by a constant that depends on $\|\mathbf{S}_p\|$ and on the ratio p/n .

LEMMA 2. *Let $\mathbf{x}_1, \dots, \mathbf{x}_n$ be i.i.d. from equation (1) with $\boldsymbol{\mu} = \mathbf{0}$ and shape matrix $\mathbf{S}_p$. Then, with probability $> 1 - \exp(-cp)$, where $c = c(\|\mathbf{S}_p\|) > 0$,*

$$(7) \quad C(\tilde{X}) \leq 2\|\mathbf{S}_p\|(1 + 2\sqrt{p/n})^2.$$

3.2. TME-based thresholding estimators. One possible approach to construct a sparse and robust estimator for the shape matrix is to add a suitable penalty to the original cost functional equation (3) of the TME. For various structural assumptions on the shape matrix, this was proposed by [Soloveychik and Wiesel \(2014\)](#) and by [Sun, Babu and Palomar \(2016\)](#).

With a sparsity inducing penalty, however, this approach in general leads to a nonconvex and potentially difficult to optimize objective. Instead, we opt for thresholding the (regularized) TME, which as discussed above, can be computed efficiently in practical polynomial time.

For a matrix $\mathbf{A} = (a_{ij})$ and threshold $t > 0$, define the entry-wise hard-thresholding operator by

$$\tau_t(\mathbf{A}) = (\mathbf{1}(|a_{ij}| > t)a_{ij}).$$

For $n > p$, where the TME $\hat{\Sigma}$ exists and by definition has unit trace, our proposed estimator for the shape matrix $\mathbf{S}_p$ takes the form

$$(8) \quad \hat{\mathbf{S}}_p = \tau_t(\hat{\Sigma}),$$

where the threshold $t = t(p, n)$ is specified below. Similarly, for general p, n , our estimator based on the regularized TME is

$$(9) \quad \hat{\mathbf{S}}_p = \tau_t \left(p \frac{\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I}}{\text{tr}(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I})} \right).$$

Note that both $\hat{\Sigma}$ in equation (8) and the argument matrix prior to thresholding in equation (9) have rank at most $\min(n, p)$.

3.3. Accuracy of the thresholded TME. The following Theorems 1 and 2, proved in Sections 5 and 6, respectively, establish the asymptotic accuracy of equations (8) and (9) as estimates of the shape matrix $\mathbf{S}_p$.

THEOREM 1. Consider a sequence $(n, p, \mathbf{S}_p)$ where $n \rightarrow \infty$, $p = p_n \rightarrow \infty$ with $p/n \rightarrow \gamma \in (0, 1)$, and $\mathbf{S}_p \in \mathcal{U}(q, s_p, M)$. For each triplet $(n, p, \mathbf{S}_p)$, let $\hat{\Sigma}$ be the TME of n i.i.d. samples $\{\mathbf{x}_i\}_{i=1}^n \subset \mathbb{R}^p$ from the elliptical distribution (1). Then there exists a constant M' depending only on γ such that for any fixed $M'' > M'$, the thresholded TME of equation (8) with threshold $t_n = M'' \sqrt{\log p/n}$, approaches $\mathbf{S}_p$ in spectral norm at a rate

$$\|\tau_{t_n}(p\hat{\Sigma}) - \mathbf{S}_p\| = \mathcal{O}_P \left(s_p \cdot \left(\frac{\log p}{n} \right)^{(1-q)/2} \right).$$

THEOREM 2. Consider a sequence $(n, p, \mathbf{S}_p)$ as in Theorem 1, here with $p/n \rightarrow \gamma \in (0, \infty)$ and with the additional assumption that $\|\mathbf{S}_p\| \leq s_{\max}$. For $\alpha > \max(0, \gamma - 1 + s_{\max}(1 + \sqrt{\gamma})^2)$, let $\hat{\Sigma}(\alpha)$ be the regularized TME of n i.i.d. samples $\{\mathbf{x}_i\}_{i=1}^n \subset \mathbb{R}^p$ from the elliptical distribution (1). Then there exists a M' depending only on γ and α such that for any fixed $M'' > M'$, the estimator of equation (9) with $t_n = M'' \sqrt{\log p/n}$, converges in spectral norm to $\mathbf{S}_p$ at rate

$$\left\| \tau_{t_n} \left(p \frac{\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I}}{\text{tr}(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I})} \right) - \mathbf{S}_p \right\| = \mathcal{O}_P \left(s_p \left(\frac{\log p}{n} \right)^{(1-q)/2} \right).$$

Several remarks regarding Theorem 2 are in place.

REMARK 1. As noted by [Bickel and Levina \(2008\)](#), page 2580, if $S_p \in \mathcal{U}(q, s_p, M)$ then $\|S_p\| \leq M^{1-q} s_p$ which may grow with p . Since we analyze the regularized TME with a *fixed* value of α , we explicitly require that $\|S_p\| \leq s_{\max}$ independent of p . If $\|S_p\|$ grows to infinity with p , then the regularization α should also grow to infinity with p , such that $\alpha > c\|S_p\|$ for some constant $c > 0$. While beyond the scope of this paper, we believe an analogue of Theorem 2 should hold in this case.

REMARK 2. The convergence rate in Theorems 1 and 2 coincides with the minimax rate for sparse covariance estimation with sub-Gaussian data, derived by [Cai and Zhou \(2012a\)](#). Since the Gaussian distribution is a particular case of an elliptical distribution, our estimators are thus minimax rate optimal. Furthermore, in light of Lemmas 1 and 2, computing the regularized TME and subsequently thresholding it, is computationally efficient.

REMARK 3. Several authors proposed the regularized Tyler's M-estimator as is (without modification), as an estimate of the shape matrix. In this case, choosing the value of the regularization constant is crucial ([Chen, Wiesel and Hero \(2011\)](#), [Couillet and McKay \(2014\)](#)). Setting the regularization parameter is also important to maximize the detection probability in various signal processing applications ([Kammoun et al. \(2018\)](#)). In contrast, as we subtract the regularization $\alpha/(1+\alpha)\mathbf{I}$ prior to thresholding, at least asymptotically, the precise value of α is unimportant, provided it is sufficiently large. This is also evident in the simulations in Section 7. From a practical perspective, we thus suggest to use a value of α as described in Lemma 1, with say $R = 1/2$, which is not only sufficient for existence but also guarantees fast convergence of the iterations that compute the regularized TME.

4. Preliminaries. In proving Theorems 1 and 2, we shall make frequent use of the following auxiliary lemmas. The first is a simple inequality. Let A, B be nonnegative random variables. Then for any $c > 0$ and $\lambda > 0$,

$$(10) \quad \Pr(AB > c) \leq \Pr(A > \lambda c) + \Pr(B > 1/\lambda).$$

Next is the following well-known result, which states that TME and regularized TME are unable to distinguish an elliptical distribution from a Gaussian one. Its proof (omitted) follows directly from the fact that (regularized) TME for data $\mathbf{x}_i$ is identical to that of data $t_i \mathbf{x}_i$, where t_i are arbitrary positive real valued numbers.

LEMMA 3. *TME or regularized TME with $\alpha > \max(0, p/n - 1)$ under an elliptical distribution with zero location vector and shape matrix S_p has the same distribution as under a zero mean Gaussian distribution with covariance S_p .*

The following two results from random matrix theory, concerning the spectral norm of a Wishart matrix and concentration of quadratic forms will also be of use; see, for example, [Davidson and Szarek \(2001\)](#), Theorem 2.13 and [Rudelson and Vershynin \(2013\)](#), Theorem 1.1.

LEMMA 4. *Let $\{\xi_i\}_{i=1}^n \subset \mathbb{R}^p$ be i.i.d. $N(0, \mathbf{I})$, and let $T_n = \frac{1}{n} \sum_i \xi_i \xi_i^T$. Then, $\mathbb{E}[\|T_n\|] \leq (1 + \sqrt{p/n})^2$ and*

$$\Pr(\|T_n\| > (1 + \sqrt{p/n} + t)^2) \leq \exp(-nt^2/2).$$

LEMMA 5. *Let $A \in \mathbb{R}^{p \times p}$ and $\xi \sim N(0, \mathbf{I})$. Then there exist absolute constants $c_1, c_2 > 0$ such that for all $\epsilon > 0$,*

$$\Pr(|\xi^T A \xi - \text{tr}(A)| > \epsilon) \leq 2 \exp\left(-c_1 \min\left\{\frac{c_2^2 \epsilon^2}{\|A\|_F^2}, \frac{c_2 \epsilon}{\|A\|}\right\}\right).$$

Finally, the following auxiliary lemma, proved in the supplemental article (Goes, Lerman and Nadler (2020)), is a slight modification of a result by Bickel and Levina (2008), page 2583.

LEMMA 6. Assume $\mathbf{B} \in \mathcal{U}(q, s_p, M)$. Let $\mathbf{A}$ be a matrix such that

$$\|\mathbf{A} - \mathbf{B}\|_{\max} \leq C_1 \sqrt{\log p/n},$$

for some $C_1 > 0$. Suppose we threshold $\mathbf{A}$ at level $t = K \sqrt{\log p/n}$, with $K > C_1$. Then there exists a constant $C_2 = C_2(C_1, K, q) < \infty$ such that

$$\|\tau_t(\mathbf{A}) - \mathbf{B}\| \leq C_2 s_p (\log p/n)^{(1-q)/2}.$$

5. Proof of Theorem 1. Let $\hat{\mathbf{S}}$ be the sample covariance of $\{\mathbf{x}_i\}_{i=1}^n$. In light of Lemma 3, we may assume that $\mathbf{x}_i$ are all i.i.d. $N(\mathbf{0}, \mathbf{S}_p)$. The proof proceeds in three steps: (i) reducing to a bound on $\|\hat{\mathbf{\Sigma}} - \hat{\mathbf{S}}\|_{\max}$; (ii) expressing $\hat{\mathbf{\Sigma}}$ as a weighted covariance matrix whose weights are all uniformly close to a constant, with high probability; and (iii) bounding $\|\hat{\mathbf{\Sigma}} - \hat{\mathbf{S}}\|_{\max}$.

5.1. *Step 1: Reduction from $\|\tau_{t_n}(\hat{\mathbf{\Sigma}}) - \mathbf{S}_p\|$ to $\|\hat{\mathbf{\Sigma}} - \hat{\mathbf{S}}\|_{\max}$.* By Lemma 6, it suffices to prove that $\|\hat{\mathbf{\Sigma}} - \mathbf{S}_p\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n})$. By the triangle inequality,

$$\|\hat{\mathbf{\Sigma}} - \mathbf{S}_p\|_{\max} \leq \|\hat{\mathbf{\Sigma}} - \hat{\mathbf{S}}\|_{\max} + \|\hat{\mathbf{S}} - \mathbf{S}_p\|_{\max}.$$

Since the proof of Theorem 1 of Bickel and Levina (2008) shows that

$$\|\hat{\mathbf{S}} - \mathbf{S}_p\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n})$$

it thus suffices to show that

$$(11) \quad \|\hat{\mathbf{\Sigma}} - \hat{\mathbf{S}}\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n}).$$

5.2. *Step 2: The weights of TME.* By Zhang, Cheng and Singer (2016), Lemma 2.1, TME can be written as a weighted covariance matrix,

$$\hat{\mathbf{\Sigma}} = p \sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^T / \text{tr} \left(\sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^T \right),$$

where the weights w_i are the unique solution of

$$(12) \quad \arg \min_{w_i > 0, \sum w_i = 1} - \sum_{i=1}^n \ln w_i + \frac{n}{p} \ln \det \left(\sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^T \right).$$

This characterization is important because of the following result.

LEMMA 7. Consider a sequence $(n, p, \mathbf{S}_p)$ where $n, p \rightarrow \infty$ with $p/n \rightarrow \gamma \in (0, 1)$, and $\mathbf{S}_p \in \mathcal{S}_{++}^p$. For every triplet $(n, p, \mathbf{S}_p)$, let $\mathbf{x}_i \stackrel{i.i.d.}{\sim} N(\mathbf{0}, \mathbf{S}_p)$ and let $\{w_i\}_{i=1}^n$ be the corresponding weights of equation (12). Then there exist positive constants C, c and c' depending only on γ such that for any $0 < \epsilon < c'$, and sufficiently large n ,

$$(13) \quad \Pr \left[\max_i |n w_i - 1| \geq \epsilon \right] \leq C n e^{-c \epsilon^2 n}.$$

The case $\mathbf{S}_p = \mathbf{I}$ was proved by Zhang, Cheng and Singer (2016), Lemma 2.2. Its generalization to an arbitrary $\mathbf{S}_p \in \mathcal{S}_{++}^p$ is proved in the supplemental article (Goes, Lerman and Nadler (2020)).

5.3. *Step 3: Bounding $\|\hat{\Sigma} - \hat{S}\|_{\max}$.* The proof concludes by applying the following lemma which establishes equation (11). Its proof is in Appendix A.2.

LEMMA 8. *Let $\hat{\Sigma}$ and $\hat{S}$ be the TME and the sample covariance matrix of $\mathbf{x}_1, \dots, \mathbf{x}_n$ i.i.d. from $N(\mathbf{0}, \mathbf{S}_p)$, where $\mathbf{S}_p \in \mathcal{U}(q, s_p, M)$ with $\text{tr}(\mathbf{S}_p) = p$. Assume that $p, n \rightarrow \infty$, with $p/n \rightarrow \gamma \in (0, 1)$. Then there exist positive constants C, c and c' that depend only on γ , such that for all $\epsilon \in (0, c')$ and n sufficiently large*

$$\Pr(\|\hat{\Sigma} - \hat{S}\|_{\max} \geq \epsilon) \leq Cne^{-c\epsilon^2 n}.$$

6. Proof of Theorem 2. We first introduce and prove a slightly modified version of Theorem 2. We then show how Theorem 2 follows from it. The modified theorem uses the following proposition, proved in Appendix A.3.

PROPOSITION 1. *Let $\mathbf{y}, \xi_1, \dots, \xi_{n-1} \in \mathbb{R}^p$ be i.i.d. $N(\mathbf{0}, \mathbf{I})$ and denote*

$$Q = Q(r) = \frac{1}{p} \mathbf{y}^T \left(\frac{1}{n} \sum_{j=1}^{n-1} \xi_j \xi_j^T + \alpha \frac{n}{p} \frac{1}{r} \mathbf{S}_p^{-1} \right)^{-1} \mathbf{y},$$

where $\mathbf{S}_p \in S_{++}^p$ with $\|\mathbf{S}_p\| \leq s_{\max}$. Assume that $\alpha > \max(0, p/n - 1 + s_{\max}(1 + \sqrt{p/n})^2)$, and define

$$r_{\min} = \frac{n}{p} \frac{\alpha}{1 + \alpha - p/n}, \quad r_{\max} = \frac{n}{p} \frac{\alpha}{1 + \alpha - p/n - s_{\max}(1 + \sqrt{p/n})^2}.$$

Then there exists a unique $r = r(p, n, \alpha, \mathbf{S}_p) \in [r_{\min}, r_{\max}]$, such that

$$(14) \quad \mathbb{E}[Q(r)] = \frac{1}{1 + \alpha - p/n},$$

where the expectation is over $\mathbf{y}$ and $\xi_1, \dots, \xi_{n-1}$.

6.1. *A reformulation of the main result.* We now introduce the modified theorem.

THEOREM 3. *Consider the setting of Theorem 2. Then there exists an M' depending only on γ and α such that for any fixed $M'' > M'$, the estimator $\tau_{t_n}(\hat{\Sigma}(\alpha) - \alpha \mathbf{I} / (1 + \alpha))$ with $t_n = M'' \sqrt{\frac{\log p}{n}}$, converges in spectral norm to a multiple of $\mathbf{S}_p$,*

$$\left\| \tau_{t_n} \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1 + \alpha} \mathbf{I} \right) - \frac{p}{n} \frac{r}{1 + \alpha} \mathbf{S}_p \right\| = \mathcal{O}_P \left(s_p \left(\frac{\log p}{n} \right)^{(1-q)/2} \right),$$

where the scalar $r = r(p, n, \alpha, \mathbf{S}_p)$ is specified in Proposition 1.

6.2. *Proof of Theorem 3.* By Lemma 3, we may assume $\mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} N(\mathbf{0}, \mathbf{S}_p)$. Following the argument in Section 5.1, combining Lemma 6 with the fact that by Proposition 1, $r < r_{\max}$, it suffices to show that

$$(15) \quad \left\| \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1 + \alpha} \mathbf{I} \right) - \frac{p}{n} \frac{r}{1 + \alpha} \hat{S} \right\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n}).$$

To establish equation (15), we first express $\hat{\Sigma}(\alpha)$ as the sum of $\frac{\alpha}{1 + \alpha} \mathbf{I}$ and weighted $\mathbf{x}_i \mathbf{x}_i^T$ terms, where the weights are the root of some equation. Next we show that this root is concentrated near the vector $r \mathbf{1}/n$, with r specified in Proposition 1.

Following the definition of the regularized TME, we write $\hat{\Sigma}(\alpha)$ as

$$(16) \quad \hat{\Sigma}(\alpha) = \frac{1}{1+\alpha} \frac{p}{n} \sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^T + \frac{\alpha}{1+\alpha} \mathbf{I},$$

where the weight vector $\mathbf{w} = (w_1, \dots, w_n)^T$ satisfies

$$(17) \quad w_i = \frac{1}{\mathbf{x}_i^T \hat{\Sigma}(\alpha)^{-1} \mathbf{x}_i} = \frac{1}{\mathbf{x}_i^T \left(\frac{1}{1+\alpha} \frac{p}{n} \sum_{j=1}^n w_j \mathbf{x}_j \mathbf{x}_j^T + \frac{\alpha}{1+\alpha} \mathbf{I} \right)^{-1} \mathbf{x}_i}.$$

Next consider the function $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ whose n components are

$$(18) \quad g(\mathbf{v})_i = v_i - \frac{1}{\mathbf{x}_i^T \left(\frac{1}{1+\alpha} \frac{p}{n} \sum_{k=1}^n v_k \mathbf{x}_k \mathbf{x}_k^T + \frac{\alpha}{1+\alpha} n \mathbf{I} \right)^{-1} \mathbf{x}_i}.$$

Comparing equation (18) to equation (17), since $\hat{\Sigma}(\alpha)$ is unique (Sun, Babu and Palomar (2014), Theorem 11), then the n nonlinear equations $g(\mathbf{v}) = \mathbf{0}$ have a unique solution, which is thus $n\mathbf{w}$. The next three lemmas state properties of g used to prove that as $p, n \rightarrow \infty$, with $p/n \rightarrow \gamma$, this root concentrates around $\mathbf{u} = r\mathbf{1}$, with r given in Proposition 1. Lemma 9 is proven in Appendix A.4 and the other two in the supplemental article (Goes, Lerman and Nadler (2020)). All three lemmas assume the setting of Theorem 3, and their generic constants depend only on γ, α and $s_{\max}$. Our analysis of the weights w_i follows the pioneering works of Couillet, Pascal and Silverstein (2014, 2015), who proved that the weights in Maronna's M-estimators converge to suitable constants, and Zhang, Cheng and Singer (2016), who derived concentration results for the weights of Tyler's M-estimator as $p, n \rightarrow \infty$ with $p/n \rightarrow \gamma < 1$.

LEMMA 9. *There exist $C, c > 0$ such that for any $\epsilon \in (0, 1)$*

$$\Pr(\|g(\mathbf{u})\|_{\infty} > \epsilon) < Cpe^{-cpe^2}.$$

LEMMA 10. *There exist $c', c_L, C, c > 0$ such that*

$$\Pr(\exists \mathbf{v} \in B_{c'}(\mathbf{u}), \|\nabla g(\mathbf{v}) - \nabla g(\mathbf{u})\|_{\max} > c_L \|\mathbf{v} - \mathbf{u}\|_{\infty}) < Cp^2 e^{-cp}.$$

LEMMA 11. *There exist $c_H, C, c > 0$ such that*

$$(19) \quad \Pr(\|(\nabla g(\mathbf{u}))^{-1}\|_{\infty} > c_H) < Cpe^{-cp}.$$

Lemmas 9 and 10 show that w.h.p. $g(\mathbf{u})$ is small and ∇g is Lipschitz near $\mathbf{u}$. These two properties are consistent with the root of g being close to $\mathbf{u}$. To rigorously prove this, following Zhang, Cheng and Singer (2016), we consider the function $f(\mathbf{v}) = (\nabla g(\mathbf{u}))^{-1} g(\mathbf{v})$. Lemma 11 shows that the matrix $(\nabla g(\mathbf{u}))^{-1}$ is w.h.p. not extremely large. Finally, the following lemma combines these properties of g to infer that its root is close to $\mathbf{u}$.

LEMMA 12. *Let $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $\mathbf{u} \in \mathbb{R}^n$ and $C > 0$. Assume that:*

1. $\nabla f(\mathbf{u}) = \mathbf{I}$;
2. $\|\nabla f(\mathbf{v}) - \nabla f(\mathbf{u})\|_{\max} \leq C \|\mathbf{v} - \mathbf{u}\|_{\infty}$ for all $\|\mathbf{v} - \mathbf{u}\|_{\infty} \leq 3\|f(\mathbf{u})\|_{\infty}$;
3. $\|f(\mathbf{u})\|_{\infty} < \min\{1/(9C), 1/3\}$.

Then there exists a $\tilde{\mathbf{v}} \in \mathbb{R}^n$ such that $f(\tilde{\mathbf{v}}) = \mathbf{0}$ and $\|\tilde{\mathbf{v}} - \mathbf{u}\|_{\infty} < 3\|f(\mathbf{u})\|_{\infty}$.

Lemma 12 is slightly stronger than Lemma 3.1 of Zhang, Cheng and Singer (2016), as it has a weaker requirement that the Lipschitz condition in Lemma 12 holds in a smaller ball $\|\mathbf{v} - \mathbf{u}\|_\infty \leq 3\|f(\mathbf{u})\|_\infty$, instead of the original requirement $\|\mathbf{v} - \mathbf{u}\|_\infty < 1$ in their Lemma 3.1. A careful inspection shows that their original proof is still valid under this weaker assumption.

To apply Lemma 12 to $f(\mathbf{v}) = (\nabla g(\mathbf{u}))^{-1}g(\mathbf{v})$, we verify that the three conditions of the lemma hold with high probability. The first condition is trivially satisfied. For the other two conditions, by Lemmas 9 and 11, w.h.p.

$$\|f(\mathbf{u})\|_\infty \leq \|(\nabla g(\mathbf{u}))^{-1}\|_\infty \cdot \|g(\mathbf{u})\|_\infty \leq c_H \epsilon.$$

Similarly, by Lemmas 10 and 11, for all $\|\mathbf{v} - \mathbf{u}\|_\infty \leq c'$, w.h.p.

$$\|\nabla f(\mathbf{v}) - \nabla f(\mathbf{u})\|_{\max} \leq \|(\nabla g(\mathbf{u}))^{-1}\|_\infty \cdot \|\nabla g(\mathbf{v}) - \nabla g(\mathbf{u})\|_{\max} \leq c_H c_L \|\mathbf{v} - \mathbf{u}\|_\infty.$$

Since for sufficiently small ϵ , $c_H \epsilon < \min\{1/(9c_L c_H), 1/3\}$, both the second and third conditions of Lemma 12 are thus satisfied with constant $C = c_L c_H$.

To conclude, with probability at least $1 - Cp^2 e^{-c p \epsilon^2}$ all three conditions of Lemma 12 hold, so there exists $\tilde{\mathbf{v}} \in \mathbb{R}^n$ such that $f(\tilde{\mathbf{v}}) = \mathbf{0}$ and $\|\tilde{\mathbf{v}} - \mathbf{u}\|_\infty \leq 3\|f(\mathbf{u})\|_\infty < 3c_H \epsilon$. Since $n\mathbf{w}$ is the unique root of $g(\mathbf{v})$ and also of $f(\mathbf{v})$,

$$(20) \quad \Pr(\|n\mathbf{w} - r\mathbf{1}\|_\infty > 3c_H \epsilon) < Cp^2 e^{-c p \epsilon^2}.$$

Next we use equation (20) to bound the LHS of equation (15). First, by equation (16),

$$\begin{aligned} & \left\| \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right) - \frac{1}{1+\alpha} \frac{p}{n} r \hat{\mathbf{S}} \right\|_{\max} \\ &= \frac{1}{1+\alpha} \frac{p}{n} \left\| \sum_{i=1}^n w_i \mathbf{x}_i \mathbf{x}_i^T - r \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right\|_{\max} \\ &\leq \frac{1}{1+\alpha} \frac{p}{n} \|n\mathbf{w} - r\mathbf{1}\|_\infty \left\| \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right\|_{\max}. \end{aligned}$$

Using this inequality and equation (10) with $\lambda = 1/[s_{\max}(1 + 2\sqrt{\gamma})^2]$,

$$\begin{aligned} & \Pr \left(\left\| \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right) - \frac{1}{1+\alpha} \frac{p}{n} r \hat{\mathbf{S}} \right\| > \epsilon \right) \\ &\leq \Pr \left(\frac{1}{1+\alpha} \frac{p}{n} \|n\mathbf{w} - r\mathbf{1}\|_\infty \|\hat{\mathbf{S}}\| > \epsilon \right) \\ &\leq \Pr \left(\|n\mathbf{w} - r\mathbf{1}\|_\infty > \epsilon \frac{n(1+\alpha)}{p s_{\max}(1 + 2\sqrt{\gamma})^2} \right) \\ &\quad + \Pr(\|\hat{\mathbf{S}}\| > s_{\max}(1 + 2\sqrt{\gamma})^2). \end{aligned}$$

Since $\hat{\mathbf{S}} = \mathbf{S}_p^{1/2} (\frac{1}{n} \sum_i \xi_i \xi_i^T) \mathbf{S}_p^{1/2}$ with $\xi_i \sim N(0, \mathbf{I})$, by Lemma 4 the second term is exponentially small in p . By equation (20), the first term is bounded by $C' p^2 e^{-c p \epsilon^2}$. Hence, equation (15) holds, which concludes the proof of Theorem 3.

6.3. *Concluding the proof of Theorem 2.* Similar to Theorems 1 and 3, to prove Theorem 2 it suffices to show that

$$(21) \quad \left\| p \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right) / \text{tr} \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right) - \hat{\mathbf{S}} \right\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n}).$$

Equation (15) combined with Proposition 1 imply that for $r > r_{\min} > 0$,

$$(22) \quad \left\| \frac{n}{p} \frac{1+\alpha}{r} \left(\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right) - \hat{\mathbf{S}} \right\|_{\max} = \mathcal{O}_P(\sqrt{\log p/n}).$$

Since $\text{tr}(\hat{\mathbf{S}})$ is tightly concentrated around p , we may replace $\hat{\mathbf{S}}$ in equation (22) by $p\hat{\mathbf{S}}/\text{tr}(\hat{\mathbf{S}})$. Equation (21) follows by the following lemma, proven in the supplemental article (Goes, Lerman and Nadler (2020)), combined with the fact that w.h.p. $\|\hat{\mathbf{S}}\|_{\max} \leq 2\|\mathbf{S}_p\|_{\max} \leq 2M$.

LEMMA 13. *Let $\mathbf{B} \in S_+^p$ with $\text{tr}(\mathbf{B}) = p$ and $\|\mathbf{B}\|_{\max} \leq b_{\max}$. Suppose that $\mathbf{A} \in S_+^p$ satisfies $\|\mathbf{A} - \mathbf{B}\|_{\max} < \epsilon \leq 1/2$. Then*

$$(23) \quad \left\| \frac{p\mathbf{A}}{\text{tr}(\mathbf{A})} - \mathbf{B} \right\|_{\max} \leq 2(1 + b_{\max})\epsilon.$$

7. Numerical experiments. Focusing on the regularized TME, we present simulations that support our theoretical analysis. Section 7.1 compares the regularized TME, the sample covariance and their thresholded versions. Section 7.2 considers the sensitivity of the proposed estimator to α . Section 7.3 demonstrates a modified estimator to cope with outliers.

7.1. Comparison of thresholded TME with covariance estimators. We considered the following shape matrix, also used by Bickel and Levina (2008):

$$\mathbf{S}_p = (s_{ij}) = (0.7^{|i-j|}).$$

Note that excluding the diagonal all rows of this matrix have ℓ_1 norm bounded by $2/(1 - 0.7) - 2 = 14/3$. Hence, by the Gershgorin disk theorem, for any p , this matrix has a finite spectral norm, $\|\mathbf{S}_p\| \leq s_{\max} = 1 + 14/3 = 17/3$. This is in accordance with our assumptions in Theorem 2.

We generated data from a Gaussian scale mixture, which is a particular case of equation (1). Here, u and ξ are independent, with $\xi \sim N(\mathbf{0}, \mathbf{I})$. We considered three different choices for the random variables u_i : (i) $u_i = 1$, so $\{\mathbf{x}_i\}_{i=1}^n$ are i.i.d. $N(\mathbf{0}, \mathbf{S}_p)$; (ii) $u_i \sim \text{Laplace}(0, 1)$, a heavy-tailed distribution with finite moments; and (iii) $u_i \sim \text{Cauchy}(0, 1)$, so the distribution does not even have a well-defined mean or covariance.

We computed four estimators for the shape matrix: (i) SampCov: the sample covariance scaled to have trace p , $p\hat{\mathbf{S}}/\text{tr}(\hat{\mathbf{S}})$; (ii) th-SampCov: the thresholded version of SampCov, $\tau_t(p\hat{\mathbf{S}}/\text{tr}(\hat{\mathbf{S}}))$; (iii) RegTME: the regularized TME, normalized to have trace p ,

$$\frac{p(\Sigma(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I})}{\text{tr}(\Sigma(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I})},$$

and (iv) th-RegTME: the thresholded version of RegTME in equation (9). We choose $\alpha = 10$, and threshold at level $t = \sqrt{(\log p)/n}$. Our stopping rule for (5) is $\|p\hat{\Sigma}_{k+1}/\text{tr}(\hat{\Sigma}_{k+1}) - p\hat{\Sigma}_k/\text{tr}(\hat{\Sigma}_k)\|_F < 10^{-12}$, or $k = 1400$ iterations.

We measured the accuracy of an estimator $\hat{\mathbf{S}}_p$ by the logarithm of its averaged relative error (LRE). That is, for 100 different realizations, we independently generated n i.i.d. samples in $\mathbb{R}^p$, and each time estimated $(\hat{\mathbf{S}}_p)_i$, where $i = 1, \dots, 100$. The LRE was then computed as follows:

$$\text{LRE} = \log \left(\frac{1}{100} \sum_{i=1}^{100} \frac{\|(\hat{\mathbf{S}}_p)_i - \mathbf{S}_p\|}{\|\mathbf{S}_p\|} \right).$$

We considered sample sizes $n \in [100, 1000]$ and the following three ratios $p/n \in \{0.5, 1, 2\}$. Figure 1 shows the LRE of the four estimators. As expected theoretically, for $u_i \equiv 1$ thresholding the sample covariance or the regularized TME yield similar errors. In contrast, for

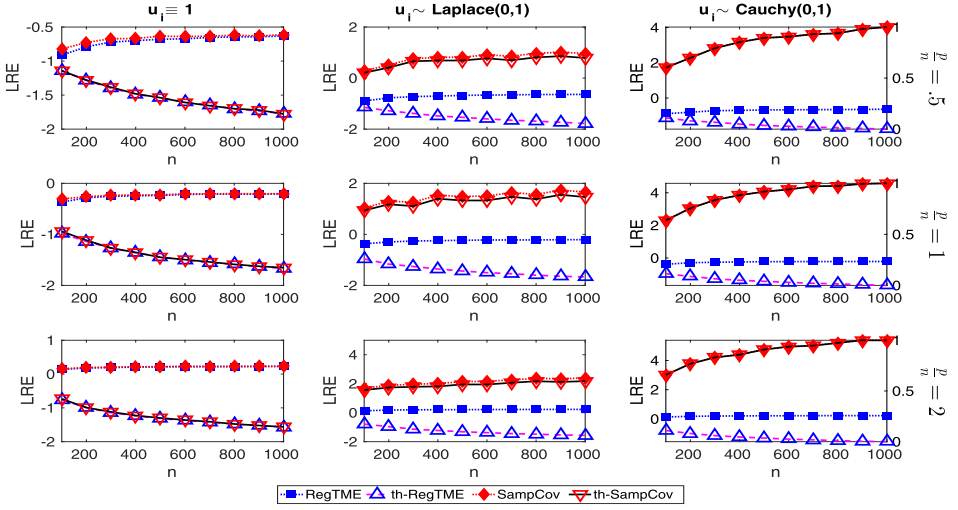


FIG. 1. Comparison of the LRE of the four estimators with data i.i.d. from a Gaussian scale mixture. The rows correspond to $p/n = 0.5, 1, 2$. The columns correspond to $u_i \equiv 1, u_i \stackrel{i.i.d.}{\sim} \text{Laplace}(0, 1)$ and $u_i \stackrel{i.i.d.}{\sim} \text{Cauchy}(0, 1)$.

heavy-tailed data the thresholded sample covariance performs poorly, whereas the thresholded regularized TME is still an accurate estimate of S_p . Note that since the regularized TME is invariant to the scaling u_i , the resulting errors of the regularized TME and its thresholded version (the blue squares and triangles) are the same for all three distributions of u_i .

7.2. Sensitivity of regularized TME to choice of α . Next we study how the error and runtime of th-RegTME depend on the regularization parameter α . We consider the Gaussian model with covariance S_p , and explore the behavior of th-RegTME for the following values of α : 0.2, 0.4, 0.6, 0.8, 1, 2, 3, ..., 20 and the following three cases: $(p, n) = (800, 400)$, $(p, n) = (800, 200)$ and $(p, n) = (400, 200)$. Even though the regularized TME does not exist if $\alpha < \max(0, p/n - 1)$, our algorithm, with a stopping criterion based on scaled matrices converged for all considered values of α . For a similar property upon scaling scatter matrices, see [Chen, Wiesel and Hero \(2011\)](#). The left panel of Figure 2 shows the LRE of th-RegTME as a function of α . The maximal LRE occurs at $p/n - 1$ and larger values of α yield slightly smaller errors, which are nearly identical for all large values of α . This is in accordance with Theorem 2, which states that asymptotically all large values of α yield the same error rate. The right panel of Figure 2 displays the logarithm of the runtime of th-RegTME as a function of α , showing a sharp increase in runtime as $p/n - 1$ approaches α .

Next we explore the behavior of th-RegTME for $p = 480$, $\alpha = 1, 2, 3, 4$ and $n = 60, 64, 68, \dots, 300$. The left panel of Figure 3 shows the error of th-RegTME as a func-

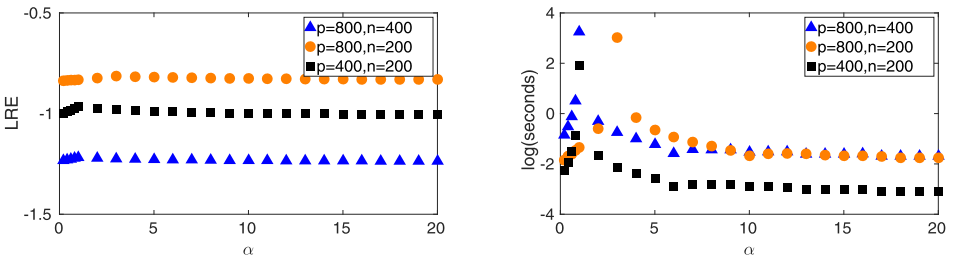


FIG. 2. LRE and log-runtime of th-RegTME on elliptical data for different choices of α and three choices of p and n .

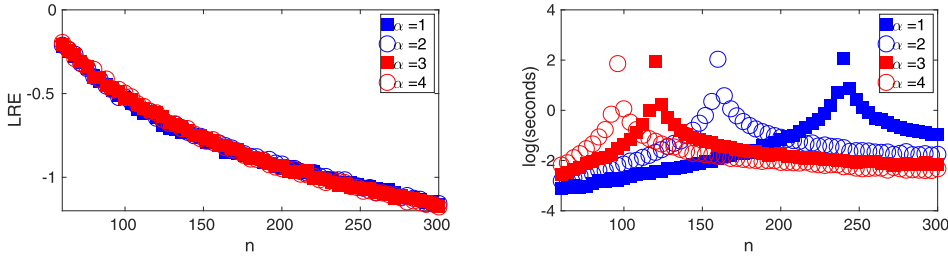


FIG. 3. *LRE and log-runtime of th-RegTME versus number of samples n , at $p = 480$ and $\alpha = 1, 2, 3, 4$.*

tion of n . Again, in accordance with theory, α has little effect on the accuracy. Of particular interest is the runtime, seen in the right panel of Figure 3. Here, we see a sharp increase in runtime as $p/n - 1$ approaches α . For $n \geq \frac{p}{\alpha+1}$, the runtime decreases as α increases.

These experiments indicate that one may generally prefer larger α , particularly for faster runtime. We propose to choose a value of α close to the bound in Lemma 1 for some $R \in (0, 1)$, which guarantees fast convergence.

7.3. Regularized TME in the presence of outliers. We conclude the numerical section with an illustrative example of the ability of the regularized TME to detect outliers, and upon their removal and thresholding, to provide a robust and accurate estimate of a sparse shape matrix. For a related rigorous study on the ability of Maronna’s M-estimator to detect outliers, see [Morales-Jimenez, Couillet and McKay \(2015\)](#).

To this end, we consider the following ϵ -contamination mixture model: $(1 - \epsilon)n$ of the observed data, the inliers, follow an elliptical distribution with the same sparse shape matrix S_p as above. The remaining ϵn of the samples, the outliers, follow an elliptical distribution with shape matrix $U(pD/\text{tr}(D))U'$, where U is a unitary matrix, uniformly distributed with Haar measure, and D is a diagonal matrix. In our first experiment, the diagonal entries d_{ii} are all i.i.d. uniformly distributed over $[1, 5]$, so the outliers are rather diffuse. In our second experiment $d_{11} = p$, $d_{22} = p/2$ and all other $d_{ii} = 1$, so the outliers are nearly on a 2-d randomly rotated subspace.

Given n samples from this ϵ -contamination model, and without knowledge of ϵ , the task is to accurately estimate the shape matrix S_p . Since both the inliers and outliers have potentially heavy-tailed distributions, it might not be possible to detect the outliers by simple schemes, such as those based on the norm of a sample or the number of its neighbors in a given radius. However, recall that by our theoretical analysis, in the absence of outliers ($\epsilon = 0$), the corresponding weights w_i in the regularized TME are all approximately equal. For $\epsilon \ll 1$, with all samples normalized to have unit norm, we thus expect the inliers to still all have similar weights, and the outliers to have quite different weights, hopefully smaller though not necessarily so. With further details in Appendix A.5, our proposed procedure for robustness to outliers is to estimate the mean and standard deviation of the inliers’ weights. Then exclude all samples whose weights are outside, say, the mean plus or minus two standard deviations, recompute the regularized TME on the remaining samples and threshold it.

Figure 4 illustrates the robustness of this procedure to outliers in two different settings. From left to right, for $\epsilon = 0.2$ and $\epsilon = 0.4$, it shows the weights of the n normalized samples $\mathbf{x}_i/\|\mathbf{x}_i\|$, sorted so the first ϵn of them are the outliers. The blue horizontal line is a robust estimate of the mean weight of the inliers, and the two red lines are this estimated mean plus and minus two standard deviations. The top row corresponds to the first outlier model with $d_{ii} \sim U[1, 5]$. The second row corresponds to our second outlier model with $D = \text{diag}(p, p/2, 1, \dots, 1)$. Note that this outlier shape matrix has a spectral norm $O(p)$, which does not satisfy our requirement that $\|D\| \leq s_{\max}$. As indeed observed empirically, the

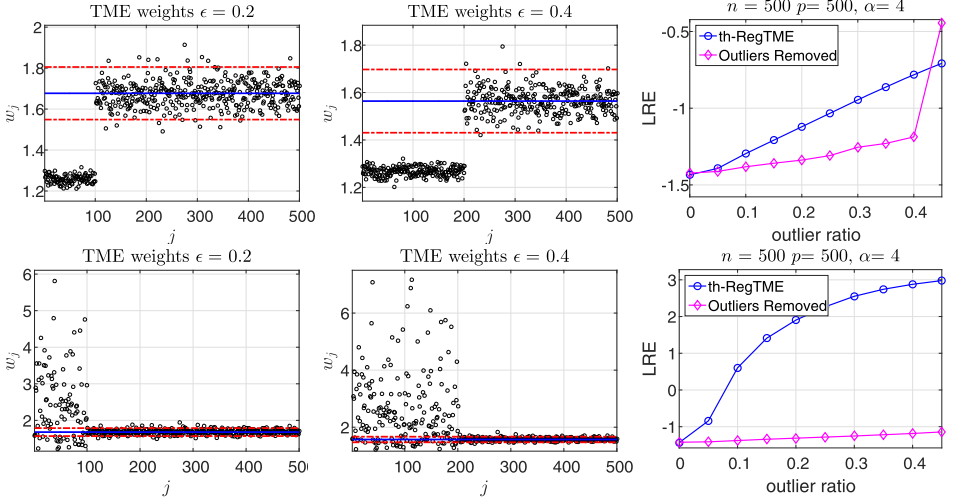


FIG. 4. The TME weights for $\epsilon = 0.2, 0.4$ and the log relative error (LRE) of thresholding the regularized TME, before and after outlier removal, versus ϵ . Top row $D_{ii} \sim U[1, 5]$. Bottom row $D = \text{diag}(p, p/2, 1, \dots, 1)$.

weights of the outliers do not so tightly concentrate around some value. Yet, our outlier exclusion procedure still succeeds to exclude most of these outliers. The error of the thresholded TME with outliers removed, compared to that of thresholding the original TME is shown in the right column of Figure 4.

This simple example illustrates the potential ability of TME to screen outliers in high-dimensional settings, at least for small contamination levels. A detailed study of this issue is an interesting topic for future research.

8. Summary and discussion. In this paper, we proposed simple estimators for the shape matrix of possibly heavy-tailed elliptical distributions, assuming the shape matrix is approximately sparse. We further analyzed their error, showing that under the spectral norm they are minimax rate optimal in a high-dimensional setting with $p/n \rightarrow \gamma$.

There are several directions for future research. One direction is to study whether our proposed approach, of thresholding a regularized TME provides accurate estimates of a sparse covariance matrix for other heavy-tailed distributions beyond the elliptical distribution. Another direction is to extend our results to the case $p = n^\beta$, with $\beta > 1$. Our current analysis assumed the regularization parameter α of TME is *fixed*, whereas if $p = n^\beta$ with $\beta > 1$, just to ensure its existence would require $\alpha \rightarrow \infty$. Handling this case thus requires extending our analysis to allow α to grow with n and p .

A question of practical interest is how to set the threshold parameter in a data-driven fashion. [Bickel and Levina \(2008\)](#), Section 3, proposed a cross validation procedure to set the threshold. Rigorously proving that this provides a good estimate in the case of (regularized) TME is an interesting topic for future research.

While our work focused on approximate sparsity of the shape matrix, robust inference under other common assumptions can also be studied. For example, one might assume that the first few leading eigenvectors of Σ are sparse, also known as sparse-PCA, or that Σ is the combination of a low rank and a sparse matrix. In particular, a robust sparse-PCA estimator may be constructed by applying a sparse-PCA procedure to Tyler's M-estimator.

Finally, another direction for future work is to develop a computationally efficient algorithm for sparse covariance estimation in the presence of a small fraction of arbitrary outliers. This setting was considered in [Chen, Gao and Ren \(2018\)](#), but without a computationally

tractable estimator. Our promising preliminary results in Section 7.3 suggest to study whether the regularized TME offers such robustness, and under which outlier models.

APPENDIX A: PROOFS

A.1. Complexity of calculating the regularized TME.

PROOF OF LEMMA 1. We arbitrarily fix a solution $\hat{\Sigma}(\alpha)$ of (4). Since $\hat{\Sigma}(\alpha)$ is invariant to scaling of the data, we assume that $\|x_i\| = 1$, $1 \leq i \leq n$. We first analyze the quantity $e_1 = \|\hat{\Sigma}(\alpha) - \hat{\Sigma}_1(\alpha)\|$. To this end, let $\lambda_{\max} = \|\hat{\Sigma}(\alpha)\|$. Taking the spectral norm in equation (4), together with the fact that $\frac{1}{x^T \hat{\Sigma}(\alpha)^{-1} x} \leq \lambda_{\max}$ for any vector x with $\|x\| = 1$,

$$\lambda_{\max} \leq \frac{1}{1+\alpha} \left\| \frac{p}{n} \sum_{i=1}^n \frac{x_i x_i^T}{x_i^T \Sigma^{-1} x_i} \right\| + \frac{\alpha}{1+\alpha} \leq \frac{1}{1+\alpha} \lambda_{\max} C(\tilde{X}) + \frac{\alpha}{1+\alpha}.$$

Equivalently, for $1+\alpha > C(\tilde{X})$,

$$\lambda_{\max} \leq \frac{\alpha}{1+\alpha - C(\tilde{X})}.$$

Combining this inequality with the fact that by equation (4) $\hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \in S_+^p$,

$$(24) \quad e_1 = \left\| \hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right\| = \lambda_{\max} - \frac{\alpha}{1+\alpha} \leq \frac{\alpha}{1+\alpha} \frac{C(\tilde{X})}{1+\alpha - C(\tilde{X})}.$$

Next we analyze the error e_k . We denote $E_k = \hat{\Sigma}(\alpha) - \hat{\Sigma}_k(\alpha)$ and write

$$\hat{\Sigma}_k(\alpha) = \hat{\Sigma}(\alpha) - E_k = \hat{\Sigma}(\alpha)^{1/2} (\mathbf{I} - \hat{\Sigma}(\alpha)^{-1/2} E_k \hat{\Sigma}(\alpha)^{-1/2}) \hat{\Sigma}(\alpha)^{1/2}.$$

Since $\hat{\Sigma}(\alpha)$ and $\hat{\Sigma}_k(\alpha)$ are invertible, so is $\mathbf{I} - \hat{\Sigma}(\alpha)^{-1/2} E_k \hat{\Sigma}(\alpha)^{-1/2}$. Let $B_k = \mathbf{I} - (\mathbf{I} - \hat{\Sigma}(\alpha)^{-1/2} E_k \hat{\Sigma}(\alpha)^{-1/2})^{-1}$ and $R_k = \hat{\Sigma}(\alpha)^{-1/2} B_k \hat{\Sigma}(\alpha)^{-1/2}$. Then

$$\hat{\Sigma}_k(\alpha)^{-1} = \hat{\Sigma}(\alpha)^{-1/2} (\mathbf{I} - B_k) \hat{\Sigma}(\alpha)^{-1/2} = \hat{\Sigma}(\alpha)^{-1} - R_k.$$

Subtracting equation (5) from equation (4) gives

$$\begin{aligned} E_{k+1} &= \frac{1}{1+\alpha} \frac{p}{n} \sum_i x_i x_i^T \left(\frac{1}{x_i^T \hat{\Sigma}(\alpha)^{-1} x_i} - \frac{1}{x_i^T \hat{\Sigma}(\alpha)^{-1} x_i - x_i^T R_k x_i} \right) \\ &= \frac{1}{1+\alpha} \frac{p}{n} \sum_i \frac{x_i x_i^T}{x_i^T \hat{\Sigma}(\alpha)^{-1} x_i} \left(1 - \frac{1}{1 - \delta_{ki}} \right), \end{aligned}$$

where $\delta_{ki} = x_i^T R_k x_i / x_i^T \hat{\Sigma}(\alpha)^{-1} x_i$.

Let $D_k = \max_{1 \leq i \leq n} |\delta_{ki}| / (1 - \delta_{ki})$. Since all terms $x_i x_i^T / x_i^T \hat{\Sigma}(\alpha)^{-1} x_i$ are positive semidefinite, the above equation implies that

$$\begin{aligned} \|E_{k+1}\| &\leq D_k \left\| \frac{1}{1+\alpha} \frac{p}{n} \sum_i \frac{x_i x_i^T}{x_i^T \hat{\Sigma}(\alpha)^{-1} x_i} \right\| \\ (25) \quad &= D_k \left\| \hat{\Sigma}(\alpha) - \frac{\alpha}{1+\alpha} \mathbf{I} \right\| = D_k e_1. \end{aligned}$$

Equation (24) gives a bound on e_1 . We now bound D_k . Since $\hat{\Sigma}(\alpha) \geq \frac{\alpha}{1+\alpha} \mathbf{I}$,

$$\|\hat{\Sigma}(\alpha)^{-1/2} E_k \hat{\Sigma}(\alpha)^{-1/2}\| \leq \|\hat{\Sigma}(\alpha)^{-1}\| e_k \leq \frac{1+\alpha}{\alpha} e_k.$$

Assume this quantity is strictly smaller than one, then

$$(26) \quad \|\mathbf{B}_k\| = \|\mathbf{I} - (\mathbf{I} - \hat{\Sigma}(\alpha)^{-1/2} \mathbf{E}_k \hat{\Sigma}(\alpha)^{-1/2})^{-1}\| \leq \frac{1+\alpha}{\alpha} \frac{e_k}{1 - \frac{1+\alpha}{\alpha} e_k}.$$

Finally, given the relation between $\mathbf{R}_k$ and $\mathbf{B}_k$,

$$|\delta_{ki}| = \frac{|\mathbf{x}_i^T \mathbf{R}_k \mathbf{x}_i|}{\mathbf{x}_i^T \hat{\Sigma}(\alpha)^{-1} \mathbf{x}_i} = \frac{|(\hat{\Sigma}(\alpha)^{-1/2} \mathbf{x}_i)^T \mathbf{B}_k (\hat{\Sigma}(\alpha)^{-1/2} \mathbf{x}_i)|}{\|\hat{\Sigma}(\alpha)^{-1/2} \mathbf{x}_i\|^2} \leq \|\mathbf{B}_k\|.$$

Thus, assuming that the bound on $\|\mathbf{B}_k\|$ in equation (26) is less than one,

$$(27) \quad D_k = \max_i \frac{|\delta_{ki}|}{1 - \delta_{ki}} \leq \frac{\|\mathbf{B}_k\|}{1 - \|\mathbf{B}_k\|} \leq \frac{1+\alpha}{\alpha} e_k \cdot \frac{1}{1 - 2\frac{1+\alpha}{\alpha} e_k}.$$

Inserting (27) and (24) into (25) yields that

$$(28) \quad \frac{e_{k+1}}{e_k} \leq \frac{C(\tilde{X})}{1 + \alpha - C(\tilde{X})} \frac{1}{1 - 2\frac{1+\alpha}{\alpha} e_k}.$$

For the proof to hold, the bound in (26) needs to be less than one, namely that $\frac{1+\alpha}{\alpha} e_k < 0.5$. For $0 < R < 1$ and $1 + \alpha > (3 + R^{-1})C(\tilde{X})$, equation (24) implies that $\frac{1+\alpha}{\alpha} e_1 < \frac{1}{2+R^{-1}}$ and combining this with equation (28) results in the estimate $e_2/e_1 < R$. Since $R < 1$, induction implies that for $k > 1$, $\frac{1+\alpha}{\alpha} e_k < \frac{1}{2+R^{-1}} < 0.5$, as required, and so equation (6) holds. Since this convergence holds with any solution of (4), this solution thus has to be unique. $\square$

PROOF OF LEMMA 2. Since the regularized TME is invariant to scaling, we may assume that all $u_i \sim \chi_p^2$, and express $\mathbf{x}_i = \mathbf{S}_p^{\frac{1}{2}} \xi_i$, where $\xi_i \sim N(\mathbf{0}, \mathbf{I})$. Let $\mathbf{U} \mathbf{D} \mathbf{U}^T$ be the eigendecomposition of $\mathbf{S}_p$. Redefining $\xi = \mathbf{U} \xi$, then $\|\mathbf{x}_i\|^2 = \xi_i^T \mathbf{D} \xi_i$ and

$$C(\tilde{X}) = \left\| \mathbf{S}_p^{\frac{1}{2}} \left(\frac{1}{n} \sum_{i=1}^n \frac{\xi_i \xi_i^T}{\frac{1}{p} \xi_i^T \mathbf{D} \xi_i} \right) \mathbf{S}_p^{\frac{1}{2}} \right\|.$$

Combining Lemma 5 with a union bound yields

$$\Pr\left(\max_i \left| \frac{1}{p} \xi_i^T \mathbf{D} \xi_i - \frac{1}{p} \text{tr}(\mathbf{D}) \right| > \epsilon\right) < 2n \exp\left(-c_1 \min\left\{ \frac{c_2^2 p^2 \epsilon^2}{\|\mathbf{D}\|_F^2}, \frac{c_2 p \epsilon}{\|\mathbf{D}\|} \right\}\right).$$

Since $\|\mathbf{D}\| = \|\mathbf{S}_p\|$ and $\|\mathbf{D}\|_F^2 \leq p \|\mathbf{S}_p\|^2$, for any fixed ϵ the above probability is exponentially small in p . Taking say $\epsilon = 1/2$ and recalling that $\text{tr}(\mathbf{D}) = p$, gives that with high probability,

$$C(\tilde{X}) \leq 2\|\mathbf{S}_p\| \cdot \left\| \frac{1}{n} \sum_{i=1}^n \xi_i \xi_i^T \right\|.$$

Equation (7) follows since by Lemma 4, w.h.p. $\|\frac{1}{n} \sum_{i=1}^n \xi_i \xi_i^T\| \leq (1 + 2\sqrt{p/n})^2$. $\square$

A.2. Proof of Lemma 8. To prove the lemma, we shall use the following auxiliary result, whose proof appears in the supplemental article (Goes, Lerman and Nadler (2020)).

LEMMA 14. Assume the setting of Lemma 8. There exist constants C, c and $c' < 1$ depending on γ such that $\forall \epsilon \in (0, c')$ and n sufficiently large,

$$(29) \quad \Pr\left(\left| \frac{P}{T_w} - 1 \right| > \epsilon\right) \leq C n e^{-c n \epsilon^2}.$$

PROOF OF LEMMA 8. By definition,

$$(30) \quad \begin{aligned} \|\hat{\Sigma} - \hat{\mathbf{S}}\|_{\max} &= \left\| \sum_{i=1}^n \left(\frac{pw_i}{T_w} - \frac{1}{n} \right) \mathbf{x}_i \mathbf{x}_i^T \right\|_{\max} \\ &\leq \left\| \frac{np\mathbf{w}}{T_w} - \mathbf{1} \right\|_{\infty} \cdot \left\| \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right\|_{\max}. \end{aligned}$$

Since $\mathbf{x}_i \sim N(0, \mathbf{S}_p)$ with $\mathbf{S}_p \in \mathcal{U}(q, s_p, M)$, then w.h.p., $\left\| \frac{1}{n} \sum \mathbf{x}_i \mathbf{x}_i^T \right\|_{\max} \leq 2M$. As for the first term on the RHS of equation (30), by the triangle inequality,

$$\left\| \frac{np\mathbf{w}}{T_w} - \mathbf{1} \right\|_{\infty} = \left\| \frac{np\mathbf{w}}{T_w} - n\mathbf{w} + n\mathbf{w} - \mathbf{1} \right\|_{\infty} \leq \|n\mathbf{w}\|_{\infty} \left| \frac{p}{T_w} - 1 \right| + \|n\mathbf{w} - \mathbf{1}\|_{\infty}.$$

Hence,

$$\Pr\left(\left\| \frac{np\mathbf{w}}{T_w} - \mathbf{1} \right\|_{\infty} > \epsilon\right) \leq \Pr\left(\|n\mathbf{w}\|_{\infty} \left| \frac{p}{T_w} - 1 \right| > \epsilon/2\right) + \Pr(\|n\mathbf{w} - \mathbf{1}\|_{\infty} > \epsilon/2).$$

Lemma 7 provides an exponential bound on the second term. For the first term, applying equation (10) with $\lambda = 2$ gives

$$\begin{aligned} \Pr\left(\|n\mathbf{w}\|_{\infty} \left| \frac{p}{T_w} - 1 \right| > \epsilon/2\right) &\leq \Pr(\|n\mathbf{w}\|_{\infty} > 2) + \Pr\left(\left| \frac{p}{T_w} - 1 \right| > \epsilon/4\right) \\ &\leq \Pr(\|n\mathbf{w} - \mathbf{1}\|_{\infty} > 1) + \Pr\left(\left| \frac{p}{T_w} - 1 \right| > \epsilon/4\right). \end{aligned}$$

By Lemmas 7 and 14, these two probabilities are exponentially small. $\square$

A.3. Proof of Proposition 1. To prove the existence of a unique $r^* = r^*(p, n, \alpha, \mathbf{S}_p)$ which satisfies equation (14), we first show that $\mathbb{E}[Q(r)]$ is strictly monotone increasing in r and then use the intermediate value theorem.

To simplify notation, let $\mathbf{T} = \frac{1}{n} \sum_{j=1}^{n-1} \xi_j \xi_j^T$ and $\beta = \beta(r) = \frac{n\alpha}{pr}$. Then

$$(31) \quad \mathbb{E}[Q(r)] = \mathbb{E}_{\xi_i}[\mathbb{E}_{\mathbf{y}}[Q(r)]] = \mathbb{E}\left[\frac{1}{p} \text{tr}((\mathbf{T} + \beta \mathbf{S}_p^{-1})^{-1})\right],$$

where the expectation is now only over the random variables ξ_i .

First, we show that for any $\mathbf{S}_p \in \mathcal{S}_p^{++}$, $\mathbb{E}[Q(r)]$ is strictly monotone increasing in r . Indeed, differentiating with respect to r and using the identity $\text{tr}(\mathbf{A}\mathbf{B}) = \text{tr}(\mathbf{B}\mathbf{A})$

$$\frac{d}{dr} \mathbb{E}[Q(r)] = \frac{n\alpha}{pr^2} \mathbb{E}\left[\frac{1}{p} \text{tr}((\mathbf{T} + \beta \mathbf{S}_p^{-1})^{-2} \mathbf{S}_p^{-1})\right].$$

Using Bhatia (1997), Problem III.6.14 and Jensen's inequality,

$$\begin{aligned} \frac{d}{dr} \mathbb{E}[Q(r)] &\geq \frac{n\alpha}{pr^2} \lambda_{\min}(\mathbf{S}_p^{-1}) \mathbb{E}\left[\frac{1}{p} \text{tr}((\mathbf{T} + \beta \mathbf{S}_p^{-1})^{-2})\right] \\ &\geq \frac{n\alpha}{pr^2} \frac{1}{s_{\max}} \mathbb{E}[1/\lambda_1(\mathbf{T} + \beta \mathbf{S}_p^{-1})^2] \\ &\geq \frac{n\alpha}{pr^2} \frac{1}{s_{\max}} \frac{1}{\mathbb{E}[\lambda_1(\mathbf{T} + \beta \mathbf{S}_p^{-1})]^2}. \end{aligned}$$

Clearly, $\lambda_1(\mathbf{T} + \beta \mathbf{S}_p^{-1}) \leq \lambda_1(\mathbf{T}) + \beta/\lambda_{\min}(\mathbf{S}_p)$. Furthermore, upon averaging over the random variables ξ_i , by Lemma 4, $\mathbb{E}[\lambda_1(\mathbf{T})] \leq (1 + \sqrt{p/n})^2$. Therefore, for any fixed $\mathbf{S}_p$, the

derivative of $\mathbb{E}[Q(r)]$ is strictly positive for any $r > 0$. Hence if there exists a solution to equation (14), then it must be unique.

Next we show that this solution must satisfy $r \geq r_{\min}$. By definition,

$$\begin{aligned} \mathbb{E}[Q(r)] &= \frac{1}{p} \sum_{j=1}^n \frac{1}{\lambda_j(\mathbf{T} + \beta \mathbf{S}_p^{-1})} \\ (32) \quad &\leq \frac{1}{p} \sum_{j=1}^n \frac{1}{\lambda_j(\beta \mathbf{S}_p^{-1})} = \frac{1}{\beta} \frac{1}{p} \sum_{j=1}^p \lambda_j(\mathbf{S}_p) = \frac{n\alpha}{pr}. \end{aligned}$$

Combining (32) with (14) implies that

$$r^*(p, n, \alpha, \mathbf{S}_p) \geq \frac{n}{p} \frac{\alpha}{1 + \alpha - p/n} = r_{\min}.$$

Finally, we bound r from above. To this end, using Jensen's inequality

$$\begin{aligned} \mathbb{E}[Q(r)] &= \frac{1}{p} \mathbb{E} \left[\sum_j \frac{1}{\lambda_j(\mathbf{T} + \beta \mathbf{S}_p^{-1})} \right] \geq \frac{1}{p} \mathbb{E} \left[\sum_j \frac{1}{\beta/\lambda_j(\mathbf{S}_p) + \|\mathbf{T}\|} \right] \\ &\geq \mathbb{E} \left[\frac{1}{\beta + s_{\max} \|\mathbf{T}\|} \right] \geq \frac{1}{\beta + s_{\max} \mathbb{E}[\|\mathbf{T}\|]}. \end{aligned}$$

By Lemma 4, $\mathbb{E}[\|\mathbf{T}\|] \leq (1 + \sqrt{p/n})^2$. Hence, for $\alpha > p/n - 1 + s_{\max}(1 + \sqrt{p/n})^2$, the solution to equation (14) satisfies

$$r^* \leq r_{\max} = \frac{n}{p} \frac{\alpha}{1 + \alpha - p/n - s_{\max}(1 + \sqrt{p/n})^2}.$$

A.4. Proof of Lemma 9. Let $\hat{\mathbf{S}} = \frac{1}{n} \sum_{k=1}^n \mathbf{x}_k \mathbf{x}_k^T$ and $\hat{\mathbf{T}} = \frac{1}{n} \sum_{k=1}^n \xi_k \xi_k^T$, where $\mathbf{x}_i = \frac{1}{\sqrt{p}} \xi_i$ and $\xi_i \stackrel{i.i.d.}{\sim} N(\mathbf{0}, \mathbf{I})$. Then equation (18) may be written as

$$\begin{aligned} \frac{1}{r} g(\mathbf{u})_i &= 1 - \frac{1/(1+\alpha)}{\frac{1}{p} \mathbf{x}_i^T (\hat{\mathbf{S}} + \beta \mathbf{I})^{-1} \mathbf{x}_i} = 1 - \frac{1/(1+\alpha)}{\frac{1}{p} \xi_i^T (\mathbf{S}_p^{-\frac{1}{2}} \hat{\mathbf{S}} \mathbf{S}_p^{-\frac{1}{2}} + \beta \mathbf{S}_p^{-1})^{-1} \xi_i} \\ (33) \quad &= 1 - \frac{1}{1+\alpha} \frac{1}{\frac{1}{p} \xi_i^T \mathbf{E} \xi_i}, \end{aligned}$$

where $\mathbf{E} = (\hat{\mathbf{T}} + \beta \mathbf{S}_p^{-1})^{-1}$ and $\beta = \alpha \frac{n}{p} \frac{1}{r}$. The quadratic form $\frac{1}{p} \xi_i^T \mathbf{E} \xi_i$ is difficult to analyze directly because $\mathbf{E}$ depends on ξ_i . To disentangle this dependency, let $\hat{\mathbf{T}}_{-i} = \frac{1}{n} \sum_{k \neq i} \xi_k \xi_k^T$, and $\mathbf{E}_{-i} = (\hat{\mathbf{T}}_{-i} + \beta \mathbf{S}_p^{-1})^{-1}$. As $\mathbf{E}^{-1}$ and $\mathbf{E}_{-i}^{-1}$ differ by a rank-one matrix $\frac{1}{n} \xi_i \xi_i^T$, by the Sherman–Morrison formula,

$$\mathbf{E} = \mathbf{E}_{-i} - \frac{1}{n} \frac{\mathbf{E}_{-i} \xi_i \xi_i^T \mathbf{E}_{-i}}{1 + \frac{1}{n} \xi_i^T \mathbf{E}_{-i} \xi_i}.$$

Therefore, denoting by Q_i the quadratic form

$$(34) \quad Q_i(r) \equiv Q_i = \frac{1}{p} \xi_i^T \mathbf{E}_{-i} \xi_i,$$

it follows that

$$(35) \quad \frac{1}{p} \xi_i^T \mathbf{E} \xi_i = Q_i - \frac{\frac{p}{n} Q_i^2}{1 + \frac{p}{n} Q_i} = \frac{Q_i}{1 + \frac{p}{n} Q_i}.$$

Plugging this expression into equation (33) gives

$$(36) \quad \frac{1}{r}g(\mathbf{u})_i = \frac{Q_i(1 + \alpha - \frac{p}{n}) - 1}{(1 + \alpha)Q_i}.$$

Next, to establish a concentration bound for $g(\mathbf{u})_i/r$, we study the concentration of Q_i . Since $\xi_i \sim N(\mathbf{0}, \mathbf{I})$ and is independent of $\mathbf{E}_{-i}$,

$$\mathbb{E}Q_i = \mathbb{E} \operatorname{tr}(\mathbf{E}_{-i})/p.$$

We first show that Q_i concentrates tightly around $\operatorname{tr}(\mathbf{E}_{-i})/p$ in view of concentration of quadratic forms. We then show that $\operatorname{tr}(\mathbf{E}_{-i})$ concentrates tightly around its mean using results about the concentration of certain functions of the eigenvalues of random matrices.

Applying Lemma 5 with $\xi = \xi_i$ and viewing the matrix $\mathbf{E}_{-i}$ as fixed,

$$\Pr\left(\left|Q_i - \frac{1}{p} \operatorname{tr}(\mathbf{E}_{-i})\right| > \epsilon\right) \leq 2 \exp\left(-c_1 \min\left\{\frac{c_2^2 p^2 \epsilon^2}{\|\mathbf{E}_{-i}\|_F^2}, \frac{c_2 p \epsilon}{\|\mathbf{E}_{-i}\|}\right\}\right),$$

where the above probability is only w.r.t. ξ_i . Next, given that $\mathbf{E}_{-i} = (\mathbf{T}_{-i} + \beta \mathbf{S}_p^{-1})^{-1}$, then $\|\mathbf{E}_{-i}\| \leq \frac{s_{\max}}{\beta}$ and $\|\mathbf{E}_{-i}\|_F^2 \leq p s_{\max}^2 / \beta^2$. Thus,

$$(37) \quad \Pr\left(\left|Q_i - \frac{1}{p} \operatorname{tr}(\mathbf{E}_{-i})\right| > \epsilon\right) \leq C \exp(-cp\epsilon^2),$$

where now the probability is over all of the ξ_k 's.

It remains to obtain a concentration inequality for $\operatorname{tr}(\mathbf{E}_{-i})/p$. To this end, consider the following $p \times (n - 1 + p)$ matrix:

$$\mathbf{Y} = \begin{pmatrix} \xi_1 & \cdots & \xi_{i-1} & \xi_{i+1} & \cdots & \xi_n & \sqrt{n\beta} \mathbf{S}_p^{-1/2} \end{pmatrix}.$$

By definition, all entries of $\mathbf{Y}$ are independent, the first $p \times (n - 1)$ are standard Gaussian random variables and the rest deterministic. Then, by Guionnet and Zeitouni (2000), Corollary 1.8b,² for any function $h : \mathbb{R} \rightarrow \mathbb{R}$ such that $h(x^2)$ is Lipschitz with constant L , and for any $\delta > 0$,

$$(38) \quad \Pr\left(\frac{1}{K} \left| \operatorname{tr} h\left(\frac{\mathbf{Y}\mathbf{Y}^T}{K}\right) - \mathbb{E} \operatorname{tr} h\left(\frac{\mathbf{Y}\mathbf{Y}^T}{K}\right) \right| > \delta\right) \leq 2 \exp\left(-\frac{\delta^2 K^2}{2L^2}\right),$$

where $K = 2p + n - 1$ and for a symmetric matrix $\mathbf{A}$ with eigenvalues λ_j , $\operatorname{tr} h(\mathbf{A}) = \sum_j h(\lambda_j)$.

Since $\mathbf{Y}\mathbf{Y}^T = n(\hat{\mathbf{T}}_{-i} + \beta \mathbf{S}_p^{-1}) = n\mathbf{E}_{-i}^{-1}$, consider the function

$$h(x) = \frac{n}{p} \cdot \frac{1}{x}$$

for which $\frac{1}{2p+n-1} \operatorname{tr} h(\mathbf{Y}\mathbf{Y}^T / (2p + n - 1)) = \operatorname{tr}(\mathbf{E}_{-i})/p$. Next note that for sufficiently large n and sufficiently small ϵ ,

$$\lambda_{\min}\left(\frac{\mathbf{Y}\mathbf{Y}^T}{2p + n - 1}\right) = \frac{n}{2p + n - 1} \lambda_{\min}(\hat{\mathbf{T}}_{-i} + \beta \mathbf{S}_p^{-1}) \geq \frac{1}{2\gamma + 1 + \epsilon} \frac{\beta}{s_{\max}} = x_0.$$

²There is a typo in the original paper. In the notation of their Corollary 1.8, $\mathbf{Z}$ should be replaced with $\mathbf{Z}/(M + N)$.

We thus apply the function h only in the interval $x \geq x_0$. The Lipschitz constant of $h(x^2)$ for n sufficiently large is bounded by

$$L \leq \left| \frac{d}{dx} h(x^2) \right|_{x=x_0} \leq 16 \frac{(\gamma + 0.5 + \epsilon)^3}{\gamma - \epsilon} \left(\frac{s_{\max}}{\beta} \right)^3 \leq 16 \frac{(\gamma + 1)^3}{\gamma} \left(\frac{s_{\max}}{\beta} \right)^3.$$

Hence, applying (38), there exists a positive constant c that depends on γ, α, r and $s_{\max}$ such that

$$(39) \quad \Pr\left(\frac{1}{p} |\text{tr}(\mathbf{E}_{-i}) - \mathbb{E} \text{tr}(\mathbf{E}_{-i})| > \delta\right) \leq 2 \exp(-cp^2 \delta^2).$$

Next, by the triangle inequality,

$$\begin{aligned} \Pr\left(\left|Q_i - \frac{\mathbb{E} \text{tr}(\mathbf{E}_{-i})}{p}\right| > \epsilon\right) &\leq \Pr\left(\left|Q_i - \frac{\text{tr}(\mathbf{E}_{-i})}{p}\right| > \frac{\epsilon}{2}\right) \\ &\quad + \Pr\left(\left|\frac{\text{tr}(\mathbf{E}_{-i})}{p} - \frac{\mathbb{E} \text{tr}(\mathbf{E}_{-i})}{p}\right| > \frac{\epsilon}{2}\right). \end{aligned}$$

Combining the above equation with equations (37) and (39), implies that at the value of r specified in Proposition 1, for which $\mathbb{E}[\text{tr}(\mathbf{E}_{-i})/p] = \frac{1}{1+\alpha-\frac{p}{n}}$,

$$(40) \quad \Pr\left(\left|Q_i - \frac{1}{1+\alpha-\frac{p}{n}}\right| > \epsilon\right) < C e^{-cp\epsilon^2}.$$

We are finally ready to establish a concentration result for $\frac{1}{r}g(\mathbf{u})$. Combining equation (36) and a union bound over all p coordinates of g ,

$$\begin{aligned} \Pr\left(\left\|\frac{1}{r}g(\mathbf{u})\right\|_{\infty} > \epsilon\right) &\leq p \Pr\left(\left|\frac{1}{r}g(\mathbf{u})_i\right| > \epsilon\right) \\ &\leq p \Pr\left(\left|\frac{Q_i(1+\alpha-p/n)-1}{(1+\alpha)Q_i}\right| > \epsilon\right). \end{aligned}$$

Applying equation (10) with $\lambda = 1$ to the equation above gives

$$\Pr\left(\left\|\frac{1}{r}g(\mathbf{u})\right\|_{\infty} > \epsilon\right) < p \Pr\left(\left|Q_i\left(1+\alpha-\frac{p}{n}\right)-1\right| > \epsilon\right) + p \Pr((1+\alpha)Q_i < 1).$$

By equation (40), the first term on the RHS is exponentially small in p . As for the second term, since $(1+\alpha)^{-1} < (1+\alpha-p/n)^{-1}$, then again by equation (40), $\Pr(Q_i < 1/(1+\alpha))$ is also exponentially small in p . The lemma thus follows from the boundedness of r from above, as established in Proposition 1.

A.5. TME with outliers. Consider an ϵ -contamination model, where $(1-\epsilon)n$ samples follow an elliptical distribution with shape matrix $\mathbf{S}_{\text{in}}$, and the remaining ϵn follow an elliptical distribution with shape matrix $\mathbf{S}_{\text{out}}$. We conjecture that under suitable assumptions, the inlier and outlier weights of the TME concentrate around two values, w_{in} and w_{out} , respectively.

For our procedure to select the inliers, we further assume that the inlier weights are approximately Gaussian distributed around w_{in} with an unknown standard deviation σ_{in} . To estimate w_{in} and σ_{in} we compute a nonparametric density estimate $\hat{f}(w)$ of all n weights (using MATLAB's `ksdensity` procedure). Then $w_{\text{in}} = \arg \max_w \hat{f}(w)$ is the weight with highest estimated density. Next, for some r we find the largest interval $[w_L, w_R]$ around w_{in} so that for $w \in [w_L, w_R]$ we have $\hat{f}(w) \geq r \max \hat{f}(w) = r \hat{f}(w_{\text{in}})$. Then, given our assumption that the weights are Gaussian distributed, $\sigma_{\text{in}} = \frac{1}{2}(w_R - w_L)/\sqrt{-2\log(r)}$. In our simulations, we used $r = 0.7$, which worked well across all different contamination levels.

Clearly, one may obtain improved estimates of these quantities, as well as the unknown ϵ , by fitting a mixture of two Gaussians to the weights. However, for our illustrative example, we opted for the above simpler procedure.

Acknowledgments. We thank the three anonymous referees for multiple suggestions that greatly improved the manuscript. We also thank Teng Zhang and Ofer Zeitouni for useful discussions, and Tony Cai, Harrison Zhou, Elizaveta Levina and Peter Bickel for correspondence regarding their papers. We also thank the IMA and the schools of the authors for supporting collaborative visits. B.N. is the incumbent William Petschek professorial chair of mathematics.

This work was supported by NSF Grants DMS-14-18386, DMS-18-21266 and GRFP-00039202, UMN Doctoral Dissertation Fellowship and the Feinberg Foundation Visiting Faculty Program Fellowship of the Weizmann Institute of Science.

SUPPLEMENTARY MATERIAL

Supplement to “Robust sparse covariance estimation by thresholding Tyler’s M-estimator” (DOI: [10.1214/18-AOS1793SUPP](https://doi.org/10.1214/18-AOS1793SUPP); .pdf). This supplement contains additional theoretical results and proofs omitted from main article.

REFERENCES

- ABRAMOVICH, Y. I. and SPENCER, N. K. (2007). Diagonally loaded normalised sample matrix inversion (LNSMI) for outlier-resistant adaptive filtering. In *Proceedings of the IEEE 32nd Intl. Conf. on Acoustics, Speech, and Signal Proc. (ICASSP)* 1105–1108.
- ANDERSON, T. W. (2003). *An Introduction to Multivariate Statistical Analysis*, 3rd ed. *Wiley Series in Probability and Statistics*. Wiley, Hoboken, NJ. [MR1990662](#)
- AVELLA-MEDINA, M., BATTEY, H. S., FAN, J. and LI, Q. (2018). Robust estimation of high-dimensional covariance and precision matrices. *Biometrika* **105** 271–284. [MR3804402](#) <https://doi.org/10.1093/biomet/asy011>
- BALAKRISHNAN, S., DU, S. S., LI, J. and SINGH, A. (2017). Computationally efficient robust sparse estimation in high dimensions. In *Conference on Learning Theory* 169–212.
- BHATIA, R. (1997). *Matrix Analysis. Graduate Texts in Mathematics* **169**. Springer, New York. [MR1477662](#) <https://doi.org/10.1007/978-1-4612-0653-8>
- BICKEL, P. J. and LEVINA, E. (2008). Covariance regularization by thresholding. *Ann. Statist.* **36** 2577–2604. [MR2485008](#) <https://doi.org/10.1214/08-AOS600>
- CAI, T. and LIU, W. (2011). Adaptive thresholding for sparse covariance matrix estimation. *J. Amer. Statist. Assoc.* **106** 672–684. [MR2847949](#) <https://doi.org/10.1198/jasa.2011.tm10560>
- CAI, T. T., REN, Z. and ZHOU, H. H. (2016). Estimating structured high-dimensional covariance and precision matrices: Optimal rates and adaptive estimation. *Electron. J. Stat.* **10** 1–59. [MR3466172](#) <https://doi.org/10.1214/15-EJS1081>
- CAI, T. T. and ZHOU, H. H. (2012a). Optimal rates of convergence for sparse covariance matrix estimation. *Ann. Statist.* **40** 2389–2420. [MR3097607](#) <https://doi.org/10.1214/12-AOS998>
- CAI, T. T. and ZHOU, H. H. (2012b). Minimax estimation of large covariance matrices under ℓ_1 -norm. *Statist. Sinica* **22** 1319–1349. [MR3027084](#)
- CAMBANIS, S., HUANG, S. and SIMONS, G. (1981). On the theory of elliptically contoured distributions. *J. Multivariate Anal.* **11** 368–385. [MR0629795](#) [https://doi.org/10.1016/0047-259X\(81\)90082-8](https://doi.org/10.1016/0047-259X(81)90082-8)
- CHEN, M., GAO, C. and REN, Z. (2018). Robust covariance and scatter matrix estimation under Huber’s contamination model. *Ann. Statist.* **46** 1932–1960. [MR3845006](#) <https://doi.org/10.1214/17-AOS1607>
- CHEN, Y., WIESEL, A. and HERO, A. O. III (2011). Robust shrinkage estimation of high-dimensional covariance matrices. *IEEE Trans. Signal Process.* **59** 4097–4107. [MR2865971](#) <https://doi.org/10.1109/TSP.2011.2138698>
- COUILLET, R., KAMMOUN, A. and PASCAL, F. (2016). Second order statistics of robust estimators of scatter. Application to GLRT detection for elliptical signals. *J. Multivariate Anal.* **143** 249–274. [MR3431431](#) <https://doi.org/10.1016/j.jmva.2015.08.021>
- COUILLET, R. and MCKAY, M. (2014). Large dimensional analysis and optimization of robust shrinkage covariance matrix estimators. *J. Multivariate Anal.* **131** 99–120. [MR3252638](#) <https://doi.org/10.1016/j.jmva.2014.06.018>

- COUILLET, R., PASCAL, F. and SILVERSTEIN, J. W. (2014). Robust estimates of covariance matrices in the large dimensional regime. *IEEE Trans. Inform. Theory* **60** 7269–7278. MR3273053 <https://doi.org/10.1109/TIT.2014.2354045>
- COUILLET, R., PASCAL, F. and SILVERSTEIN, J. W. (2015). The random matrix regime of Maronna's M -estimator with elliptically distributed samples. *J. Multivariate Anal.* **139** 56–78. MR3349480 <https://doi.org/10.1016/j.jmva.2015.02.020>
- DAVIDSON, K. R. and SZAREK, S. J. (2001). Local operator theory, random matrices and Banach spaces. In *Handbook of the Geometry of Banach Spaces, Vol. I* 317–366. North-Holland, Amsterdam. MR1863696 [https://doi.org/10.1016/S1874-5849\(01\)80010-3](https://doi.org/10.1016/S1874-5849(01)80010-3)
- DÜMBGEN, L. (1998). On Tyler's M -functional of scatter in high dimension. *Ann. Inst. Statist. Math.* **50** 471–491. MR1664575 <https://doi.org/10.1023/A:1003573311481>
- DÜMBGEN, L., NORDHAUSEN, K. and SCHUHMACHER, H. (2016). New algorithms for M -estimation of multivariate scatter and location. *J. Multivariate Anal.* **144** 200–217. MR3434950 <https://doi.org/10.1016/j.jmva.2015.11.009>
- DÜMBGEN, L., PAULY, M. and SCHWEIZER, T. (2015). M -functionals of multivariate scatter. *Stat. Surv.* **9** 32–105. MR3324616 <https://doi.org/10.1214/15-SS109>
- EL KAROUI, N. (2008). Operator norm consistent estimation of large-dimensional sparse covariance matrices. *Ann. Statist.* **36** 2717–2756. MR2485011 <https://doi.org/10.1214/07-AOS559>
- FALK, M. (2002). The sample covariance is not efficient for elliptical distributions. *J. Multivariate Anal.* **80** 358–377. MR1889781 <https://doi.org/10.1006/jmva.2000.1983>
- FANG, K. T., KOTZ, S. and NG, K. W. (1990). *Symmetric Multivariate and Related Distributions. Monographs on Statistics and Applied Probability* **36**. CRC Press, London. MR1071174 <https://doi.org/10.1007/978-1-4899-2937-2>
- FRAHM, G. (2004). Generalized elliptical distributions: Theory and applications. Ph.D. thesis, Univ. Köln.
- FRAHM, G. and JAEKEL, U. (2007). Tyler's M -estimator, random matrix theory, and generalized elliptical distributions with applications to finance. Discussion Papers in Econometrics and Statistics, No. 2/07, Institute of Econometrics and Statistics, Univ. Cologne.
- FRAHM, G. and JAEKEL, U. (2010). A generalization of Tyler's M -estimators to the case of incomplete data. *Comput. Statist. Data Anal.* **54** 374–393. MR2756433 <https://doi.org/10.1016/j.csda.2009.08.019>
- GOES, J., LERMAN, G. and NADLER, B. (2020). Supplement to “Robust sparse covariance estimation by thresholding Tyler's M -estimator.” <https://doi.org/10.1214/18-AOS1793SUPP>.
- GUINNET, A. and ZEITOUNI, O. (2000). Concentration of the spectral measure for large matrices. *Electron. Commun. Probab.* **5** 119–136. MR1781846 <https://doi.org/10.1214/ECP.v5-1026>
- HAN, F., LU, J. and LIU, H. (2014). Robust scatter matrix estimation for high dimensional distributions with heavy tails. Technical Report, Princeton Univ., Princeton, NJ.
- HUBER, P. J. and RONCHETTI, E. M. (2009). *Robust Statistics*, 2nd ed. *Wiley Series in Probability and Statistics*. Wiley, Hoboken, NJ. MR2488795 <https://doi.org/10.1002/9780470434697>
- KAMMOUN, A., COUILLET, R., PASCAL, F. and ALOUINI, M.-S. (2018). Optimal design of the adaptive normalized matched filter detector using regularized Tyler estimators. *IEEE Trans. Aerosp. Electron. Syst.* **54** 755–769. MR3456682 <https://doi.org/10.1109/TSP.2015.2494858>
- KELKER, D. (1970). Distribution theory of spherical distributions and a location-scale parameter generalization. *Sankhyā Ser. A* **32** 419–438. MR0287628
- KENT, J. T. and TYLER, D. E. (1988). Maximum likelihood estimation for the wrapped Cauchy distribution. *J. Appl. Stat.* **15** 247–254.
- KENT, J. T. and TYLER, D. E. (1991). Redescending M -estimates of multivariate location and scatter. *Ann. Statist.* **19** 2102–2119. MR1135166 <https://doi.org/10.1214/aos/1176348388>
- LAM, C. and FAN, J. (2009). Sparsistency and rates of convergence in large covariance matrix estimation. *Ann. Statist.* **37** 4254–4278. MR2572459 <https://doi.org/10.1214/09-AOS720>
- MARDIA, K. V., KENT, J. T. and BIBBY, J. M. (1979). *Multivariate Analysis*. Academic Press, London. MR0560319
- MARONNA, R. A. (1976). Robust M -estimators of multivariate location and scatter. *Ann. Statist.* **4** 51–67. MR0388656
- MARONNA, R. A. and YOHAI, V. J. (2017). Robust and efficient estimation of multivariate scatter and location. *Comput. Statist. Data Anal.* **109** 64–75. MR3603641 <https://doi.org/10.1016/j.csda.2016.11.006>
- MORALES-JIMENEZ, D., COUILLET, R. and MCKAY, M. R. (2015). Large dimensional analysis of robust M -estimators of covariance with outliers. *IEEE Trans. Signal Process.* **63** 5784–5797. MR3405886 <https://doi.org/10.1109/TSP.2015.2460225>
- NORDHAUSEN, K. and TYLER, D. E. (2015). A cautionary note on robust covariance plug-in methods. *Biometrika* **102** 573–588. MR3394276 <https://doi.org/10.1093/biomet/asv022>

- OLLILA, E. and KOIVUNEN, V. (2003). Robust antenna array processing using M -estimators of pseudo-covariance. In *Proceedings of the IEEE 14th International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC)* 2659–2663.
- OLLILA, E. and TYLER, D. E. (2012). Distribution-free detection under complex elliptically symmetric clutter distribution. In *IEEE 7th Sensor Array and Multichannel Signal Processing Workshop, (SAM)* 413–416.
- OLLILA, E. and TYLER, D. E. (2014). Regularized M -estimators of scatter matrix. *IEEE Trans. Signal Process.* **62** 6059–6070. [MR3281544](#) <https://doi.org/10.1109/TSP.2014.2360826>
- PASCAL, F., CHITOUR, Y. and QUEK, Y. (2014). Generalized robust shrinkage estimator and its application to STAP detection problem. *IEEE Trans. Signal Process.* **62** 5640–5651. [MR3273520](#) <https://doi.org/10.1109/TSP.2014.2355779>
- ROTHMAN, A. J., LEVINA, E. and ZHU, J. (2009). Generalized thresholding of large covariance matrices. *J. Amer. Statist. Assoc.* **104** 177–186. [MR2504372](#) <https://doi.org/10.1198/jasa.2009.0101>
- RUDELSON, M. and VERSHYNIN, R. (2013). Hanson–Wright inequality and sub-Gaussian concentration. *Electron. Commun. Probab.* **18** no. 82, 9. [MR3125258](#) <https://doi.org/10.1214/ECP.v18-2865>
- SIRKIÄ, S., TASKINEN, S. and OJA, H. (2007). Symmetrised M -estimators of multivariate scatter. *J. Multivariate Anal.* **98** 1611–1629. [MR2370110](#) <https://doi.org/10.1016/j.jmva.2007.06.005>
- SOLOVEYCHIK, I. and WIESEL, A. (2014). Tyler’s covariance matrix estimator in elliptical models with convex structure. *IEEE Trans. Signal Process.* **62** 5251–5259. [MR3268109](#) <https://doi.org/10.1109/TSP.2014.2348951>
- SUN, Y., BABU, P. and PALOMAR, D. P. (2014). Regularized Tyler’s scatter estimator: Existence, uniqueness, and algorithms. *IEEE Trans. Signal Process.* **62** 5143–5156. [MR3260359](#) <https://doi.org/10.1109/TSP.2014.2348944>
- SUN, Y., BABU, P. and PALOMAR, D. P. (2016). Robust estimation of structured covariance matrix for heavy-tailed elliptical distributions. *IEEE Trans. Signal Process.* **64** 3576–3590. [MR3515702](#) <https://doi.org/10.1109/TSP.2016.2546222>
- TYLER, D. E. (1987a). A distribution-free M -estimator of multivariate scatter. *Ann. Statist.* **15** 234–251. [MR0885734](#) <https://doi.org/10.1214/aos/1176350263>
- TYLER, D. E. (1987b). Statistical analysis for the angular central Gaussian distribution on the sphere. *Biometrika* **74** 579–589. [MR0909362](#) <https://doi.org/10.1093/biomet/74.3.579>
- WIESEL, A. (2012). Unified framework to regularized covariance estimation in scaled Gaussian models. *IEEE Trans. Signal Process.* **60** 29–38. [MR2932100](#) <https://doi.org/10.1109/TSP.2011.2170685>
- WIESEL, A. and ZHANG, T. (2014). Structured robust covariance estimation. *Found. Trends Signal Process.* **8** 127–216. [MR3438278](#) <https://doi.org/10.1561/20000000053>
- ZHANG, T., CHENG, X. and SINGER, A. (2016). Marčenko–Pastur law for Tyler’s M -estimator. *J. Multivariate Anal.* **149** 114–123. [MR3507318](#) <https://doi.org/10.1016/j.jmva.2016.03.010>