

Outlier-Robust High-Dimensional Sparse Estimation via Iterative Filtering

Ilias Diakonikolas*
University of Wisconsin-Madison
`ilias@cs.wisc.edu`

Sushrut Karmalkar†
UT Austin
`s.sushrut@gmail.com`

Daniel Kane‡
University of California, San Diego
`dakane@ucsd.edu`

Eric Price§
UT Austin
`ecprice@cs.utexas.edu`

Alistair Stewart¶
Web3 Foundation
`stewart.al@gmail.com`

November 18, 2019

Abstract

We study high-dimensional sparse estimation tasks in a robust setting where a constant fraction of the dataset is adversarially corrupted. Specifically, we focus on the fundamental problems of robust sparse mean estimation and robust sparse PCA. We give the first practically viable robust estimators for these problems. In more detail, our algorithms are sample and computationally efficient and achieve near-optimal robustness guarantees. In contrast to prior provable algorithms which relied on the ellipsoid method, our algorithms use spectral techniques to iteratively remove outliers from the dataset. Our experimental evaluation on synthetic data shows that our algorithms are scalable and significantly outperform a range of previous approaches, nearly matching the best error rate without corruptions.

*Supported by NSF Award CCF-1652862 (CAREER) and a Sloan Research Fellowship.

†Supported by NSF Award CNS-1414023.

‡Supported by NSF Award CCF-1553288 (CAREER) and a Sloan Research Fellowship

§Supported in part by NSF Award CCF-1751040 (CAREER).

¶Part of this work was performed while the author was a postdoctoral researcher at USC.

1 Introduction

1.1 Background

The task of leveraging sparsity to extract meaningful information from high-dimensional datasets is a fundamental problem of significant practical importance, motivated by a range of data analysis applications. Various formalizations of this general problem have been investigated in statistics and machine learning for at least the past two decades, see, e.g., [HTW15] for a recent textbook on the topic. This paper focuses on the *unsupervised setting* and in particular on estimating the parameters of a high-dimensional distribution under sparsity assumptions. Concretely, we study the problems of *sparse mean estimation* and *sparse PCA* under natural data generating models.

The classical setup in statistics is that the data was generated by a probabilistic model of a given type. This is a simplifying assumption that is only approximately valid, as real datasets are typically exposed to some source of contamination. The field of robust statistics [Hub64, HR09, HRRS86] aims to design estimators that are *robust* in the presence of model misspecification. In recent years, designing computationally efficient robust estimators for high-dimensional settings has become a pressing challenge in a number of applications. These include the analysis of biological datasets, where natural outliers are common [RPW⁺02, PLJD10, LAT⁺08] and can contaminate the downstream statistical analysis, and *data poisoning attacks* [BNJT10], where even a small fraction of fake data (outliers) can substantially degrade the learned model [BNL12, SKL17].

This discussion motivates the design of robust estimators that can tolerate a *constant* fraction of adversarially corrupted data. We will use the following model of corruptions (see, e.g., [DKK⁺16]):

Definition 1.1. *Given $0 < \varepsilon < 1/2$ and a family of distributions $\mathcal{D}$ on $\mathbb{R}^d$, the adversary operates as follows: The algorithm specifies some number of samples N , and N samples $X_1, X_2, \dots, X_N$ are drawn from some (unknown) $D \in \mathcal{D}$. The adversary is allowed to inspect the samples, removes εN of them, and replaces them with arbitrary points. This set of N points is then given to the algorithm. We say that a set of samples is ε -corrupted if it is generated by the above process.*

Our model of corruptions generalizes several other robustness models, including Huber’s contamination model [Hub64] and the malicious PAC model [Val85, KL93].

In the context of robust sparse mean estimation, we are given an ε -corrupted set of samples from an unknown mean Gaussian distribution $\mathcal{N}(\mu, I)$, where $\mu \in \mathbb{R}^d$ is assumed to be k -sparse, and the goal is to output a hypothesis vector $\hat{\mu}$ that approximates μ in ℓ_2 -norm. In the context of robust sparse PCA (in the spiked covariance model), we are given an ε -corrupted set of samples from $\mathcal{N}(\mathbf{0}, \rho vv^T)$, where $v \in \mathbb{R}^d$ is assumed to be k -sparse and the goal is to approximate v . In both settings, we would like to design computationally efficient estimators with sample complexity $\text{poly}(k, \log d, 1/\varepsilon)$, i.e., close to the information theoretic minimum, that achieve near-optimal error guarantees.

Until recently, even for the simplest high-dimensional parameter estimation settings, no polynomial time robust learning algorithms with dimension-independent error guarantees were known. Two concurrent works [DKK⁺16, LRV16] made the first progress on this front for the unsupervised setting. Specifically, [DKK⁺16, LRV16] gave the first polynomial time algorithms for robustly learning the mean and covariance of high-dimensional Gaussians and other models. These works focused on the dense regime and as a result did not obtain algorithms with sublinear sample complexity in the sparse setting. Building on [DKK⁺16], more recent work [BDLS17] obtained sample efficient polynomial time algorithms for the robust *sparse* setting, and in particular for the problems of robust sparse mean estimation and robust sparse PCA studied in this paper. These algorithms are based the unknown convex programming methodology of [DKK⁺16] and in particular *inherently*

rely on the ellipsoid algorithm. Moreover, the separation oracle required for the ellipsoid algorithm turns out to be another convex program — corresponding to an SDP to solve sparse PCA. As a consequence, the running time of these algorithms, while polynomially bounded, is impractically high.

1.2 Our Results and Techniques

The main contribution of this paper is the design of significantly faster robust estimators for the aforementioned high-dimensional sparse problems. More specifically, our algorithms are iterative and each iteration involves a simple spectral operation (computing the largest eigenvalue of an approximate matrix). Our algorithms achieve the same error guarantee as [BDLS17] with similar sample complexity. At the technical level, we enhance the *iterative filtering methodology* of [DKK⁺16] to the sparse setting, which we believe is of independent interest and could lead to faster algorithms for other robust sparse estimation tasks as well.

For robust sparse mean estimation, we show:

Theorem 1.2 (Robust Sparse Mean Estimation). *Let $D \sim \mathcal{N}(\mu, I)$ be a Gaussian distribution on $\mathbb{R}^d$ with unknown k -sparse mean vector μ , and $\varepsilon > 0$. Let S be an ε -corrupted set of samples from D of size $N = \tilde{\Omega}(k^2 \log(d)/\varepsilon^2)$. There exists an algorithm that, on input S , k , and ε runs in polynomial time returns $\hat{\mu}$ such that with probability at least $2/3$ it holds $\|\hat{\mu} - \mu\|_2 = O(\varepsilon \sqrt{\log(1/\varepsilon)})$.*

Some comments are in order. First, the sample complexity of our algorithm is asymptotically the same as that of [BDLS17], and matches the lower bound of [DKS17] against Statistical Query algorithms for this problem. The major advantage of our algorithm over [BDLS17] is that while their algorithm made use of the ellipsoid method, ours uses only spectral techniques and is scalable.

For robust sparse PCA in the spiked covariance model, we show:

Theorem 1.3 (Robust Sparse PCA). *Let $D \sim \mathcal{N}(\mathbf{0}, I + \rho vv^T)$ be a Gaussian distribution on $\mathbb{R}^d$ with spiked covariance for an unknown k -sparse unit vector v , and $0 < \rho < O(1)$. For $\varepsilon > 0$, let S be an ε -corrupted set of samples from D of size $N = \Omega(k^4 \log^4(d/\varepsilon)/\varepsilon^2)$. There exists an algorithm that, on input S , k , and ε , runs in polynomial time and returns $\hat{v} \in \mathbb{R}^d$ such that with probability at least $2/3$ we have that $\|\hat{v}\hat{v}^T - vv^T\|_F = O(\varepsilon \log(1/\varepsilon)/\rho)$.*

The sample complexity upper bound in Theorem 1.3 is somewhat worse than the information theoretic optimum of $\Theta(k^2 \log d/\varepsilon^2)$. While the ellipsoid-based algorithm of [BDLS17] achieves near-optimal sample complexity (within logarithmic factors), our algorithm is practically viable as it only uses spectral operations. We also note that the sample complexity in our above theorem is not known to be optimal for our algorithm. It seems quite plausible, via a tighter analysis, that our algorithm in fact has near-optimal sample complexity as well.

For both of our algorithms, in the most interesting regime of $k \ll \sqrt{d}$, the running time per iteration is dominated by the $O(Nd^2)$ computation of the empirical covariance matrix. The number of iterations is at most εN , although it typically is much smaller, so both algorithms take at most $O(\varepsilon N^2 d^2)$ time.

1.3 Related Work

There is extensive literature on exploiting sparsity in statistical estimation (see, e.g., [HTW15]). In this section, we summarize the related work that is directly related to the results of this paper. Sparse mean estimation is arguably one of the most fundamental sparse estimation tasks and is closely related to the Gaussian sequence model [Tsy08, Joh17]. The task of sparse PCA in the

spiked covariance model, initiated in [Joh01], has been extensively investigated (see Chapter 8 of [HTW15] and references therein). In this work, we design algorithms for the aforementioned problems that are robust to a constant fraction of outliers.

Learning in the presence of outliers is an important goal in statistics studied since the 1960s [Hub64]. See, e.g., [HR09, HRRS86] for book-length introductions in robust statistics. Until recently, all known computationally efficient high-dimensional estimators could tolerate a negligible fraction of outliers, even for the task of mean estimation. Recent work [DKK⁺16, LRV16] gave the first efficient robust estimators for basic high-dimensional unsupervised tasks, including mean and covariance estimation. Since the dissemination of [DKK⁺16, LRV16], there has been a flurry of research activity on computationally efficient robust learning in high dimensions [BDLS17, CSV17, DKK⁺17, DKS17, DKK⁺18a, SCV18, DKS18b, DKS18a, HL18, KSS18, PSBR18, DKK⁺18b, KKM18, DKS19, LSLC18a, CDKS18, CDG18, CDGW19].

In the context of robust sparse estimation, [BDLS17] obtained sample-efficient and polynomial time algorithms for robust sparse mean estimation and robust sparse PCA. The main difference between [BDLS17] and the results of this paper is that the [BDLS17] algorithms use the ellipsoid method (whose separation oracle is an SDP). Hence, these algorithms are prohibitively slow for practical applications. More recent work [LSLC18b] gave an iterative method for robust sparse mean estimation, which however requires multiple solutions to a convex relaxation for sparse PCA in each iteration. Finally, [LLC19] proposed an algorithm for robust sparse mean estimation via iterative trimmed hard thresholding. While this algorithm seems practically viable in terms of runtime, it can only tolerate $1/(\sqrt{k} \log(nd))$ – i.e., *sub-constant* – fraction of corruptions.

1.4 Paper Organization

In Section 2, we describe our algorithms and provide a detailed sketch of their analysis. In Section 6, we report detailed experiments demonstrating the performance of our algorithms on synthetic data in various parameter regimes. Due to space limitations, the full proofs of correctness for our algorithms can be found in the full version of this paper.

2 Algorithms

In this section, we describe our algorithms in tandem with a detailed outline of the intuition behind them and a sketch of their analysis. Due to space limitations, the proof of correctness is deferred to the full version of our paper.

At a high-level, our algorithms use the iterative filtering methodology of [DKK⁺16]. The main idea is to iteratively remove a small subset of the dataset, so that eventually we have removed all the important outliers and the standard estimator (i.e., the estimator we would have used in the noiseless case) works. Before we explain our new ideas that enhance the filtering methodology to the sparse setting, we provide a brief technical description of the approach.

Overview of Iterative Filtering. The basic idea of iterative filtering [DKK⁺16] is the following: In a given iteration, carefully pick some test statistic (such as $v \cdot x$ for a well-chosen v). If there were no outliers, this statistic would follow a nice distribution (with good concentration properties). This allows us to do some sort of statistical hypothesis testing of the “null hypothesis” that each x_i is an inlier, rejecting it (and believing that x_i is an outlier) if $v \cdot x_i$ is far from the expected distribution. Because there are a large number of such hypotheses, one uses a procedure reminiscent of the Benjamini-Hochberg procedure [BH95] to find a candidate set of outliers with low *false discovery rate* (FDR), i.e., a set with more outliers than inliers in expectation. This procedure looks for a

threshold T such that the fraction of points with test statistic above T is at least a constant factor more than it “should” be. If such a threshold is found, those points are mostly outliers and can be safely removed. The key goal is to judiciously design a test statistic such that either the outliers aren’t particularly important—so the naive empirical solution is adequate—or at least one point will be filtered out.

In other words, the goal is to find a test statistic such that, if the distribution of the test statistic is “close” to what it would be in the outlier-free world, then the outliers cannot perturb the answer too much. An additional complication is that the test statistics depend on the data (such as $v \cdot x$, where v is the principal component of the data) making the distribution on inliers also nontrivial. This consideration drives the sample complexity of the algorithms.

In the algorithms we describe below, we use a specific parameterized notion of a good set. We define these precisely in the supplementary material, briefly, any large enough sample drawn from the uncorrupted distribution will satisfy the structural properties required for the set to be good.

We now describe how to design such test statistics for our two sparse settings.

Notation Before we describe our algorithms, we set up some notation. We define $h_k : \mathbb{R}^d \rightarrow \mathbb{R}^d$ to be the thresholding operator that keeps the k entries of v with the largest magnitude and sets the rest to 0. For a finite set S , we will use $a \in_u S$ to mean that a is chosen uniformly at random from S . For $M \in \mathbb{R}^d \times \mathbb{R}^d$ and $U \subseteq [d]$, let M_U denote the matrix M restricted to the $U \times U$ submatrix.

Robust Sparse Mean Estimation. Here we briefly describe the motivation and analysis of Algorithm 1, describing a single iteration of our filter for the robust sparse mean setting.

In order to estimate the k -sparse mean μ , it suffices to ensure that our estimate μ' has $|v \cdot (\mu' - \mu)|$ small for any $2k$ -sparse unit vector v . The now-standard idea in robust statistics [DKK⁺16] is that if a small number of corrupted samples suffice to cause a large change in our estimate of $v \cdot \mu$, then this must lead to a substantial increase in the sample variance of $v \cdot x$, which we can detect.

Thus, a very basic form of a robust algorithm might be to compute a sample covariance matrix $\tilde{\Sigma}$, and let v be the $2k$ -sparse unit vector that maximizes $v^T \tilde{\Sigma} v$. If this number is close to 1, it certifies that our estimate μ' — obtained by truncating the sample mean to its k -largest entries — is a good estimate of the true mean μ . If not, this will allow us to filter our sample set by throwing away the values where $v \cdot x$ is furthest from the true mean. This procedure guarantees that we have removed more corrupted samples than uncorrupted ones. We then repeat the filter until the empirical variance in every sparse direction is close to 1.

Unfortunately, the optimization problem of finding the optimal v is computationally challenging, requiring a convex program. To circumvent the need for a convex program, we notice that $v^T \tilde{\Sigma} v - 1 = (\tilde{\Sigma} - I) \cdot (vv^T)$ is large only if $\tilde{\Sigma} - I$ has large entries on the $(2k)^2$ non-zero entries of vv^T . Thus, if the $4k^2$ largest entries of $\tilde{\Sigma} - I$ had small ℓ_2 -norm, this would certify that no such bad v existed and would allow us to return the truncated sample mean. In case these entries have large ℓ_2 -norm, we show that we can produce a filter that removes more bad samples than good ones. Let A be the matrix consisting of the large entries of $\tilde{\Sigma}$ (for the moment assume that they are all off diagonal, but this is not needed). We know that the sample mean of $p(x) = (x - \mu')^T A (x - \mu') = \tilde{\Sigma} \cdot A = \|A\|_F^2$. On the other hand, if μ' approximates μ on the $O(k^2)$ entries in question, we would have that $\|p\|_2 = \|A\|_F$. This means that if $\|A\|_F$ is reasonably large, an ε -fraction of corrupted points changed the mean of p from 0 to $\|A\|_F^2 = \|A\|_F \|p\|_2$. This means that many of these errors must have had $|p(x)| \ll \|A\|_F / \varepsilon \|p\|_2$. This becomes very unlikely for good samples if $\|A\|_F$ is much larger than ε (by standard results on the concentration of Gaussian polynomials). Thus, if μ' is

approximately μ on these $O(k^2)$ coordinates, we can produce a filter. To ensure this, we can use existing filter-based algorithms to approximate the mean on these $O(k^2)$ coordinates. This results in Algorithm 1. For the analysis, we note that if the entries of A are small, then $v^T(\tilde{\Sigma} - I)v$ must be small for any unit k -sparse v , which certifies that the truncated sample mean is good. Otherwise, we can filter the samples using the first kind of filter. This ensures that our mean estimate is sufficiently close to the true mean that we can then filter using the second kind of filter.

It is not hard to show that the above works if we are given sufficiently many samples, but to obtain a tight analysis of the sample complexity, we need a number of subtle technical ideas. The detailed analysis of the sample complexity is deferred to the full version of our paper.

Robust Sparse PCA Here we briefly describe the motivation and analysis of Algorithm 2, describing a single iteration of our filter for the sparse PCA setting.

Note that estimating the k -sparse vector v is equivalent to estimating $\mathbf{E}[XX^T - I] = vv^T$. In fact, estimating $\mathbf{E}[XX^T - I]$ to error ε in Frobenius norm allows one to estimate v within error ε in ℓ_2 -norm. Thus, we focus on the task of robustly approximating the mean of $Y = XX^T - I$.

Our algorithm is going to take advantage of one fact about X that even errors cannot hide: that $\mathbf{Var}[v \cdot X]$ is large. This is because removing uncorrupted samples cannot reduce the variance by much more than an ε -fraction, and adding samples can only increase it. This means that an adversary attempting to fool our algorithm can only do so by creating other directions where the variance is large, or simply by adding other large entries to the sample covariance matrix in order to make it hard to find this particular k -sparse eigenvector. In either case, the adversary is creating large entries in the empirical mean of Y that should not be there. This suggests that the largest entries of the empirical mean of Y , whether errors or not, will be of great importance.

These large entries will tell us where to focus our attention. In particular, we can find the k^2 largest entries of the empirical mean of Y and attempt to filter based on them. When we do so, one of two things will happen: Either we remove bad samples and make progress or we verify that these entries ought to be large, and thus must come from the support of v . In particular, when we reach the second case, since the adversary cannot shrink the empirical variance of $v \cdot X$ by much, almost all of the entries on the support of v must remain large, and this can be captured by our algorithm.

The above algorithm works under a set of deterministic conditions on the good set of samples that are satisfied with high probability with $\text{poly}(k) \log(d)/\varepsilon^2$ samples. Our current analysis does not establish the information-theoretically optimal sample size of $O(k^2 \log(d)/\varepsilon^2)$, though we believe that this is plausible via a tighter analysis.

We note that a naive implementation of this algorithm will achieve error $\text{poly}(\varepsilon)$ in our final estimate for v , while our goal is to obtain $\tilde{O}(\varepsilon)$ error. To achieve this, we need to overcome two difficulties: First, when trying to filter Y on subsets of its coordinates, we do not know the true variance of Y , and thus cannot expect to obtain $\tilde{O}(\varepsilon)$ error. This is fixed with a bootstrapping method similar to that in [Kan18] to estimate the covariance of a Gaussian. In particular, we do not know $\mathbf{Var}[Y]$ a priori, but after we run the algorithm, we obtain an approximation to v , which gives an approximation to $\mathbf{Var}[Y]$. This in turn lets us get a better approximation to v and a better approximation to $\mathbf{Var}[Y]$; and so on.

3 Preliminaries

We will use the following notation and definitions.

Basic Notation For $n \in \mathbb{Z}_+$, let $[n] \stackrel{\text{def}}{=} \{1, 2, \dots, n\}$. Throughout this paper, for $v = (v_1, \dots, v_d) \in \mathbb{R}^d$, we will use $\|v\|_2$ to denote its Euclidean norm. If $M \in \mathbb{R}^{d \times d}$, we will use $\|M\|_2$ to denote its spectral norm, $\|M\|_F$ to denote its Frobenius norm, and $\text{tr}[M]$ to denote its trace. We will also let $\preceq$ and $\succeq$ denote the PSD ordering on matrices. For a finite multiset S , we will write $X \in_u S$ to denote that X is drawn from the empirical distribution defined by S . Given finite multisets S and S' we let $\Delta(S, S')$ be the size of the symmetric difference of S and S' divided by the cardinality of S .

For $v \in \mathbb{R}^d$ and $S \subseteq [d]$, let v_S be the vector with $(v_S)_i = v_i$, $i \in S$, and $(v_S)_i = 0$ otherwise. We denote by $h_k(v)$ the thresholding operator that keeps the k entries of v with largest magnitude (breaking ties arbitrarily) and sets the rest to 0. For $M \in \mathbb{R}^{d \times d}$ and $U \subseteq [d]$, let M_U denote the matrix M restricted to the $U \times U$ sub-matrix. For $W \subseteq [d] \times [d]$, then we will use $M_{(W)}$ to denote the matrix M restricted to the elements whose entries are in W .

Let δ_{ij} denote the Kronecker delta function. We will denote $\text{erfc}(z) = (2/\sqrt{\pi}) \int_z^\infty e^{-t^2} dt$. The notation $\tilde{O}(\cdot)$ and $\tilde{\Omega}(\cdot)$ hides logarithmic factors in the argument.

4 Robust Sparse Mean Estimation

Algorithm 1 Robust Sparse Mean Estimation via Iterative Filtering

1: **procedure** ROBUST-SPARSE-MEAN(S, k, ε, τ)
input: A multiset S such that there exists an (ε, k, τ) -good set G with $\Delta(G, S) \leq 2\varepsilon$.
output: Multiset S' or vector $\hat{\mu}$ satisfying Proposition 4.4.

2: Compute the sample mean $\tilde{\mu} = \mathbf{E}_{X \in_u S}[X]$ and the sample covariance matrix $\tilde{\Sigma}$, i.e., $\tilde{\Sigma} = (\tilde{\Sigma}_{i,j})_{1 \leq i, j \leq d}$ with $\tilde{\Sigma}_{i,j} = \mathbf{E}_{X \in_u S}[(X_i - \tilde{\mu}_i)(X_j - \tilde{\mu}_j)]$.

3: Let $U \subseteq [d] \times [d]$ be the set of the k largest magnitude entries of the diagonal of $\tilde{\Sigma} - I$ and the largest magnitude $k^2 - k$ off-diagonal entries, with ties broken so that if $(i, j) \in U$ then $(j, i) \in U$.

4: **if** $\|(\tilde{\Sigma} - I)_{(U)}\|_F \leq O(\varepsilon \log(1/\varepsilon))$ **then return** $\hat{\mu} := h_k(\tilde{\mu})$.

5: Set $U' = \{i \in [d] : (i, j) \in U\}$.

6: Compute the largest eigenvalue λ^* of $(\tilde{\Sigma} - I)_{U'}$ and a corresponding unit eigenvector v^* .

7: **if** $\lambda^* \geq \Omega(\varepsilon \sqrt{\log(1/\varepsilon)})$ **then:** Let $\delta_\ell := 3\sqrt{\varepsilon \lambda^*}$. Find $T > 0$ such that

$$\Pr_{X \in_u S}[|v^* \cdot (X - \tilde{\mu})| \geq T + \delta_\ell] \geq 9 \cdot \text{erfc}(T/\sqrt{2}) + \frac{3\varepsilon^2}{T^2 \ln(k \ln(Nd/\tau))}.$$

8: **return** the multiset $S' = \{x \in S : |v^* \cdot (x - \tilde{\mu})| \leq T + \delta_\ell\}$.

9: Let $p(x) = \left((x - \tilde{\mu})^T (\tilde{\Sigma} - I)_{(U)}^T (x - \tilde{\mu}) - \text{Tr}((\tilde{\Sigma} - I)_{(U)}) \right) / \|(\tilde{\Sigma} - I)_{(U)}\|_F$.

10: Find $T > 6$ such that

$$\Pr_{X \in_u S}[|p(X)| \geq T] \geq 9 \exp(-T/4) + 3\varepsilon^2/(T \ln^2 T).$$

11: **return** the multiset $S' = \{x \in S : |p(x)| \leq T\}$.

In this section, we prove correctness of Algorithm 1 establishing Theorem 1.2. For completeness, we restate a formal version of this theorem:

Theorem 4.1. *Let $D \sim \mathcal{N}(\mu, I)$ be an identity covariance Gaussian distribution on $\mathbb{R}^d$ with unknown k -sparse mean vector μ , and $\varepsilon, \tau > 0$. Let S be an ε -corrupted set of samples from D of size*

$N = \tilde{\Omega}(k^2 \log(d/\tau)/\varepsilon^2)$. There exists an efficient algorithm that, on input S , k , ε , and τ , returns a mean vector $\hat{\mu}$ such that with probability at least $1 - \tau$ it holds $\|\hat{\mu} - \mu\|_2 = O(\varepsilon \sqrt{\log(1/\varepsilon)})$.

4.1 Proof of Theorem 4.1

In this section, we describe and analyze our algorithm establishing Theorem 4.1. We start by formalizing the set of deterministic conditions on the good data under which our algorithm succeeds:

Definition 4.2. Fix $0 < \varepsilon, \tau < 1$ and $k \in \mathbb{Z}_+$. A multiset G of points in $\mathbb{R}^d$ is (ε, k, τ) -good with respect to $\mathcal{N}(\mu, I)$ if, for $X \in_u G$ and $Y \sim \mathcal{N}(\mu, I)$, the following conditions hold:

- (i) For all $i \in [d]$, $|\mathbf{E}_{X \in_u G}[X_i] - \mu_i| \leq \varepsilon/k$, and for all $i, j \in [d]$, $|\mathbf{E}_{X \in_u G}[(X_i - \mu_i)(X_j - \mu_j)] - \delta_{ij}| \leq \varepsilon/k$.
- (ii) For all $x \in G$ and $i \in [d]$, we have $|x_i - \mu_i| \leq O(\sqrt{\log(d|G|/\tau)})$.
- (iii) For all $2k^2$ -sparse unit vectors $v \in \mathbb{R}^d$, we have that:
 - (a) $|\mathbf{E}_{X \in_u G}[v \cdot (X - \mu)]| \leq O(\varepsilon)$,
 - (b) $|\mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2] - 1| \leq O(\varepsilon)$, and
 - (c) For all $T \geq 6$, $\Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \leq 3 \cdot \operatorname{erfc}(T/\sqrt{2}) + \varepsilon^2 / (T^2 \ln(k \ln(d|G|/\tau)))$.
- (iv) For all homogeneous* degree-2 polynomials p with $\mathbf{Var}_{\mathcal{N}(\mu, I)}[p(Y)] = 1$ and at most k^2 terms, we have that:
 - (a) $|\mathbf{E}_{X \in_u G}[p(X)] - \mathbf{E}_{\mathcal{N}(\mu, I)}[p(Y)]| \leq O(\varepsilon \sqrt{\mathbf{Var}_{\mathcal{N}(\mu, I)}[p(Y)]}) = O(\varepsilon)$, and
 - (b) For all $T \geq 5$, $\Pr_{X \in_u G}[|p(X) - \mathbf{E}_{\mathcal{N}(\mu, I)}[p(Y)]| \geq T] \leq 3 \exp(-T/4) + \varepsilon^2 / (T \ln^2 T)$.

Our first lemma says that a sufficiently large set of samples from $\mathcal{N}(\mu, I)$ is good with high probability:

Lemma 4.3. A set of $N = \tilde{\Omega}(k^2 \log(d/\tau)/\varepsilon^2)$ samples from $\mathcal{N}(\mu, I)$ is (ε, k, τ) -good (with respect to $\mathcal{N}(\mu, I)$) with probability at least $1 - \tau$.

Our algorithm iteratively applies the procedure ROBUST-SPARSE-MEAN (Algorithm 1). The crux of the proof is the following performance guarantee of ROBUST-SPARSE-MEAN:

Proposition 4.4. Algorithm 1 has the following performance guarantee: On input a multiset S of N points in $\mathbb{R}^d$ such that $\Delta(G, S) \leq 2\varepsilon$, where G is an (ε, k, τ) -good set with respect to $\mathcal{N}(\mu, I)$, procedure ROBUST-SPARSE-MEAN returns one of the following:

1. A mean vector $\hat{\mu}$ such that $\|\hat{\mu} - \mu\|_2 = O(\varepsilon \sqrt{\log(1/\varepsilon)})$, or
2. A multiset $S' \subset S$ satisfying $\Delta(G, S') \leq \Delta(G, S) - \varepsilon/N$.

We note that our overall algorithm terminates after at most $2N$ iterations of Algorithm 1, in which case it returns a candidate mean vector satisfying the first condition of Proposition 4.4. Note that the initial ε -corrupted set S satisfies $\Delta(G, S) \leq 2\varepsilon$. If $S^{(i)} \subset S$ is the multiset returned after the i -th iteration, then we have that $0 \leq \Delta(S^{(i)}, G) \leq 2\varepsilon - i(\varepsilon/N)$.

In the rest of this section, we prove Proposition 4.4.

We start by showing the first part of Proposition 4.4. Note that Algorithm 1 outputs a candidate mean vector only if $\|(\hat{\Sigma} - I)_{(U)}\|_F \leq O(\varepsilon \log(1/\varepsilon))$. We start with the following lemma:

*Recall that a degree- d polynomial is called homogeneous if its non-zero terms are all of degree exactly d .

Lemma 4.5. *If $\|(\tilde{\Sigma} - I)_{(U)}\|_F \leq O(\varepsilon \log(1/\varepsilon))$, then for any $T \subseteq [d]$ with $|T| \leq k$, we have $\|\tilde{\mu}_T - \mu_T\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$.*

Proof. Fix $T \subseteq [d]$ with $|T| \leq k$. By definition, $\|(\tilde{\Sigma} - I)_T\|_F$ is the Frobenius norm of the corresponding sub-matrix on $T \times T$. Note that this is the ℓ_2 -norm of a set of k diagonal entries and $k^2 - k$ off-diagonal entries of $\tilde{\Sigma} - I$. By construction, U is the set that maximizes this norm, and therefore

$$\|(\tilde{\Sigma} - I)_T\|_2 \leq \|(\tilde{\Sigma} - I)_T\|_F \leq \|(\tilde{\Sigma} - I)_{(U)}\|_F \leq O(\varepsilon \log(1/\varepsilon)).$$

Given this bound, we leverage a proof technique from [DKK⁺16] showing that a bound on the spectral norm of the covariance implies a ℓ_2 -error bound on the mean. This implication is not explicitly stated in [DKK⁺16], but follows directly from the arguments in Section 5.1.2 of that work. In particular, the analysis of the “small spectral norm” case in that section shows that $\|\tilde{\mu}_T - \mu_T\|_2 \leq O\left(\sqrt{\varepsilon}\|(\tilde{\Sigma} - I)_T\|_2^{1/2} + \varepsilon\sqrt{\log(1/\varepsilon)}\right)$, from which the desired claim follows. This completes the proof of Lemma 4.5. $\square$

Given Lemma 4.5, the correctness of the sparse mean approximation output in Step 4 of Algorithm 1 follows from the following corollary:

Corollary 4.6. *Let $\hat{\mu} = h_k(\tilde{\mu})$. If $\|(\tilde{\Sigma} - I)_{(U)}\|_F \leq O(\varepsilon \log(1/\varepsilon))$, then $\|\hat{\mu} - \mu\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$.*

Proof. For vectors x, y , let N_x denote the set of coordinates on which x is non-zero and $N_{x|y}$ denote the set of coordinates on which x is non-zero and y is zero. Setting $T = N_\mu$ and $T = N_{\hat{\mu}|\mu}$ in Lemma 4.5, we get that $\|\tilde{\mu}_{N_\mu} - \mu\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$ and $\|\tilde{\mu}_{N_{\hat{\mu}|\mu}}\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$.

If $\tilde{\mu}$ has k or fewer non-zero coordinates, then $\hat{\mu} = \tilde{\mu}$ and $\|\hat{\mu} - \mu\|_2 = \|\tilde{\mu}_{N_\mu \cup N_{\hat{\mu}|\mu}} - \mu\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$ and we are done. Otherwise, $\hat{\mu}$ has exactly k non-zero coordinates and so $|N_{\hat{\mu}|\mu}| \leq |N_{\hat{\mu}|\mu}|$. Since the nonzero coordinates of $\hat{\mu}$ are the k largest magnitude coordinates of $\tilde{\mu}$, for any $i \in N_{\mu|\hat{\mu}}$ and $j \in N_{\hat{\mu}|\mu}$, we have that $|\tilde{\mu}_i| \leq |\tilde{\mu}_j|$. Since $\|\tilde{\mu}_{N_{\hat{\mu}|\mu}}\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$, at least one coordinate $j \in N_{\hat{\mu}|\mu}$ must have $\tilde{\mu}_j^2 \leq O(\varepsilon^2 \log(1/\varepsilon))/|N_{\hat{\mu}|\mu}|$. Therefore, for any $i \in N_{\mu|\hat{\mu}}$, we have that $\tilde{\mu}_i^2 \leq O(\varepsilon^2 \log(1/\varepsilon))/|N_{\hat{\mu}|\mu}|$.

Thus, we have

$$\|\tilde{\mu}_{N_{\mu|\hat{\mu}}}\|_2^2 = \sum_{i \in N_{\mu|\hat{\mu}}} \tilde{\mu}_i^2 \leq \frac{|N_{\mu|\hat{\mu}}| \cdot O(\varepsilon^2 \log(1/\varepsilon))}{|N_{\hat{\mu}|\mu}|} \leq O(\varepsilon^2 \log(1/\varepsilon)),$$

where the second inequality used that $|N_{\mu|\hat{\mu}}| \leq |N_{\hat{\mu}|\mu}|$.

Since $\|\tilde{\mu}_{N_\mu} - \mu\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$, by the triangle inequality we have that $\|\mu_{N_{\mu|\hat{\mu}}}\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$. Finally, we have that

$$\|\mu - \hat{\mu}\|_2^2 = \|\mu_{N_\mu \cap N_{\hat{\mu}}} - \hat{\mu}_{N_\mu \cap N_{\hat{\mu}}}\|_2^2 + \|\mu_{N_{\mu|\hat{\mu}}}\|_2^2 + \|\tilde{\mu}_{N_{\hat{\mu}|\mu}}\|_2^2 \leq O(\varepsilon^2 \log(1/\varepsilon)),$$

concluding the proof. $\square$

Lemma 4.5 and Corollary 4.6 give the first part of Proposition 4.4.

We now analyze the complementary case that $\|(\tilde{\Sigma} - I)_{(U)}\|_F = \Omega(\varepsilon \log(1/\varepsilon))$. In this case, we apply two different filters, a linear filter (Steps 5-8), and a quadratic filter (Steps 9-11). To prove the second part of Proposition 4.4, we will show that at least one of these two filters: (i) removes at least one point, and (ii) it removes more corrupted than uncorrupted points.

The analysis in the case of the linear filter follows by a reduction to the linear filter in [DKK⁺16] for the non-sparse setting (see Proposition 5.5 in Section 5.1 of that work). More specifically, the linear filter in Steps 5-8 is essentially identical to the linear filter of [DKK⁺16] restricted to the $2k^2 \times 2k^2$ matrix $\tilde{\Sigma}_{U'}$. We note that Definition 4.2 implies that every restriction to $2k^2$ coordinates satisfies the properties of the good set in the sense of [DKK⁺16] (Definition 5.2(i)-(ii) of that work). This implies that the analysis of the linear filter from [DKK⁺16] holds in our case, establishing the desired properties. Since the linear filter removes more corrupted points than uncorrupted points, it will remove at most a 2ε fraction of the points over all the iterations.

If the condition of the linear filter does not apply, i.e., if $\|(\tilde{\Sigma} - I)_{U'}\|_2 \leq O(\varepsilon \log(1/\varepsilon))$, the aforementioned analysis in [DKK⁺16] implies $\|\tilde{\mu}_{U'} - \mu_{U'}\|_2 \leq O(\varepsilon \sqrt{\log(1/\varepsilon)})$. In this case, we show that the second filter behaves appropriately.

Let $p(x)$ be the polynomial considered in the quadratic filter. We start with the following technical lemma analyzing the expectation and variance of $p(x)$ under various distributions:

Lemma 4.7. *The following hold true:*

- (i) $\mathbf{E}_{\mathcal{N}(\tilde{\mu}, I)}[p(Y)] = 0$ and $\mathbf{Var}_{\mathcal{N}(\tilde{\mu}, I)}[p(Y)] = 1$.
- (ii) $\mathbf{E}_S[p(X)] = \|(\tilde{\Sigma} - I)_{(U)}\|_F$.
- (iii) $|\mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)]| \leq O(\varepsilon^2 \log(1/\varepsilon))$ and $\mathbf{Var}_{\mathcal{N}(\mu, I)}[p(Z)] = 1 + O(\varepsilon^2 \log(1/\varepsilon))$.

Proof. Let $A = \frac{(\tilde{\Sigma} - I)}{\|(\tilde{\Sigma} - I)\|_F}$ and $p(x) := (x - \tilde{\mu})^T A_{(U)}(x - \tilde{\mu}) - \text{Tr}[A_{(U)}]$. We have

$$\mathbf{E}_{\mathcal{N}(\tilde{\mu}, I)}[(Y - \tilde{\mu})^T A_{(U)}(Y - \tilde{\mu})] = \text{Tr}[A_{(U)} \mathbf{E}_{\mathcal{N}(\tilde{\mu}, I)}[(Y - \tilde{\mu})(Y - \tilde{\mu})^T]] = \text{Tr}[A_{(U)} I] = \text{Tr}[A_{(U)}] .$$

Therefore, $\mathbf{E}_{\mathcal{N}(\tilde{\mu}, I)}[p(Y)] = \text{Tr}[A_{(U)}] - \text{Tr}[A_{(U)}] = 0$. Similarly,

$$\mathbf{E}_S[(X - \tilde{\mu})^T A_{(U)}(X - \tilde{\mu})] = \mathbf{E}_S[\text{Tr}[A_{(U)}(X - \tilde{\mu})(X - \tilde{\mu})^T]] = \text{Tr}[A_{(U)} \mathbf{E}_S[(X - \tilde{\mu})(X - \tilde{\mu})^T]] = \text{Tr}[A_{(U)} \tilde{\Sigma}] ,$$

and so

$$\mathbf{E}_S[p(X)] = \text{Tr}[A_{(U)}(\tilde{\Sigma} - I)] = \text{Tr}[A_{(U)} A] \|(\tilde{\Sigma} - I)_{(U)}\|_F = \|A_U\|_F \|(\tilde{\Sigma} - I)_{(U)}\|_F = \|(\tilde{\Sigma} - I)_{(U)}\|_F .$$

We have thus shown (ii) and the first part of (i).

We now proceed to show the first part of (iii). Note that

$$\begin{aligned} \mathbf{E}_{\mathcal{N}(\mu, I)}[(Z - \tilde{\mu})(Z - \tilde{\mu})^T] &= \mathbf{E}_{\mathcal{N}(\mu, I)}[(Z - \mu)(Z - \mu)^T] + \mathbf{E}_{\mathcal{N}(\mu, I)}[(\tilde{\mu} - \mu)(Z - \tilde{\mu})^T] + \mathbf{E}_{\mathcal{N}(\mu, I)}[(Z - \mu)(\tilde{\mu} - \mu)^T] \\ &= I + (\tilde{\mu} - \mu)(\tilde{\mu} - \mu)^T + 0. \end{aligned}$$

Thus, we can write

$$\mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)] = \text{Tr}[A_{(U)}(\mathbf{E}_{\mathcal{N}(\mu, I)}[(Z - \tilde{\mu})(Z - \tilde{\mu})^T] - I)] = (\tilde{\mu} - \mu)^T A_{(U)}(\tilde{\mu} - \mu) ,$$

and so

$$|\mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)]| \leq \|\tilde{\mu}_{U'} - \mu_{U'}\|_2^2 \|A_{(U)}\|_2 \leq \|\tilde{\mu}_{U'} - \mu_{U'}\|_2^2 \leq O(\varepsilon^2 \log(1/\varepsilon)) .$$

This proves all the statements about expectations.

We now analyze the variance of $p(x)$ for Y and Z . Since $A_{(U)}$ is symmetric, we can write $A_{(U)} = O^T \Lambda O$ for an orthogonal matrix O and a diagonal matrix Λ . Note that $Y' = O(Y - \tilde{\mu})$ is

distributed as $\mathcal{N}(0, I)$. Under these substitutions, $p(Y) = \sum_i \Lambda_{ii} Y_i'^2 - \sum \Lambda_{ii}$, the variance of which is the same as the variance of $p'(Y) = \sum_i \Lambda_{ii} Y_i'^2$ and so

$$\mathbf{Var}_{\mathcal{N}(\tilde{\mu}, I)}[p'(Y)] = \sum_i \Lambda_{ii}^2 \mathbf{Var}_{\mathcal{N}(0, I)}[Y_i'^2] = \|\Lambda\|_F^2 = \|A_{(U)}\|_F^2 = 1.$$

Similarly, to estimate the variance of $p(Z)$ we see we just need to estimate the variance of $p'(Z) = (Z - \tilde{\mu})^T O^T \Lambda O (Z - \tilde{\mu}) = (Z - \mu + \mu - \tilde{\mu})^T O^T \Lambda O (Z - \mu + \mu - \tilde{\mu})$ where $(Z - \mu) \sim N(0, I)$. Let $Z' = O^T(Z - \mu)$ and $\nu = O(\mu - \tilde{\mu})$.

This gives us

$$\begin{aligned} \mathbf{Var}_{\mathcal{N}(\mu, I)}[p'(Z)] &= \mathbf{E}_{\mathcal{N}(0, I)}[(Z' \Lambda Z' - \mathbf{E}_{\mathcal{N}(0, I)}[Z' \Lambda Z'] + 2\nu^T \Lambda (Z' - \nu))^2] \\ &= \mathbf{E}_{\mathcal{N}(0, I)}[(Z' \Lambda Z' - \mathbf{E}_{\mathcal{N}(0, I)}[Z' \Lambda Z'])^2] + 4\mathbf{E}_{\mathcal{N}(0, I)}[(\nu^T \Lambda (Z' - \nu))^2] \\ &\quad + 4\mathbf{E}_{\mathcal{N}(0, I)}[(\nu^T \Lambda (Z' - \nu))(Z' \Lambda Z' - \mathbf{E}_{\mathcal{N}(0, I)}[Z' \Lambda Z'])] \\ &\leq 1 + 8\mathbf{E}_{\mathcal{N}(0, I)}[(\nu^T \Lambda (Z' - \nu))^2] \\ &\leq 1 + 8\nu^T \Lambda^2 \nu + 8(\nu^T \Lambda \nu)^2 \\ &\leq 1 + 8\|\tilde{\mu}_{U'} - \mu_{U'}\|_2^2 \|\Lambda^2\|_2 + 8\|\tilde{\mu}_{U'} - \mu_{U'}\|_2^4 \|\Lambda\|_2^2 \\ &\leq 1 + O(\varepsilon^2 \log(1/\varepsilon)) \end{aligned}$$

This completes the proof of Lemma 4.7. $\square$

Suppose that we find a threshold $T > 0$ such that Step 10 of the algorithm holds, i.e., the quadratic filter applies. Then we can show that Step 11 removes more bad points than good points. This follows from standard arguments, by combining Definition 4.2(iv)(b) with our upper bound for $\mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)]$ from Lemma 4.7. Let $S = G \cup E \setminus L$. By Definition 4.2(iv)(b), for the good set G , we have that for $T > 4$, $\mathbf{Pr}_{X \in_u G} [|p(X) - \mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)]| \geq T] \leq 3 \exp(-T/4) + \varepsilon^2/(T \ln^2 T)$. Lemma 4.7(iii) implies that $|\mathbf{E}_{\mathcal{N}(\mu, I)}[p(Z)]| \leq O(\varepsilon^2 \log(1/\varepsilon))$. Therefore, shifting T by $\varepsilon^2 \log(1/\varepsilon)$ we obtain the following corollary:

Corollary 4.8. *We have that:*

- (i) $|\mathbf{E}_{X \in_u G}[p(X)]| \leq O(\varepsilon)$ and,
- (ii) For $T \geq 6$, $\mathbf{Pr}_{X \in_u G} [|p(X)| \geq T] \leq (3 + O(\varepsilon)) \exp(-T/4) + (1 + O(\varepsilon)) (\varepsilon^2/(T \ln^2 T))$.

Condition (ii) implies that the fraction of points in G that violate the quadratic filter condition is less than $1/2$ the fraction of points in S that violate the same condition. Therefore, the quadratic filter removes more corrupted points than uncorrupted ones.

It remains to show that if Algorithm 1 does not terminate in Step 4 and the linear filter does not apply, then the quadratic filter necessarily applies. To establish this, we need a couple more technical lemmas. We first show that the expectation of $p(x)$ over the set of good samples that are removed is small:

Lemma 4.9. *We have that $|L| \cdot |\mathbf{E}_{X \in_u L}[p(X)]| \leq |S| \cdot O(\varepsilon \log(1/\varepsilon))$.*

Proof. Since $L \subset G$ and $|G| = O(|S|)$, for $T \geq 6$ we have

$$|L| \cdot \mathbf{Pr}_{X \in_u L} [|p(X)| \geq T] \leq |G| \cdot \mathbf{Pr}_{X \in_u G} [|p(X)| \geq T] \leq O(|S|(\exp(-T/4) + \varepsilon^2/(T \ln^2 T))) ,$$

where we used Corollary 4.8. Thus, we obtain that for $\varepsilon < O(1)$.

$$\begin{aligned}
|L| \cdot |\mathbf{E}_{X \in_u L}[p(X)]| &\leq |L| \cdot \mathbf{E}_{X \in_u L}[|p(X)|] \\
&= \int_0^\infty |L| \cdot \Pr_{X \in_u L}[|p(X)| \geq T] dT \\
&\leq \int_0^{3 \ln(1/\varepsilon)} |L| dT + \int_{3 \ln(1/\varepsilon)}^\infty O(|S|(\exp(-T/4) + \varepsilon^2/(T \ln^2 T))) dT \\
&\leq O(|S|\varepsilon \log(1/\varepsilon)) + O(|S|\varepsilon) + O(|S|\varepsilon^2/\log \log(1/\varepsilon)) \\
&= O(|S|\varepsilon \log(1/\varepsilon)) ,
\end{aligned}$$

where we used the fact that $|L| = O(\varepsilon|S|)$ and that the derivative of $1/\ln x$ is $1/x \ln^2 x$. This completes the proof of Lemma 4.9. $\square$

By a similar argument, we can show that if the quadratic filter does not apply, then the remaining points in E contribute a small amount to the expectation of $p(x)$.

Lemma 4.10. *Suppose that for all $T \geq 6$, we have $\Pr_{X \in_u S}[|p(X)| \geq T] \leq 9 \exp(-T/4) + 3\varepsilon^2/(T \ln^2 T)$. Then, we have that $|E| \cdot |\mathbf{E}_{X \in_u E}[p(X)]| \leq O(|S|\varepsilon \log(1/\varepsilon))$.*

By combining the above, we obtain the following corollary, completing the analysis of our algorithm:

Corollary 4.11. *If we reach Step 10 of Algorithm 1, then there exists a $T \geq 6$ such that $\Pr_{X \in_u S}[|p(X)| \geq T] \geq 9 \exp(-T/4) + 3\varepsilon^2/(T \ln^2 T)$.*

Proof. Suppose for a contradiction that no such T exists. Using Corollary 4.8, Lemmas 4.9 and 4.10, we obtain that

$$\begin{aligned}
|S| \cdot \|\widetilde{\Sigma} - I\|_F &= |S| \cdot \mathbf{E}_S[p(X)] = |G| \cdot \mathbf{E}_{X \in_u G}[p(X)] + |E| \cdot \mathbf{E}_{X \in_u E}[p(X)] - |L| \cdot \mathbf{E}_{X \in_u L}[p(X)] \\
&= O(|S|\varepsilon \log(1/\varepsilon)) .
\end{aligned}$$

This is a contradiction, as if this was the case, Algorithm 1 would have returned in Step 4. $\square$

5 Robust Sparse PCA

In this section, we prove correctness of Algorithm 2 establishing Theorem 1.3, which we restate for completeness:

Theorem 5.1. *Let $D \sim \mathcal{N}(\mathbf{0}, I + \rho v v^T)$ be a centered Gaussian distribution on $\mathbb{R}^d$ with spiked covariance $\Sigma = I + \rho v v^T$ for an unknown k -sparse unit vector v , and $0 < \rho < O(1)$ a real number. For some $\varepsilon > 0$, let S be an ε -corrupted set of samples from D of size $N = \Omega(k^4 \log^4(d/\varepsilon)/\varepsilon^2)$. There exists an algorithm that, on input S , k , and ε , runs in polynomial time and returns $w \in \mathbb{R}^d$ such that with probability at least $2/3$ we have that $\|w w^T - v v^T\|_F = O\left(\frac{\varepsilon \log(1/\varepsilon)}{\rho}\right)$.*

We will require some additional notation. For any $M \in \mathbb{R}^{d \times d}$, define $\text{vec}(M) \in \mathbb{R}^{d^2}$ to be a canonical flattening of this vector and $\gamma(x) \in \mathbb{R}^{d^2}$ to be $\text{vec}(x x^T - I)$. Also let $\gamma_A(x) = \gamma(x_A)$ and $\text{vec}_B(M) = \text{vec}(M_B)$, where $A \subset [d] \times [d]$ and $x, M \in \mathbb{R}^{d \times d}$.

As is standard with such robust statistics arguments, we will need to assume that the uncorrupted set of good samples G has some desired properties. In particular, we will make use of the following notion of a good set:

Algorithm 2 Robust Sparse PCA via Iterative Filtering

1: **procedure** ROBUST-SPARSE-PCA($S, k, \tilde{\Sigma}, \varepsilon, \delta, \tau$)
input: A multiset S , an estimate of the true covariance $\tilde{\Sigma}$, a real number $\delta \in \mathbb{R}$.
output: A multiset S' or matrix Σ' satisfying Proposition 5.4.
2: For any $x \in \mathbb{R}^d$ define $\gamma(x) := \text{vec}(xx^T - I) \in \mathbb{R}^{d^2}$.
3: Compute $\tilde{\mu} := \mathbf{E}_S[\gamma(x)]$, $\hat{\mu} = h_{k^2}(\mu)$ and $Q := \text{Supp}(\hat{\mu})$.
4: Compute

$$M_Q := \mathbf{E}_S[(\gamma(x) - \tilde{\mu})(\gamma(x) - \tilde{\mu})^T]_{Q \times Q} \in \mathbb{R}^{k^2} \times \mathbb{R}^{k^2}$$

5: Let λ, v^* be the maximum eigenvalue and corresponding eigenvector of $M_Q - \text{Cov}_{X \sim \mathcal{N}(0, \tilde{\Sigma})}(\gamma(x)_Q)$.
6: **if** $\lambda < C \cdot (\delta + \varepsilon \log^2(1/\varepsilon))$, where C is a sufficiently large constant **then**
7: Compute w , the largest eigenvector of $\text{mat}(\tilde{\mu})_Q$. **return** $ww^T + I$.
8: Let $\hat{\mu} = \text{median}(\{\gamma(x) \cdot v^* \mid x \in S\})$. Find a number $T > \log(1/\varepsilon)$ satisfying

$$\Pr_S[|\gamma(x)_Q \cdot v^* - \hat{\mu}| > CT + 3] > \frac{\varepsilon}{T^2 \log^2(T)}.$$

return $S' = \{x \in S \mid |(\gamma(x)_Q \cdot v^*) - \hat{\mu}| < T\}$.

Definition 5.2. Define a set $G \subset \mathbb{R}^n$ to be (ε, k) -good for $\mathcal{N}(0, I + \rho vv^T)$ and $\rho > 0$ if the following hold for every $Q \subset [d] \times [d]$

1. For some sufficiently large constant C and for every $i \in [d]$ and $x \in G$, $|x_i| \leq C \sqrt{\log(d|G|)}$.
2. $\|(\mathbf{E}_G[xx^T] - I - \rho vv^T)_Q\|_F \leq \varepsilon$
3. For all $w \in \mathbb{R}^{k^2}$,

$$\mathbf{Var}_G[\gamma_Q(x) \cdot w] = (1 \pm \varepsilon) \mathbf{Var}_{\mathcal{N}(0, I + \rho vv^T)}[\gamma_Q(x) \cdot w]$$

4. For C a sufficiently large constant, and for all $w \in \mathbb{R}^{k^2}$ satisfying $\|w\|_2 = 1$, and all $T > \log(1/\varepsilon)$

$$\Pr_G[|\gamma_Q(x) \cdot w - \rho \text{vec}_Q(vv^T) \cdot w| > CT] < \frac{\varepsilon}{T^2 \log^2(T)}$$

We note that given a sufficiently large set of independent samples from X that the above conditions hold with high probability.

Lemma 5.3. If G is a set of $N = Ck^4 \log^4(d/\varepsilon)/\varepsilon^2$ samples drawn from $\mathcal{N}(0, I + \rho vv^T)$, for C a sufficiently large constant. Then G is (ε, k) -good with probability at least $2/3$.

We think in fact that we should be able to produce a good set with substantially fewer samples.

Conjecture 1. There exists an $N = k^2 \text{polylog}(d/\varepsilon)/\varepsilon^2$ so that if G is a set of N samples drawn from $\mathcal{N}(0, I + \rho vv^T)$, then G is (ε, k) -good with probability at least $1 - 1/d$.

We can now proceed with the proof of our main Theorem. In particular, our algorithm will follow quickly from the existence of the following subroutine:

Proposition 5.4. *Let G be an (ε, k) -good set for $\mathcal{N}(0, \Sigma)$ with $\Sigma = I + \rho v v^T$ with v a unit length, k -sparse vector and $0 < \rho < 1$. There exists an algorithm (Algorithm 2) that given a matrix $\tilde{\Sigma}$ and a set S with $\|\tilde{\Sigma} - \Sigma\|_F \leq \delta$ and $\Delta(S, G) \leq \varepsilon|G|$ returns either a matrix Σ' with $\|\Sigma' - \Sigma\|_F = O(\sqrt{\varepsilon\delta} + \varepsilon \log(1/\varepsilon))$ or a subset $T \subset S$ with $\Delta(T, G) < \Delta(S, G)$.*

Our main theorem follows from iteratively applying the Proposition. The error stabilizes at δ with $\delta = O(\sqrt{\varepsilon\delta} + \varepsilon \log(1/\varepsilon))$, which implies that $\delta = O(\varepsilon \log(1/\varepsilon))$. We begin by analyzing what happens when our algorithm returns a matrix. We first note that if we pass the filter, then $\tilde{\mu}_Q$ will be approximately correct.

Lemma 5.5. *With the notation as in Algorithm 2, we have that $\|\tilde{\mu}_Q - \text{vec}_Q(\Sigma - I)\|_2 = O(\sqrt{\varepsilon\lambda} + \sqrt{\varepsilon\delta} + \varepsilon \log(1/\varepsilon))$.*

Proof. Let $\|\tilde{\mu}_Q - \text{vec}_Q(\Sigma - I)\|_2 = a$ and $S = (G \setminus L) \cup E$. We wish to show that

$$\left\| \sum_{x \in S} (xx^T - \Sigma)_Q \right\|_2 = O(\sqrt{\varepsilon\delta} + \varepsilon \log(1/\varepsilon))|G|.$$

By the triangle inequality, the left hand side above is at most

$$\left\| \sum_{x \in G} (xx^T - \Sigma)_Q \right\|_2 + \left\| \sum_{x \in L} (xx^T - \Sigma)_Q \right\|_2 + \left\| \sum_{x \in E} (xx^T - \Sigma)_Q \right\|_2. \quad (1)$$

Since G is a good set, by Condition 2, we have that the first term is $O(\varepsilon|G|)$. We now bound the second term. Since $\Sigma = I + \rho v v^T$ and $\gamma(x) = \text{vec}(xx^T - I)$ the second term is at most the supremum over unit vectors $w \in \mathbb{R}^{k^2}$ of

$$\sum_{x \in L} (w \cdot \gamma_Q(x) - \rho w \cdot \text{vec}_Q(vv^T))$$

Using the fact that for any random variable $\mathbf{E}_D[X] = \int_0^\infty \mathbf{Pr}_D[X > t] dt$, this is at most

$$\int_0^\infty |\{x \in L : |(w \cdot \gamma_Q(x) - \rho w \cdot \text{vec}_Q(vv^T))| > t\}| dt.$$

Since $L \subset G$ and $|L| = \varepsilon|G|$ this is at most

$$\int_0^{C \log(1/\varepsilon)} \varepsilon|G| + \int_{C \log(1/\varepsilon)}^\infty \varepsilon / ((t/C)^2 \log^2(t/C)) |G| dt = O(\varepsilon \log(1/\varepsilon)|G|),$$

where the bound on the second term above is by Condition 4 of the definition of a good set.

We can bound the final term in 1 by Cauchy-Schwartz as

$$(\varepsilon|G|)^{1/2} \left(\sum_{x \in E} (w \cdot (xx^T - \Sigma)_Q)^2 \right)^{1/2}.$$

To bound this we note that

$$\sum_{x \in S} (w \cdot (xx^T - \Sigma)_Q)^2 \leq |S|(\mathbf{Var}_S[w \cdot \gamma_Q(x)] + a^2).$$

We also know that

$$\mathbf{Var}_G[w \cdot \gamma_Q(x)] = \mathbf{Var}_{\mathcal{N}(0, \rho v v^T + I)}[w \cdot \gamma_Q(x)] + O(\varepsilon) = \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)] + O(\varepsilon + \delta).$$

Thus subtracting both sides by $\mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)]$ and scaling by G gives us

$$\sum_{x \in G} ((w \cdot (x x^T - \Sigma)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)]) = O(\varepsilon + \delta)|G|.$$

However, since v^* is the eigenvector corresponding to the largest eigenvalue, we also have that

$$\mathbf{Var}_S[w \cdot \gamma_Q(x)] - \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)] \leq \lambda + a^2.$$

Combining with the above and using the fact that $S = (G \setminus L) \cup E$. we have that

$$\begin{aligned} \sum_{x \in E} ((w \cdot (x x^T - \Sigma)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)]) &\leq \sum_{x \in L} ((w \cdot (x x^T - \Sigma)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)]) \\ &\quad + |G|O(\varepsilon + \delta + \lambda + a^2). \end{aligned}$$

However, we can bound

$$\sum_{x \in L} ((w \cdot (x x^T - \Sigma)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \tilde{\Sigma})}[w \cdot \gamma_Q(x)])$$

by

$$O(|L|) + \int_0^\infty |\{x \in L : |(w \cdot \gamma_Q(x) - \rho w \cdot \text{vec}_Q(v v^T))| > t\}| 2t dt.$$

As before, this is at most

$$\int_0^{C \log(1/\varepsilon)} 2t\varepsilon |G| dt + \int_{C \log(1/\varepsilon)}^\infty \varepsilon / ((t/C)^2 \log^2(t/C)) |G| 2t dt = O(\varepsilon \log^2(1/\varepsilon)|G|).$$

Thus, the final term in our sum is at most

$$(\varepsilon |G|)^{1/2} O(\varepsilon \log^2(1/\varepsilon)|G| + (a^2 + \varepsilon + \delta + \lambda)|G|)^{1/2}$$

Therefore, we have that

$$a = O(a\sqrt{\varepsilon} + \sqrt{\varepsilon\lambda} + \sqrt{\varepsilon\delta} + \varepsilon \log(1/\varepsilon)),$$

from which we conclude our result. □

Given this, we would like to show that Σ' is close to Σ . In particular, we have:

Lemma 5.6. *Suppose that $A = \mathbf{E}_{x \in_u S}[x x^T - I]$ and Q the set of its k^2 largest entries. If $\|(A - \rho v v^T)_Q\|_F = \eta$ then for w a normalized, principle eigenvector of A_Q we have that ρw is within $O(\eta + \varepsilon \log(1/\varepsilon))$ of either ρv or $-\rho v$.*

Before we begin with the proof, we make an important observation:

Lemma 5.7. *In the notation above, for any set of entries R defining a $k^2 \times k^2$ submatrix, $(A+I)_R \geq ((\rho vv^T + I) - O(\varepsilon \log(1/\varepsilon))I)_R$, as self-adjoint operators.*

Proof. Note that $A + I = \mathbf{E}_{x \in_u S}[xx^T] \geq (1 - \varepsilon)\mathbf{E}_{x \in_u G \setminus L}[xx^T]$. By Property 2 of what it means to be a good set, $\mathbf{E}_{x \in_u G}[xx^T] = \rho vv^T + I + O(\varepsilon)$. Thus, it suffices to show that for any unit vector u with support of size at most k^2 that $|L|/|G|\mathbf{E}_{x \in_u L}[(x \cdot u)^2] = O(\varepsilon \log(1/\varepsilon))$. This follows easily from Property 4. $\square$

We are now ready to prove Lemma 5.6.

Proof. Let R be the support of vv^T . Note that A has larger total L^2 mass on Q than it does on R . Therefore,

$$\|A_{R \setminus Q}\|_F \leq \|A_{Q \setminus R}\|_F \leq \|(A - \rho vv^T)_Q\|_F = \eta.$$

Let $B = (\rho vv^T)_{R \setminus Q}$. We note that with respect to Frobenius norm:

$$A_Q = A_{Q \cap R} + A_{Q \setminus R} = A_{Q \cap R} + O(\eta) = (\rho vv^T)_{Q \cap R} + O(\eta) = (\rho vv^T - B) + O(\eta).$$

We also note that this is $A_{Q \cap R} + O(\eta) = A_R + O(\eta)$. Combining this with the above lemma, we have that

$$(\rho vv^T - B + I) + O(\eta) \geq (\rho vv^T + I) - O(\varepsilon \log(1/\varepsilon))I.$$

Rearranging, we find that $B \leq O(\eta + \varepsilon \log(1/\varepsilon))I$. But we note that the sign of the i, j entry of B is the same as the sign of $v_i v_j$ or 0. This means that B is similar to a matrix with non-negative entries, and thus by The Perron–Frobenius Theorem, the largest eigenvalue of B is positive, and hence $\|B\|_2 = O(\eta + \varepsilon \log(1/\varepsilon))$. Therefore, we have that

$$\|A_Q - \rho vv^T\|_2 \leq \|A_Q - (\rho vv^T - B)\|_2 + \|B\|_2 = O(\eta + \varepsilon \log(1/\varepsilon)).$$

Note that unless ε and η are sufficiently small, there is nothing to prove. Otherwise, we have that $v \cdot A_Q v \geq \rho - O(\eta + \varepsilon \log(1/\varepsilon))$, so w will be an eigenvector with some eigenvalue $\lambda > \rho/2$. Since $\|A_Q - \rho vv^T\|_2 < \rho/2$, this means that w must have a non-trivial component in the v -direction. Assume that w is proportional to $v + u$ with u orthogonal to v . Then we have that

$$\lambda(v + u) = \lambda w = A_Q w = A_Q(v + u) = \rho v + O(\eta + \varepsilon \log(1/\varepsilon)).$$

Taking the perpendicular to v component above, we have that $\|u\|_2 = O(\eta + \varepsilon \log(1/\varepsilon))$, and this completes our proof. $\square$

Finally, note that

$$\begin{aligned} \|vv^T - ww^T\|_F^2 &= \|vv^T\|_F^2 + \|ww^T\|_F^2 - 2\text{tr}(vv^T ww^T) \\ &= 2 - 2(v \cdot w)^2 \leq 2 - 2|v \cdot w| = \|v \pm w\|_2^2 \\ &= \frac{\|\rho v \pm \rho w\|_2^2}{\rho^2} \leq O\left(\left(\frac{\eta + \varepsilon \log(1/\varepsilon)}{\rho}\right)^2\right) \end{aligned}$$

Thus, plugging in $\eta = O(\sqrt{\varepsilon \lambda} + \sqrt{\varepsilon \delta} + \varepsilon \log(1/\varepsilon))$ above, we find that $\|vv^T - ww^T\|_F = O(\frac{\sqrt{\varepsilon \delta} + \varepsilon \log(1/\varepsilon)}{\rho})$.

We have left to analyze what happens when our algorithm returns a set S' . It is easy to see by Conditions 2 and 3 that only $1/3$ of the elements of G have $(\gamma(x)_Q - \rho \text{vec}(vv^T)_Q) \cdot v^* > 3$. Therefore, we have that $\hat{\mu}$ is within 3 of $\rho v^* \cdot (\text{vec}(vv^T)_Q)$. From this and Condition 4 it is easy to see that if C is sufficiently large (even compared to the C in Condition 4), that less than half of the elements of S with $|\text{vec}(xx^T) \cdot v^*| > CT + 3$ will be in G , and thus $\Delta(S', G) < \Delta(S, G)$.

All that remains is to show that such a threshold T exists. To do this consider

$$\mathbf{Var}_S[v^* \cdot \text{vec}(xx^T)_Q] - \mathbf{Var}_{\mathcal{N}(0, \Sigma)}[v^* \cdot \text{vec}(xx^T)_Q].$$

This is $O(\delta) + \lambda$. On the other hand, since translating a random variable should not change its variance, we see.

$$\begin{aligned} \mathbf{Var}_S[v^* \cdot \text{vec}(xx^T)_Q] &= \mathbf{E}_S[(v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2] - \mathbf{E}_S[v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q]^2 \\ &= \mathbf{E}_S[(v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2] + O(\varepsilon\lambda + \varepsilon\delta + \varepsilon^2 \log^2(1/\varepsilon)) \end{aligned}$$

by Lemma 5.5. Thus,

$$\mathbf{E}_S[(v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2] \geq \mathbf{Var}_{\mathcal{N}(0, \Sigma)}[v^* \cdot \text{vec}(xx^T)_Q] + \lambda/2.$$

Now by Conditions 2 and 3 we have that

$$\sum_{x \in G} ((v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \Sigma)}[v^* \cdot \text{vec}(xx^T)_Q]) = O(|G|\varepsilon).$$

By arguments from the proof of Lemma 5.5, we also have that

$$\sum_{x \in L} ((v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2 - \mathbf{Var}_{\mathcal{N}(0, \Sigma)}[v^* \cdot \text{vec}(xx^T)_Q]) = O(|G|\varepsilon \log^2(1/\varepsilon)).$$

Thus, we must have

$$\sum_{x \in E} (v^* \cdot \text{vec}(xx^T - \rho vv^T)_Q)^2 \gg |G|\lambda.$$

However, this is at most

$$O\left(|E| + \int_0^\infty |\{x \in E : |v^* \cdot \text{vec}(xx^T)_Q - \hat{\mu}| > CT + 3\}| t dt\right).$$

If there is no such threshold, this is at most

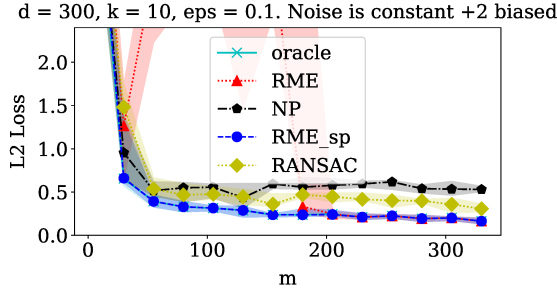
$$O\left(|E| + \int_0^{\log(1/\varepsilon)} |E| t dt + \int_{\log(1/\varepsilon)}^\infty \varepsilon/(t^2 \log^2(t)) t dt\right) = O(\varepsilon \log^2(1/\varepsilon)|G|),$$

which is a contradiction. This completes our proof.

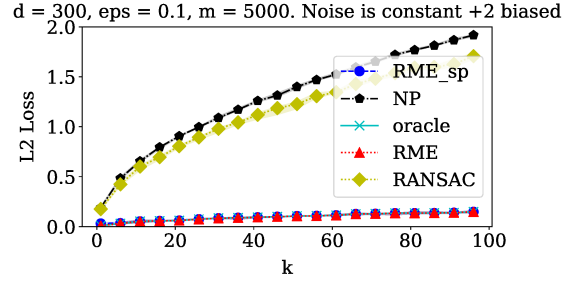
6 Experiments

For every experiment, we run 10 trials and plot the median value of the measurement. We shade the interquartile range around each measurement as a measure of the confidence of that measurement.

Each experiment was run on a computer with a 2.7 GHz Intel Core i5 processor with an 8GB 1867 MHz DDR3 RAM.

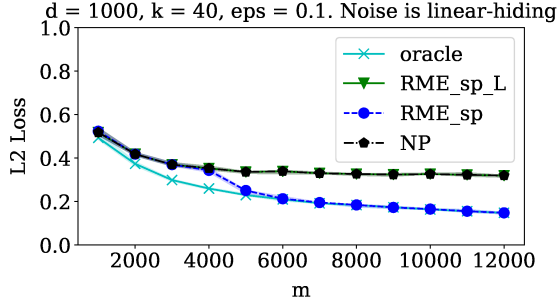


(a) Unlike RANSAC, our algorithm RME_sp can filter out the noise and match the oracle's performance. RME also matches the oracle, but needs more samples.

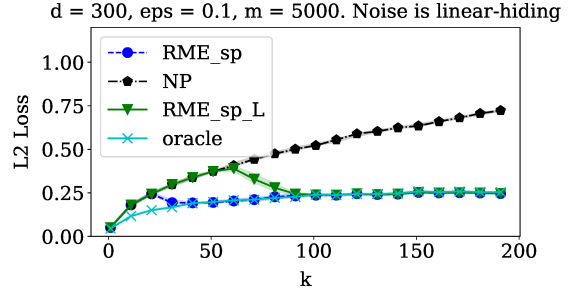


(b) For fixed m , as k increases, RANSAC and NP both diverge from RME and RME_sp.

Figure 1: Constant-bias noise is easy for our algorithm, since it is caught by the linear filter.

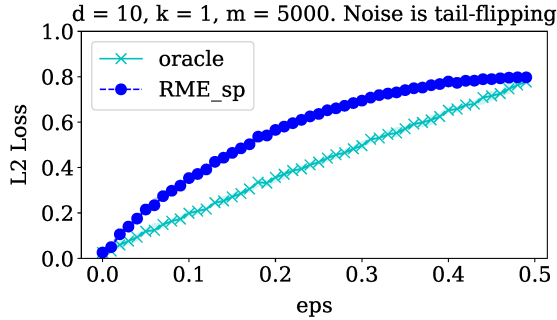


(a) With sufficiently many samples, the quadratic filter can filter out the noise, matching the oracle. The linear filter alone does not, even with a large number of samples.

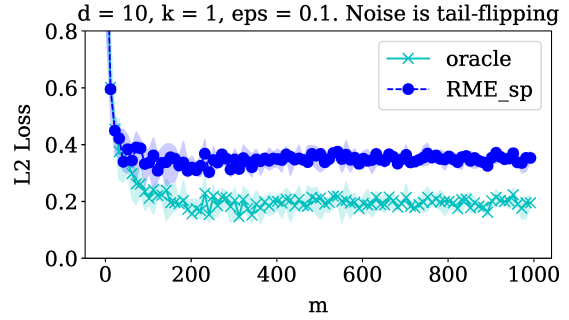


(b) For $k \ll \sqrt{d}$, the linear filter alone does not filter out the noise, leading to an $\varepsilon\sqrt{k}$ dependence for RME_sp.L. Our algorithm RME_sp nearly matches oracle.

Figure 2: The linear-hiding noise model shows that the quadratic filter is necessary.

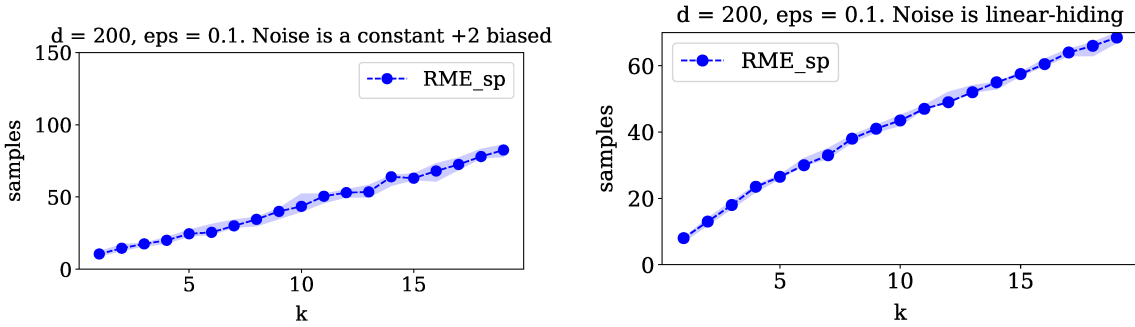


(a) This noise model gives $\Omega(\varepsilon\sqrt{\log(1/\varepsilon)})$ error to the oracle, and RME_sp is at most twice this.



(b) This gap persists regardless of m .

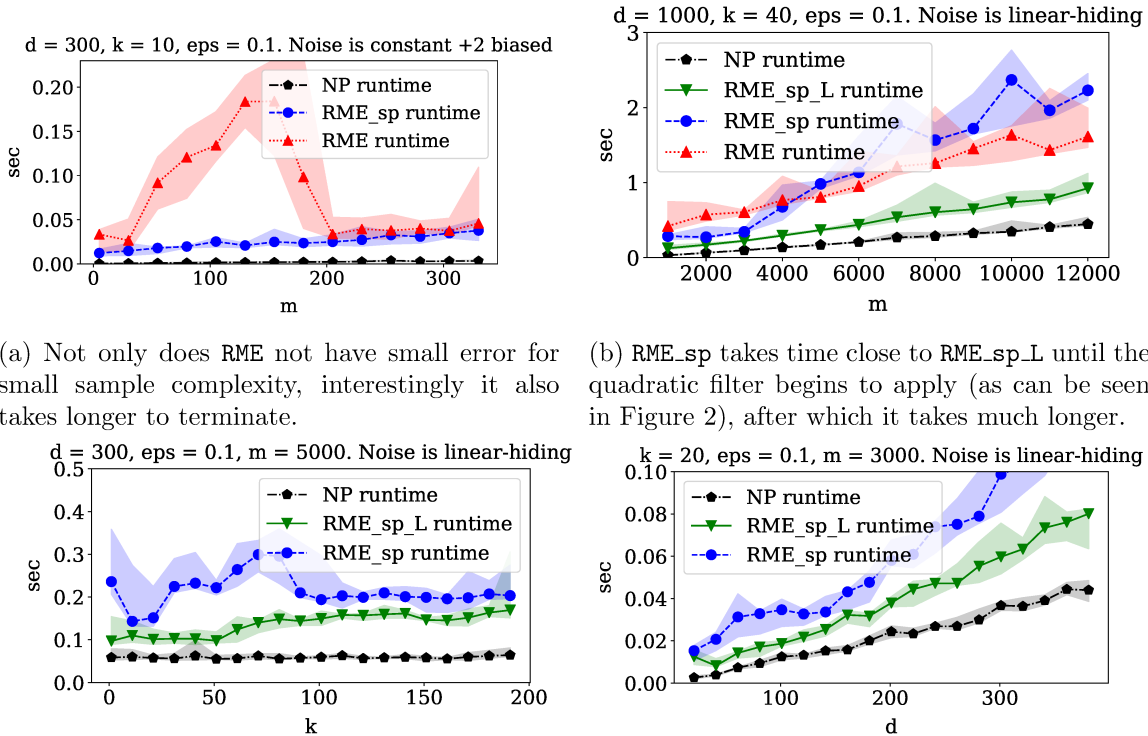
Figure 3: The flipping noise model demonstrates that the error can remain $\Omega(\varepsilon\sqrt{\log(1/\varepsilon)})$.



(a) The constant noise model is easy to remove and does not take many samples.

(b) The linear-hiding noise model is harder and requires more samples to get the same guarantee.

Figure 4: Sample complexity required to do well—in this case, 70% of errors being less than 1.2—depends on the noise model.



(a) Not only does RME not have small error for small sample complexity, interestingly it also takes longer to terminate.

(b) RME_sp takes time close to RME_sp_L until the quadratic filter begins to apply (as can be seen in Figure 2), after which it takes much longer.

(c) The runtimes for our sparse algorithms does not change very much as we increase k for the linear-hiding noise.

(d) The runtimes for our sparse algorithms increases with d linearly the case of linear-hiding noise.

Figure 5: Runtimes for robust mean estimation.

6.1 Robust Sparse Mean Estimation

The performance of robust estimation algorithms depend heavily on the noise model. The “hard” noise distributions for one algorithm may be easy for a different algorithm, if that one can identify and filter out the outliers. We therefore consider three different synthetic data distributions: two that demonstrate the $\varepsilon\sqrt{k}$ worst-case performance of other algorithms, and one that demonstrates the $\varepsilon\sqrt{\log(1/\varepsilon)}$ performance of our full algorithm.

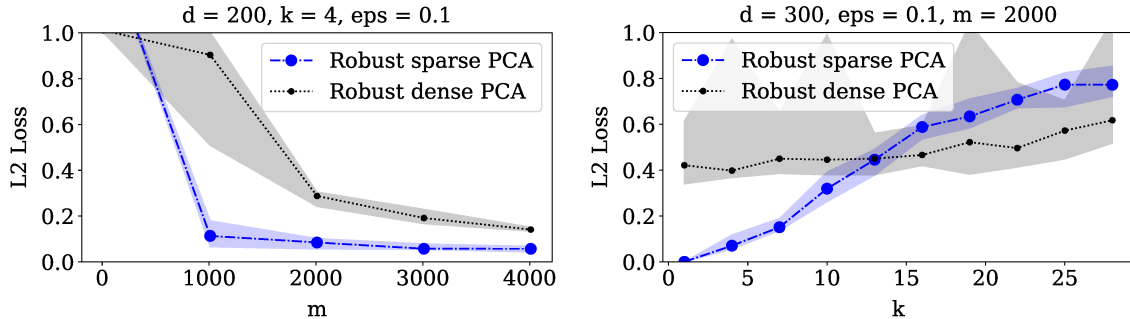
The algorithms we consider are **RME_sp**, our algorithm; **RME_sp_L**, a version of our algorithm with only the linear filter and not the quadratic one; **NP**, the “naive pruning” algorithm that drops samples with obviously-outlier coordinates, then outputs the empirical mean; **oracle**, which is told exactly which coordinates are inliers and outputs their empirical mean; **RME**, which applies the non-sparse robust mean estimation algorithm of [DKK⁺17]; and **RANSAC**, which computes the mean of a randomly chosen set of points, half the size of the entire set. One mean is preferred to another if it has more points in a ball of radius $\sqrt{d} + \sqrt{d}$ around it. For algorithms that have non-sparse outputs, we sparsify to the largest k coordinates before measuring the ℓ_2 distance to the true mean.

Our distributions are:

- **Constant-bias noise.** Noise that biases every coordinate consistently (e.g., if the outliers add 2 to every coordinate, or set every coordinate to $\mu_i + 1$) is difficult for naive algorithms (such as coordinate-wise median, **NP**, **RANSAC**) to deal with, but ideal for the linear filter. In Figure 1 we consider the noise that adds 2 to every coordinate.
- **Linear-hiding noise.** To demonstrate that the quadratic filter in our algorithm is necessary, we use the following data distribution. The inliers are drawn from $\mathcal{N}(0, I)$. The outliers are evenly split between two types: $\mathcal{N}(1_S, I)$ for some size- k set S , and $\mathcal{N}(0, 2I - I_S)$. The diagonal of the empirical covariance does not reveal S , so our linear filter fails to prune anything, leading to $\varepsilon\sqrt{k}$ error for **RME_sp_L**; the quadratic filter successfully removes all the outliers. This is shown in Figure 2.
- **Flipping noise.** For both those types of noise, with sufficiently many samples our final algorithm will prune out essentially all the outliers; there also exist noise models where $\Omega(\varepsilon\sqrt{\log(1/\varepsilon)})$ noise will remain at all times. In Figure 3 we demonstrate this for the noise model that picks a k -sparse direction v , and replaces the ε fraction of points furthest in the $-v$ direction with points in the $+v$ direction. In fact, for this noise even the **oracle** method also has $\Omega(\varepsilon\sqrt{\log(1/\varepsilon)})$ error from the missing points, but our algorithm has twice the error from the unfilterable added points.

Discussion. Matching our theoretical results, with sufficiently many samples the worst-case performance of **RME_sp** seems to be within a constant factor of the $O(\varepsilon\sqrt{\log(1/\varepsilon)})$ worst-case performance of **oracle**. This is not true for the naive algorithms **NP**, **RANSAC**, or the simplification **RME_sp_L** of our algorithm, which all have an $\varepsilon\sqrt{k}$ dependence. While our theoretical results show that $\tilde{O}(k^2)$ samples suffice, the empirical results given in Figure 4 are consistent with $\tilde{O}(k)$ being sufficient.

Our algorithm runs much faster than the ellipsoid based approach. For instance for $k = 10, d = 300, m = 50$ for the case of constant-biased noise our algorithm takes time 0.015 seconds to finish. In comparison the very first iteration for the SDP-based solution takes 10 seconds to solve with CVXOPT; the full ellipsoid-based algorithm, if implemented, would take many times that.



(a) The natural dense algorithm RDPCA requires more samples than the sparse algorithm to get error < 0.1

(b) For a fixed m , RSPCA performs better than RDPCA when $k < \sqrt{d}$ and then performs worse until coming close to RDPCA. Note that the variance of RSPCA is smaller than that of RDPCA.

Figure 6: Sample complexity of RSPCA is better than RDPCA for smaller sparsity.

6.2 Robust Sparse PCA

In Figure 6 we compare our robust sparse PCA algorithm RSPCA to a dense algorithm RDPCA for robust PCA. RDPCA looks at the empirical covariance matrix and then in the direction of maximum variance robustly estimates standard deviation. The algorithm then filters points using a modified version of the linear filter from [DKK⁺17] and hence requires a sample complexity of $\tilde{O}(d)$. For this algorithm, we only consider a single simple noise model. We draw outlier samples from $\mathcal{N}(0, I + uu^T)$ where u has disjoint support from the true vector v .

The sparse algorithm seems to perform better than the dense algorithm for k up to roughly $\sqrt{d}$; this is better than what we can prove, which is that it should be better up to at least $d^{1/4}$.

7 Conclusion

In this paper we have presented iterative filtering algorithms for two natural robust sparse estimation tasks: sparse mean estimation and sparse PCA. In both cases, our algorithm achieves a near-optimal $\tilde{O}(\varepsilon)$ error with a sample complexity primarily dependent on the sparsity k , and only logarithmically on the ambient dimension d . Our theoretical results are comparable to those of [BDLS17], but our algorithm only uses simple spectral techniques rather than the ellipsoid algorithm. This makes our algorithm quite feasible to implement. Our implementations perform essentially as expected: in sparse settings they require significantly fewer samples than dense robust estimation, and have accuracy avoiding the $\sqrt{k}$ dependence of other techniques like RANSAC.

References

- [BDLS17] S. Balakrishnan, S. S. Du, J. Li, and A. Singh. Computationally efficient robust sparse estimation in high dimensions. In *Proc. 30th Annual Conference on Learning Theory (COLT)*, pages 169–212, 2017.
- [BH95] Y. Benjamini and Y. Hochberg. Controlling the false discovery rate: a practical and powerful approach to multiple testing. *Journal of the Royal statistical society: series B (Methodological)*, 57(1):289–300, 1995.

- [BNJT10] M. Barreno, B. Nelson, A. D. Joseph, and J. D. Tygar. The security of machine learning. *Machine Learning*, 81(2):121–148, 2010.
- [BNL12] B. Biggio, B. Nelson, and P. Laskov. Poisoning attacks against support vector machines. In *Proceedings of the 29th International Conference on Machine Learning, ICML 2012*, 2012.
- [CDG18] Y. Cheng, I. Diakonikolas, and R. Ge. High-dimensional robust mean estimation in nearly-linear time. *CoRR*, abs/1811.09380, 2018. Conference version in SODA 2019, p. 2755–2771.
- [CDGW19] Y. Cheng, I. Diakonikolas, R. Ge, and D. P. Woodruff. Faster algorithms for high-dimensional robust covariance estimation. In *Conference on Learning Theory, COLT 2019*, pages 727–757, 2019.
- [CDKS18] Y. Cheng, I. Diakonikolas, D. M. Kane, and A. Stewart. Robust learning of fixed-structure Bayesian networks. In *Proc. 33rd Annual Conference on Neural Information Processing Systems (NIPS)*, 2018.
- [CSV17] M. Charikar, J. Steinhardt, and G. Valiant. Learning from untrusted data. In *Proc. 49th Annual ACM Symposium on Theory of Computing (STOC)*, pages 47–60, 2017.
- [DKK⁺16] I. Diakonikolas, G. Kamath, D. M. Kane, J. Li, A. Moitra, and A. Stewart. Robust estimators in high dimensions without the computational intractability. In *Proc. 57th IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 655–664, 2016. Journal version in SIAM Journal on Computing, 48(2), p. 742–864, 2019.
- [DKK⁺17] I. Diakonikolas, G. Kamath, D. M. Kane, J. Li, A. Moitra, and A. Stewart. Being robust (in high dimensions) can be practical. In *Proc. 34th International Conference on Machine Learning (ICML)*, pages 999–1008, 2017.
- [DKK⁺18a] I. Diakonikolas, G. Kamath, D. M. Kane, J. Li, A. Moitra, and A. Stewart. Robustly learning a Gaussian: Getting optimal error, efficiently. In *Proc. 29th Annual Symposium on Discrete Algorithms (SODA)*, pages 2683–2702, 2018.
- [DKK⁺18b] I. Diakonikolas, G. Kamath, D. M Kane, J. Li, J. Steinhardt, and A. Stewart. Sever: A robust meta-algorithm for stochastic optimization. *arXiv preprint arXiv:1803.02815*, 2018.
- [DKS17] I. Diakonikolas, D. M. Kane, and A. Stewart. Statistical query lower bounds for robust estimation of high-dimensional Gaussians and Gaussian mixtures. In *Proc. 58th IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 73–84, 2017.
- [DKS18a] I. Diakonikolas, D. M. Kane, and A. Stewart. Learning geometric concepts with nasty noise. In *Proc. 50th Annual ACM Symposium on Theory of Computing (STOC)*, pages 1061–1073, 2018.
- [DKS18b] I. Diakonikolas, D. M. Kane, and A. Stewart. List-decodable robust mean estimation and learning mixtures of spherical Gaussians. In *Proc. 50th Annual ACM Symposium on Theory of Computing (STOC)*, pages 1047–1060, 2018.

- [DKS19] I. Diakonikolas, W. Kong, and A. Stewart. Efficient algorithms and lower bounds for robust linear regression. In *Proc. 30th Annual Symposium on Discrete Algorithms (SODA)*, 2019.
- [HL18] S. B. Hopkins and J. Li. Mixture models, robustness, and sum of squares proofs. In *Proc. 50th Annual ACM Symposium on Theory of Computing (STOC)*, pages 1021–1034, 2018.
- [HR09] P. J. Huber and E. M. Ronchetti. *Robust statistics*. Wiley New York, 2009.
- [HRRS86] F. R. Hampel, E. M. Ronchetti, P. J. Rousseeuw, and W. A. Stahel. *Robust statistics. The approach based on influence functions*. Wiley New York, 1986.
- [HTW15] T. Hastie, R. Tibshirani, and M. Wainwright. *Statistical Learning with Sparsity: The Lasso and Generalizations*. Chapman & Hall/CRC, 2015.
- [Hub64] P. J. Huber. Robust estimation of a location parameter. *Ann. Math. Statist.*, 35(1):73–101, 03 1964.
- [Joh01] I. M. Johnstone. On the distribution of the largest eigenvalue in principal components analysis. *The Annals of Statistics*, 29(2):295–327, 2001.
- [Joh17] I. M. Johnstone. Gaussian estimation: Sequence and wavelet models. Available at http://statweb.stanford.edu/~imj/GE_08_09_17.pdf, 2017.
- [Kan18] D. M. Kane. Robust covariance estimation. *Talk given at TTIC Workshop on Computational Efficiency and High-Dimensional Robust Statistics*, 2018. Available at <http://www.iliadiakonikolas.org/tti-robust/Kane-Covariance.pdf>.
- [KKM18] A. Klivans, P. Kothari, and R. Meka. Efficient algorithms for outlier-robust regression. In *Proc. 31st Annual Conference on Learning Theory (COLT)*, pages 1420–1430, 2018.
- [KL93] M. J. Kearns and M. Li. Learning in the presence of malicious errors. *SIAM Journal on Computing*, 22(4):807–837, 1993.
- [KSS18] P. K. Kothari, J. Steinhardt, and D. Steurer. Robust moment estimation and improved clustering via sum of squares. In *Proc. 50th Annual ACM Symposium on Theory of Computing (STOC)*, pages 1035–1046, 2018.
- [LAT⁺08] J.Z. Li, D.M. Absher, H. Tang, A.M. Southwick, A.M. Casto, S. Ramachandran, H.M. Cann, G.S. Barsh, M. Feldman, L.L. Cavalli-Sforza, and R.M. Myers. World-wide human relationships inferred from genome-wide patterns of variation. *Science*, 319:1100–1104, 2008.
- [LLC19] L. Liu, T. Li, and C. Caramanis. High dimensional robust estimation of sparse models via trimmed hard thresholding. *CoRR*, abs/1901.08237, 2019.
- [LM00] B. Laurent and P. Massart. Adaptive estimation of a quadratic functional by model selection. *Ann. Statist.*, 28(5):1302–1338, 10 2000.
- [LRV16] K. A. Lai, A. B. Rao, and S. Vempala. Agnostic estimation of mean and covariance. In *Proc. 57th IEEE Symposium on Foundations of Computer Science (FOCS)*, 2016.

- [LSLC18a] L. Liu, Y. Shen, T. Li, and C. Caramanis. High dimensional robust sparse regression. *arXiv preprint arXiv:1805.11643*, 2018.
- [LSLC18b] L. Liu, Y. Shen, T. Li, and C. Caramanis. High dimensional robust sparse regression. *CoRR*, abs/1805.11643, 2018.
- [PLJD10] P. Paschou, J. Lewis, A. Javed, and P. Drineas. Ancestry informative markers for fine-scale individual assignment to worldwide populations. *Journal of Medical Genetics*, 47:835–847, 2010.
- [PSBR18] A. Prasad, A. S. Suggala, S. Balakrishnan, and P. Ravikumar. Robust estimation via robust gradient estimation. *arXiv preprint arXiv:1802.06485*, 2018.
- [RPW⁺02] N. Rosenberg, J. Pritchard, J. Weber, H. Cann, K. Kidd, L.A. Zhivotovsky, and M.W. Feldman. Genetic structure of human populations. *Science*, 298:2381–2385, 2002.
- [SCV18] J. Steinhardt, M. Charikar, and G. Valiant. Resilience: A criterion for learning in the presence of arbitrary outliers. In *Proc. 9th Innovations in Theoretical Computer Science Conference (ITCS)*, pages 45:1–45:21, 2018.
- [SKL17] J. Steinhardt, P. Wei Koh, and P. S. Liang. Certified defenses for data poisoning attacks. In *Advances in Neural Information Processing Systems 30*, pages 3520–3532, 2017.
- [Tsy08] A. B. Tsybakov. *Introduction to Nonparametric Estimation*. Springer Publishing Company, Incorporated, 2008.
- [Val85] L. Valiant. Learning disjunctions of conjunctions. In *Proceedings of the Ninth International Joint Conference on Artificial Intelligence*, pages 560–566, 1985.

A Proof of Lemma 4.3

Let G be a set of $N = \tilde{\Omega}(k^2 \log(d/\tau)/\varepsilon^2)$ i.i.d. samples drawn from $\mathcal{N}(\mu, I)$. We will show that each of Conditions (i)-(iv) hold with probability at least $1 - \tau/5$. The lemma then follows by a union bound.

Proof of (i): To establish (i), let $\mu^G := \mathbf{E}_{X \in_u G}[X]$ and note that the random variable $N\mu^G$ is distributed as $\mathcal{N}(N \cdot \mu, N \cdot I)$. Hence, μ^G has independent coordinates with $N\mu_i^G \sim \mathcal{N}(N \cdot \mu_i, N)$. By standard Gaussian tail bounds, we have that $\mathbf{Pr} \left[|N(\mu_i^G - \mu_i)| \geq T\sqrt{N} \right] \leq 2 \cdot \exp(-T^2/2)$. Setting $T/\sqrt{N} = \varepsilon/k$ gives that $\mathbf{Pr} \left[|\mu_i^G - \mu_i| \geq \varepsilon/k \right] \leq 2 \cdot \exp(-N\varepsilon^2/(2k^2)) \leq \tau/(10d)$. By a union bound over all $i \in [d]$, it follows that

$$\mathbf{Pr} \left[\exists i \in [d] : |\mu_i^G - \mu_i| \geq \varepsilon/k \right] \leq \tau/10.$$

This completes the proof of the first part of (i).

For the second part of (i), we will show that with probability at least $1 - \tau/10$ we have that for all $i, j \in [d]$, $|\mathbf{E}[(X_i - \mu_i)(X_j - \mu_j)] - \delta_{ij}| \leq \varepsilon/k$. We will need the following simple technical fact:

Fact A.1 (see, e.g., [LM00]). *Let Y_i be iid standard univariate Gaussians and $a_i \geq 0$, $i \in [m]$. If $Z = \sum_{i=1}^m a_i(Y_i^2 - 1)$, then for any $x \geq 0$ the following hold:*

$$\Pr [Z \geq 2\|a\|_2\sqrt{x} + 2\|a\|_\infty x] \leq \exp(-x), \quad (2)$$

and

$$\Pr [Z \leq -2\|a\|_2\sqrt{x}] \leq \exp(-x). \quad (3)$$

We start with the case that $i = j$. Note that the random variable $N \cdot \mathbf{E}_{X \in_u G} [(X_i - \mu_i)^2]$ follows a χ^2 -distribution with N degrees of freedom, i.e., it is the sum of N independent squared standard Gaussians. An application of Equation (2) implies that for all $x \geq 0$ we have:

$$\Pr \left[\left| N \cdot \mathbf{E}_{X \in_u G} [(X_i - \mu_i)^2] - N \right| \geq 2\sqrt{Nx} + 2x \right] \leq \exp(-x).$$

Setting $x := N\varepsilon^2/(9k^2)$, we get that

$$\Pr \left[\left| \mathbf{E}_{X \in_u G} [(X_i - \mu_i)^2] - 1 \right| \geq 2\varepsilon/(3k) + 2\varepsilon^2/(9k^2) \right] \leq \exp(-N\varepsilon^2/(9k^2)) \leq \tau/(10d^2).$$

We now analyze the case that $i \neq j$. Let $Y \sim \mathcal{N}(\mu, I)$. Note that for $i \neq j$, $i, j \in [d]$, we have that

$$(Y_i - \mu_i)(Y_j - \mu_j) = \left(\frac{(Y_i - \mu_i)}{2} + \frac{(Y_j - \mu_j)}{2} \right)^2 - \left(\frac{(Y_i - \mu_i)}{2} - \frac{(Y_j - \mu_j)}{2} \right)^2.$$

Since $\frac{(Y_i - \mu_i)}{2} + \frac{(Y_j - \mu_j)}{2}$ and $\frac{(Y_i - \mu_i)}{2} - \frac{(Y_j - \mu_j)}{2}$ are independent and distributed as $\mathcal{N}(0, 1/2)$, for $i \neq j$, the random variable $N \cdot \mathbf{E}_{X \in_u G} [(X_i - \mu_i)(X_j - \mu_j)]$ is distributed as the difference of a sum of N independent squared zero-mean Gaussians with variance $1/2$, and another such sum. This random variable has expectation 0 and once again, by Equation (2) applied with $a_i = 1/2$, it follows that

$$\Pr \left[\left| N \cdot \mathbf{E}_{X \in_u G} [(X_i - \mu_i)(X_j - \mu_j)] \right| \geq 2\sqrt{Nx} + x \right] \leq \exp(-x).$$

Setting $x := N\varepsilon^2/(9k^2)$ as above gives that

$$\Pr [|\mathbf{E}_{X \in_u G} [(X_i - \mu_i)(X_j - \mu_j)]| \geq \varepsilon/k] \leq \tau/(10d^2).$$

A union bound over all $i, j \in [d]$ implies that

$$\Pr [\exists i, j \in [d] : |\mathbf{E}_{X \in_u G} [(X_i - \mu_i)(X_j - \mu_j)] - \delta_{ij}| \geq \varepsilon/k] \leq \tau/10.$$

This gives the second part of (i). By a union bound, Condition (i) holds with probability at least $1 - \tau/5$.

Proof of (ii): For $Y \sim \mathcal{N}(\mu, I)$, the standard Gaussian tail bound gives $\Pr [|Y_i - \mu_i| \geq T] \leq 2\exp(-T^2/2)$. Setting $T = \sqrt{2\ln(10Nd/\tau)}$ implies that $\Pr [|Y_i - \mu_i| \geq T] \leq \tau/(10Nd)$. By a union bound, the desired upper bound holds for all $i \in [d]$ and all N samples with probability at least $1 - \tau/10$.

Proof of (iii): To establish (iii), we first prove that Conditions (iii)(a)-(c) hold for any fixed unit vector v and threshold $T \geq 4$ with sufficiently high probability, and then take a union bound over a net of $2k^2$ -sparse unit vectors and thresholds.

To avoid clutter in the notation, we will denote $\delta \stackrel{\text{def}}{=} \frac{\varepsilon^2}{\ln(k \ln(Nd/\tau))}$, so that the second term in the RHS of Condition (iii)(c) is equal to δ/T^2 .

We start by proving the following claim:

Claim A.2. *For any unit vector v in $\mathbb{R}^d$ and threshold $T \geq 4$ with probability at least $1 - \exp\left(-\Omega\left(\frac{N\delta}{\log(1/\delta)}\right)\right)$, we have that (a) $|\mathbf{E}_{X \in_u G}[v \cdot (X - \mu)]| \leq O(\varepsilon)$, (b) $|\mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2] - 1| \leq O(\varepsilon)$, and (c) $\Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \leq (5/2) \cdot \text{erfc}(T/\sqrt{2}) + \delta/(2T^2)$.*

Proof. To prove (a), note that for each fixed unit vector $v \in \mathbb{R}^d$, $N\mathbf{E}_{X \in_u G}[v \cdot (X - \mu)]$ is distributed as $\mathcal{N}(0, N)$. By standard Gaussian tail bounds, we have that

$$\Pr[|\mathbf{E}_{X \in_u G}[v \cdot (X - \mu)]| \geq \varepsilon] \leq 2 \cdot \exp(-N\varepsilon^2/2) \ll \exp(-\Omega(N\delta/\log(1/\delta))) ,$$

where the last inequality follows from the fact that $\delta \ll \varepsilon^2$.

To prove (b), note that for each fixed unit vector $v \in \mathbb{R}^d$ the random variable $N \cdot \mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2]$ follows a χ^2 -distribution with parameter N . By Equation (2), we get

$$\Pr\left[|N \cdot \mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2] - N| \geq 2\sqrt{Nx} + 2x\right] \leq \exp(-x) ,$$

for $x \geq 0$. Applying the above inequality for $x := N\varepsilon^2/9$, we get

$$\Pr\left[|\mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2] - 1| \geq 2\varepsilon/3 + (2/9)\varepsilon^2\right] \leq \exp(-N\varepsilon^2/9) .$$

To prove (c), we start by noting that, for any fixed unit vector v and $Y \sim \mathcal{N}(\mu, I)$, $v \cdot (Y - \mu)$ is a standard univariate Gaussian, and therefore $\Pr[|v \cdot (Y - \mu)| \geq T] = 2\text{erfc}(T/\sqrt{2})$. Let

$$Q(T) \stackrel{\text{def}}{=} (5/2)\text{erfc}(T/\sqrt{2}) + \delta/(2T^2) .$$

Observe that $N \cdot \Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T]$ is a sum of N independent Bernoulli random variables each with mean $2\text{erfc}(T/\sqrt{2})$. An application of the Chernoff bound and the fact that $Q(T) \geq (5/4)[2\text{erfc}(T/\sqrt{2})]$ gives that $\Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \geq Q(T)$ holds with probability at most $\exp\left(-\frac{NQ(T)}{60}\right)$.

We choose T' to satisfy $\text{erfc}(T') = \delta^2/(4T'^4)$, which implies that $T' = \Theta(\sqrt{\ln(1/\delta)})$. We break the analysis into two cases: $T \leq T'$ or $T > T'$.

If $T \leq T'$, then $Q(T) \geq Q(T') \geq \delta/(2T'^2) = \Omega(\frac{\delta}{\log(1/\delta)})$ and the above upper bound of $\exp\left(-\frac{NQ(T)}{60}\right)$ on the desired probability gives (c).

If $T > T'$, we have that $\text{erfc}(T/\sqrt{2}) \leq \delta^2/(4T^4)$. In this case, we require a more precise version of the Chernoff bound, which bounds from above the probability of the event $\Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \geq Q(T)$ by $\exp(-N \cdot D_{KL}(Q(T)||2\text{erfc}(T/\sqrt{2})))$, where $D_{KL}(p||q)$ denotes the KL-divergence between the Bernoulli random variables with probabilities p and q .

Let $p = \delta/(2T^2)$, $q = 2\text{erfc}(T/\sqrt{2})$, and note that $q \leq p^2$ or $p/q \geq q^{-1/2}$. Then, since $T > 4$, we see $p < 1/32$ and we can now bound from below the KL-divergence by $\delta/10$, as follows:

$$\begin{aligned} D_{KL}(Q(T)||q) &\geq D_{KL}(p||q) = p \ln(p/q) + (1-p) \ln((1-p)/(1-q)) \\ &\geq p \ln(p/q) - \ln(1-p) \geq p (\ln(p/q) - 1 - p) \\ &\geq \delta/(2T^2) \cdot (\ln(1/q^{1/2}) - 1 - p) \\ &\geq \delta/(2T^2) \cdot (T^2/8 - T^2/16) \geq \delta/16 , \end{aligned}$$

Thus, we have that $\Pr_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \geq Q(T)$ with probability at most $\exp(-\Omega(N\delta))$ in this case. This completes the proof of (c).

By a union bound, all events hold with probability at least $1 - \exp\left(-\Omega\left(\frac{N\delta}{\log(1/\delta)}\right)\right)$, completing the proof of Claim A.2. $\square$

We now define a cover over all $2k^2$ -sparse vectors as well as the possible values of T , and take a union bound over the product. To this end, let

$$R \stackrel{\text{def}}{=} \Theta\left(k \cdot \sqrt{\log(Nd/\tau)}\right)$$

be such that by (ii) we have $\|x - \mu\|_\infty \leq \frac{R}{\sqrt{2k}}$, for $x \in G$.

For each set $U \subseteq [d]$ of coordinates of size $2k^2$, let $\mathcal{C}_U$ be an ε/R^2 -cover, in ℓ_2 -norm, of the set of unit vectors supported on U (i.e., with all non-zero coordinates in U). Such a cover exists with $|\mathcal{C}_U| \leq O(R^2/\varepsilon)^{2k^2}$. Let $\mathcal{C}$ be the union of $\mathcal{C}_U$ over all sets U of coordinates of size $2k^2$. Then we have that

$$|\mathcal{C}| \leq \binom{d}{2k^2} \cdot O(R^2/\varepsilon)^{2k^2} \leq O(dR^2/\varepsilon)^{2k^2}.$$

Let $\mathcal{T} := \{\sqrt{i\varepsilon} \mid i \in \mathbb{Z}_+, 0 \leq i \leq R^2/\varepsilon^2\}$ be a net over thresholds T . Note that $|\mathcal{C}| \cdot |\mathcal{T}| \leq O(dR^2/\varepsilon)^{2k^2+2}$. By a union bound, Claim A.2 holds for all $v \in \mathcal{C}$ and $T \in \mathcal{T}$ except with probability at most

$$\begin{aligned} & O\left((dk^2/\varepsilon) \log(Nd/\tau)\right)^{2k^2+2} \cdot \exp(\Omega(-N\delta/\log(1/\delta))) \\ &= \exp\left(O(k^2 \log(dk \log(d/\tau)/\varepsilon)) - \Omega(N\varepsilon^2/\log^3(k/\varepsilon \log(d/\tau)))\right) \leq \tau/10, \end{aligned}$$

where we used the fact that $N = \tilde{\Omega}(k^2 \log(d/\tau)/\varepsilon^2)$. It remains to prove (iii) assuming this event holds.

By definition, for any k^2 -sparse unit vector $v \in \mathbb{R}^d$, there exists a $v' \in \mathcal{C}$ such that $\|v' - v\|_2 \leq \varepsilon/R^2$ and such that $v' - v$ is also k^2 -sparse. Thus, for any $x \in G$, we have

$$\begin{aligned} |v \cdot (x - \mu) - v' \cdot (x - \mu)| &\leq \|v' - v\|_1 \|x - \mu\|_\infty \\ &\leq \sqrt{2k} \|v' - v\|_2 R / \sqrt{2k} \leq \varepsilon/R. \end{aligned}$$

Therefore, for the mean we have that $|\mathbf{E}_{X \in_u G}[v \cdot X]| \leq |\mathbf{E}_{X \in_u G}[v' \cdot X]| + \frac{\varepsilon}{R} \leq O(\varepsilon)$. This gives Condition (iii)(a).

To establish Condition (iii)(b), we note that for any $x \in G$, we have

$$\begin{aligned} |(v \cdot (x - \mu))^2 - (v' \cdot (x - \mu))^2| &\leq O(|v \cdot (x - \mu) - v' \cdot (x - \mu)| (|v \cdot (x - \mu)| + |v' \cdot (x - \mu)|)) \\ &\leq O(\varepsilon/R) \cdot O(k \cdot R/k) \\ &\leq O(\varepsilon), \end{aligned}$$

where the second line uses the fact that $|v \cdot (x - \mu)| \leq \|v\|_1 \|x - \mu\|_\infty \leq k \|x - \mu\|_\infty \leq R$. Therefore, we have that $|\mathbf{E}_{X \in_u G}[(v \cdot (X - \mu))^2] - 1| \leq |\mathbf{E}_{X \in_u G}[(v' \cdot (X - \mu))^2] - 1| + O(\varepsilon) = O(\varepsilon)$. This gives Condition (iii)(b).

We now prove Condition (iii)(c). Consider the event $\{x \in G : |v \cdot (x - \mu)| \geq T\}$ for $T \geq \sqrt{2 \ln(1/\varepsilon) + 2}$. First note that this event is contained in the event $\{x \in G : |v' \cdot (x - \mu)| \geq T - \varepsilon/R\}$. Moreover, note that the event is empty, unless $T \leq \|v\|_1 \|x - \mu\|_\infty \leq R$, in which case $(T - \varepsilon/R)^2 \geq$

$T^2 - 2\varepsilon$. Therefore, by the definition of $\mathcal{T}$, there is a $T' \in \mathcal{T}$ with $T^2 - 2\varepsilon \leq T'^2 \leq (T - \varepsilon/R)^2$. Then we have

$$\begin{aligned}
\Pr_{X \in_u G} [|v \cdot (X - \mu)| \geq T] &\leq \Pr_{X \in_u G} [|v' \cdot (X - \mu)| \geq T - \varepsilon/R] \\
&\leq \Pr_{X \in_u G} [|v' \cdot (X - \mu)| \geq T'] \\
&\leq 5\text{erfc}(T')/2 + \delta/(2T'^2) \\
&\leq 5\text{erfc}\left(\sqrt{T^2 - 2\varepsilon}\right)/2 + \delta/(2(T^2 - 2\varepsilon)) \\
&= (5/(2\sqrt{2\pi})) \int_{\sqrt{T^2 - 2\varepsilon}}^{\infty} \exp(-x^2/2)dx + \delta/T^2 \\
&= (5/(2\sqrt{2\pi})) \int_T^{\infty} \exp(-(y^2 - 2\varepsilon)/2)(y/\sqrt{y^2 - 2\varepsilon})dy + \delta/T^2 \\
&= (5/(2\sqrt{2\pi})) \int_T^{\infty} \exp(\varepsilon) \exp(-y^2/2)(1 + O(\varepsilon))dy + \delta/T^2 \\
&\leq (5/(2\sqrt{2\pi})) \int_T^{\infty} (1 + O(\varepsilon)) \exp(-y^2/2)dy + \delta/T^2 \\
&\leq 3\text{erfc}(T/\sqrt{2}) + \delta/T^2,
\end{aligned}$$

where the third line follows from Claim A.2(c) applied for (v', T') . This completes the proof of Condition (iii)(c).

Proof of (iv): At a high-level, the proof is similar to that of Condition (iii) above. We start by proving that Conditions (iv)(a)-(b) hold for any fixed degree-2 polynomial and threshold T with sufficiently high probability, and then take a union bound over a net of k^2 -sparse $p(x)$ and T .

Note that a homogeneous degree-2 polynomial can be written as $p(x) = (x - \mu)^T A (x - \mu)$, for a symmetric matrix A , in which case we have $\mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)] = \text{Tr}(A)$ and $\mathbf{Var}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)] = \|A\|_F^2$.

We start by establishing the following claim:

Claim A.3. *Let $Y \sim \mathcal{N}(\mu, I)$. Given a homogeneous degree-2 polynomial $p(x)$ with $\mathbf{Var}[p(Y)] = 1$ and T with $4 \leq T \leq R \stackrel{\text{def}}{=} \Theta(k \cdot \sqrt{\log(Nd/\tau)})$, we have that: (a) $|\mathbf{E}_{X \in_u G}[p(X)] - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \leq O(\varepsilon)$, and (b) $\Pr_{X \in_u G}[|p(X) - \mathbf{E}[p(Y)]| \geq T] \leq 2 \exp(-T/4) + \varepsilon^2/(2T \ln^2 T)$, except with probability at most $\exp(-\Omega(N\varepsilon^2/\ln^2(R/\varepsilon)))$.*

Proof. By diagonalizing A , we can write $p(Y) = c + \sum_{i=1}^d a_i Z_i^2$, where the Z_i are independent and distributed as $\mathcal{N}(0, 1)$ and c, a_i are real coefficients with $\sum_i a_i^2 = \|A\|_F^2 = 2$. Note that $N\mathbf{E}[p(X)]$ is a sum of Nd independent squared Gaussians, each of which has variance at most $\|A\|_F^2 = 1$ and the ℓ_2 -norm of all their variances is $\sqrt{N}\|A\|_F = \sqrt{N}$. Equation (2) gives that $\Pr[N\mathbf{E}_{X \in_u G}[p(X)] - N\mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)] \geq 2\sqrt{Nx} + 2x] \leq \exp(-x)$, for $x \geq 0$. Taking $x := N\varepsilon^2$, we obtain that

$$\Pr[|\mathbf{E}_{X \in_u G}[p(X)] - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \geq 2\varepsilon + 2\varepsilon^2] \leq \exp(-N\varepsilon^2).$$

This shows (a).

We proceed to prove (b). By Equation (2) applied for a single sample, we have that $\Pr_{Y \sim \mathcal{N}(\mu, I)}[|p(Y) - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \geq 2\sqrt{x} + 2x] \leq \exp(-x)$ for $x \geq 0$. Taking $x := T$ for $T \geq 4$, we have $2\sqrt{T} \leq T$, and so

$$\Pr_{Y \sim \mathcal{N}(\mu, I)}[|p(Y) - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \geq T] \leq \exp(-T/4).$$

Note that $N\mathbf{Pr}_{X \in_u G}[|p(X) - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \geq T]$ is a sum of N independent Bernoulli random variables each with expectation at most $\exp(-T/4)$. Let

$$Q(T) \stackrel{\text{def}}{=} 2\exp(-T/4) + \varepsilon^2/(2T \ln^2 T).$$

Since $Q(T) \geq 2\mathbf{Pr}[|p(Y) - \mathbf{E}[p(Y)]| \geq T]$, by the multiplicative Chernoff bound we have that $\mathbf{Pr}_{X \in_u G}[|p(X) - \mathbf{E}_{Y \sim \mathcal{N}(\mu, I)}[p(Y)]| \geq T] \leq Q(T)$, except with probability at most $\exp(-NQ(T)/6)$.

Let T' be such that $\exp(-T'/6) = \varepsilon^2/(2T' \ln^2(T'))$. Note that $T' = \Theta(\log(1/\varepsilon))$. For $T \leq T'$, we have that $Q(T) \geq \varepsilon^2/(T' \ln^2 T')$, and so $\exp(-NQ(T)/6) \geq \exp(\Omega(-N\varepsilon/\log(1/\varepsilon)(\log \log(1/\varepsilon))^2))$.

For $T \geq T'$, note that $\varepsilon^2/(2T \ln^2(T)) \geq \exp(-T/6)$. Again we need to use a more explicit version of the Chernoff bound, which gives that $\mathbf{Pr}_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \geq Q(T)$ with probability at most $\exp(-ND_{KL}(Q(T) \parallel \exp(-T/4)))$.

When $T' \leq T \leq R$, $p = \varepsilon^2/(2T \ln^2 T)$, and $q = 2\exp(-T/4)$, we obtain

$$\begin{aligned} D_{KL}(Q(T) \parallel q) &\geq D_{KL}(p \parallel q) = p \ln(p/q) + (1-p) \ln((1-p)/(1-q)) \\ &\geq p \ln(p/q) - \ln(1-p) \\ &\geq p(\ln(p/q) - 1 - p) \\ &= (\varepsilon^2/(2T \ln^2 T))(\ln(p/\exp(-T/4)) - 1 - p) \\ &\geq (\varepsilon^2/(2T \ln^2 T))(\ln(\exp(T/6)) - 1 - p) \\ &\geq (\varepsilon^2/(2T \ln^2 T)) \cdot (T/7) \\ &\geq \varepsilon^2/(14 \ln^2 T) \geq \varepsilon^2/(14 \ln^2 R), \end{aligned}$$

Where we used the fact that $4 \leq T \leq R$, and the fact that if $p < 1/8$, then $\ln(1-p) < p(1-p)$. Thus, it follows that $\mathbf{Pr}_{X \in_u G}[|v \cdot (X - \mu)| \geq T] \geq Q(T)$ with probability at most $\exp(-\Omega(N\varepsilon^2/\ln^2 R))$ in this case. In either case, by a union bound, the claim holds except with probability $\exp(-\Omega(N\varepsilon^2/\ln^2(R/\varepsilon)))$. This completes the proof of (b) and of Claim A.3. $\square$

It remains to construct a cover of k^2 -sparse homogeneous degree-2 polynomials which have at most k^2 terms and $\mathbf{Var}[p(Y)] = 1$. Let U be the set of k^2 monomials $x_i x_j$, for $1 \leq i, j \leq d$. We construct a cover $\mathcal{C}_U$ of polynomials with terms only in the monomials in U as follows: We take a cover of unit vectors in $\mathbb{R}^{k^2}$ to within ℓ_2 -norm ε/R^2 and use the coordinates of each vector as the coefficients of the corresponding monomial. Thus, we can take $|\mathcal{C}_U| = 2^{O(k^2)}$. Then we let $\mathcal{C}$ be the union of $\mathcal{C}_U$ for all sets of k^2 monomials U . We therefore have that $|\mathcal{C}| \leq \binom{d}{k^2} \cdot O(R^2/\varepsilon)^{k^2} \leq O(dR^2/\varepsilon)^{k^2}$.

Let $\mathcal{T} = \{i\varepsilon : i \in \mathbb{Z}_+, 0 \leq i \leq R^2/\varepsilon^2\}$. Thus, $|\mathcal{C}| \cdot |\mathcal{T}| \leq O(dR^2/\varepsilon)^{k^2+1}$. By a union bound, Claim A.3 holds for all $p \in \mathcal{C}$ and $T \in \mathcal{T}$, except with probability at most

$$\begin{aligned} &O(dk^2 \log(Nd/\tau)/\varepsilon)^{k^2+1} \cdot \exp(-\Omega(N\varepsilon^2/\ln^2(R/\varepsilon))) \\ &= \exp(O(k^2 \log(dk \log(N/\tau)/\varepsilon)) - \Omega(N\varepsilon^2/\ln^2(k \log(Nd/\tau)/\varepsilon))) \leq \tau/10, \end{aligned}$$

where we used the fact that $N = \tilde{\Omega}(k^2 \log(d/\tau)/\varepsilon^2)$. It remains to prove (iv) assuming this event holds.

Consider any homogeneous degree-2 polynomial $p(x)$ with at most k^2 terms and $\mathbf{Var}[p(Y)] = 1$. By construction of the cover, there is a polynomial $p'(x) \in \mathcal{C}$ such that the total number of monomials appearing in either $p(x)$ or $p'(x)$ is at most k^2 , and if we write $p(x) = (x - \mu)^T A (x - \mu)$ and $p'(x) = (x - \mu)^T A' (x - \mu)$ for symmetric matrices A, A' , then $\|A - A'\|_F \leq \varepsilon/R^2$. Let U' be

the set of coordinates appearing in either $p(x)$ and $p'(x)$ and note that $|U'| \leq 2k^2$. For $x \in G$, we have

$$\begin{aligned} |p(x) - p'(x)| &= |(x - \mu)^T (A - A')(x - \mu)| \\ &= |(x - \mu)_{U'}^T (A - A')(x - \mu)_{U'}| \\ &\leq \|(x - \mu)_{U'}\|_\infty^2 \|A - A'\|_2 \\ &\leq R^2 \cdot \varepsilon / R^2 \leq \varepsilon. \end{aligned}$$

Therefore, we have that $|\mathbf{E}[p(X)] - \mathbf{E}[p(Y)]| \leq |\mathbf{E}[p'(X)] - \mathbf{E}[p'(Y)]| + 2\varepsilon \leq O(\varepsilon)$, since Claim A.3 holds for $p'(x)$. We have thus established Condition (iv)(a).

To show Condition (iv)(b), consider the event $\{x \in G : |p(x)| \geq T\}$ for $T > 5$. Let $T' \in \mathcal{T}$ be such that $T - 2\varepsilon \leq T' \leq T - \varepsilon$. Then, $|p(x)| \geq T$ implies that $|p'(x)| \geq T'$, and therefore

$$\begin{aligned} \Pr_{X \in_u G}[|p(X)| \geq T] &\leq \Pr_{X \in_u G}[|p'(X)| \geq T'] \\ &\leq 2 \exp(-T'/3) + \varepsilon^2 / (2T' \ln(T')^2) \\ &\leq 2 \exp(-(T - 2\varepsilon)/3) + \varepsilon^2 / (2(T - 2\varepsilon) \ln^2(T - 2\varepsilon)) \\ &\leq 3 \exp(-T/4) + \varepsilon^2 / (T \ln^2 T), \end{aligned}$$

where the second line follows from Claim A.3(b) for (p', T') . This completes the proof of Condition (iv)(b).

The proof of Lemma 4.3 is now complete.

B Proof of Lemma 5.3

Condition 1 follows from standard gaussian concentration bounds. To see that Condition 2 holds, we prove entrywise closeness of the matrices involved. We will use the following standard concentration inequality

Lemma B.1. *For any degree q , n -variate polynomial f*

$$\Pr_{A \sim \mathcal{N}(0, \Sigma)}[|f(A) - \mathbb{E}[f(A)]| > \tau] \lesssim e^{-\left(\frac{\tau^2}{R \cdot \mathbf{Var}_{A \sim \mathcal{N}(0, \Sigma)}[f(A)]}\right)^{1/q}}$$

where R is some universal constant.

Entries of $xx^T - (I + \rho vv^T)$ are degree 2 polynomials of Gaussians, and thus so is their mean over G . Hence, Lemma B.1 implies that for any $(i, j) \in [d] \times [d]$ that

$$\Pr[|\mathbf{E}_G[x_i x_j] - (\delta_{i,j} + \rho v_i v_j)| > \varepsilon/k] \lesssim \exp\left(-(N\varepsilon^2/Rk^2)^{1/2}\right).$$

Taking a union bound over i, j shows that with high probability $\mathbf{E}_G[xx^T]$ has each entry within ε/k of that of $\rho vv^T + I$, and this immediately implies Condition 2.

Condition 3 holds via a similar argument. Observe that it is sufficient to consider the case $\|w\|_2 = 1$ and sample enough points to satisfy

$$|\mathbf{E}_G[(x_i x_j - \delta_{i,j} - \rho v_i v_j)(x_k x_l - \delta_{k,l} - \rho v_k v_l)] - \mathbf{E}_{\mathcal{N}(0, I + \rho vv^T)}[(x_i x_j - \delta_{i,j} - \rho v_i v_j)(x_k x_l - \delta_{k,l} - \rho v_k v_l)]| \leq \frac{\varepsilon}{k^2}.$$

Then the spectral norm of the covariance matrix of $\gamma(x)$ for any $Q \times Q$ submatrix will also be bounded by ε . Note that this is just the probability that a degree-4 polynomial in Gaussian inputs

deviates too much from its mean, and thus by Lemma B.1 the probability that the above fails to hold for any (i, j, k, l) is at most

$$\exp\left(-(N\varepsilon^2/Rk^4)^{1/4}\right).$$

Taking a union bound over (i, j, k, l) yields our result.

Finally, for Condition 4, we note that (perhaps changing the constant C), it suffices to prove it for all $\binom{d^2}{k}$ possible Q 's and for all w in a cover of the unit ball of $\mathbb{R}^{k^2}$ (which will have size $2^{O(k^2)}$ and for T powers of 2 less than or equal to $k \log(dN)$ (since by Condition 1 $|\text{vec}(xx^T)_Q| = O(k \log(dN))$ for all $x \in G$). Once we have fixed Q, w and T , $\gamma_Q(x) \cdot w - \rho \text{vec}(vv^T)_Q \cdot w$ is a mean 0, variance $O(1)$, degree-2 polynomial so by Lemma B.1, the probability that it is more than CT is at most e^{-2T} . Then the probability that at least $\varepsilon N/(T^2 \log^2(T))$ of our x 's have this property is at most

$$\begin{aligned} \binom{N}{\varepsilon N/(T^2 \log^2(T))} \exp(-2T(\varepsilon N)/(T^2 \log^2(T))) &\leq \left(\frac{N e^{-2T}}{e \varepsilon N/(T^2 \log^2(T))} \right)^{(\varepsilon N)/(T^2 \log^2(T))} \\ &\leq \exp(-\Omega(T \varepsilon N/(T^2 \log^2(T)))) \\ &\leq \exp(-\Omega(\varepsilon N/T \log^2(T))) \\ &\leq \exp(-\Omega(k^3 \log(d/\varepsilon))). \end{aligned}$$

Taking a union bound over Q, w, T completes the proof.